

Preface to the Joint Proceedings of the ACM IUI 2026 Workshops

Antonela Tommasel¹, Patricia Kahr², Karthik Dinakar³ and Werner Geyer⁴

¹ISISTAN, CONICET-UNCPBA, JKU, Argentina, Austria

²University of Zurich, Switzerland

³Pienso Inc., United States of America

⁴Oracle, United States of America

Abstract

This volume collects the proceedings of five workshops co-located with the 31st Annual ACM Conference on Intelligent User Interfaces (IUI 2026), held in Paphos, Cyprus, on March 23–26, 2026. The workshops addressed diverse and timely topics at the intersection of artificial intelligence and human–computer interaction: agentic AI systems, uncertainty communication, intelligent health interfaces, human–AI experiences in extended reality, and human oversight of AI. The five workshops together accepted 42 papers out of 69 combined submissions.

Keywords

intelligent user interfaces, human-computer interaction, agentic AI, uncertainty communication, health user interfaces, extended reality, human oversight, workshop proceedings

1. Overview

This volume collects the proceedings of five workshops held at the 31st ACM Conference on Intelligent User Interfaces (IUI 2026) in Paphos, Cyprus. Together, the workshops reflect a shared concern at the heart of the IUI community: as AI systems grow more capable and more embedded in everyday life, how do we keep people meaningfully in control? **AgentCraft** examined the design and evaluation of agentic AI systems, asking how developers and end-users can maintain intelligible oversight of autonomous plans and actions. **CURE** investigated how interfaces can communicate uncertainty effectively, helping users calibrate trust in AI outputs that are often opaque or overconfident. **HealthIUI** explored the particular challenges that arise when intelligent interfaces enter health and wellness contexts, where personalization, privacy, and trust carry high stakes. **SHAPEXR** pushed into embodied territory, studying how extended reality environments can serve as both a testbed and an application domain for AI-powered, human-centric interaction. Finally, **AI CHAOS!** confronted the governance layer directly, asking what meaningful human oversight of high-risk AI systems actually requires – in design, methodology, and policy. Across 42 accepted papers and a rich program of keynotes, activities, and discussions, these workshops collectively advance the IUI community’s understanding of how to build AI systems that remain transparent, trustworthy, and accountable. Table 1 summarizes submission and acceptance statistics.

2. Workshop 1: AgentCraft 2026

Scope and Topics

Ambitious efforts are underway to build AI agents powered by large language models (LLMs) across a spectrum of domains, including programming, productivity, creative work, and enterprise automation. Although many frameworks have emerged to aid in the agent development process, major challenges remain regarding agent autonomy, reasoning processes, unpredictable behaviors, and consequential

Joint Proceedings of the ACM IUI Workshops 2026, March 23–26, 2026, Paphos, Cyprus

✉ antonela.tommasel@jku.at (A. Tommasel); patricia.kahr@uzh.ch (P. Kahr); karthik@pienso.com (K. Dinakar); werner.geyer@gmail.com (W. Geyer)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

Table 1

Submission and acceptance statistics for the five IUI 2026 workshops.

Workshop	Submitted	Accepted	Rate	PC Size
AgentCraft 2026	8	5	63%	17
CURE 2026	6	4	67%	30
HealthIUI 2026	30	14	47%	28
SHAPEXR 2026	11	7	64%	25
AI CHAOS! 2026	14	12	86%	9
Total	69	42	—	—

actions. Developers struggle to comprehend and debug their agent’s behavior in systematic ways. They also struggle to determine where in an agent’s workflow human input and oversight are needed. Intelligent interfaces that enable meaningful oversight of agentic plans, decisions, and actions are needed to foster transparency, build trust, and manage complexity. Further, evaluations of agentic AI systems require moving past automated machine learning benchmarks toward human-centric evaluations that demonstrate real-world usefulness, reliability, and safety. This workshop brought together researchers, developers, and designers to explore approaches to address these challenges of building human-centered and trustworthy agentic AI systems. We explored novel user interfaces for (1) mixed-initiative collaboration including developers and end-users, during agent development and deployment, (2) design patterns for debugging agent behaviors, (3) strategies for determining when developer control and oversight are needed, and (4) evaluation methods that ground agent performance in real-world impact. The workshop convened a diverse community from academia and industry in the form of a half-day workshop at IUI 2026.

Organizing Committee

- Karthik Dinakar, Pienso Inc., USA
- Justin D. Weisz, Snorkle Labs, USA
- Henry Lieberman, MIT CSAIL, USA
- Werner Geyer, Oracle, USA

Paper Selection

AgentCraft 2026 received 8 submissions. Each submission was reviewed by at least two program committee members. The program committee accepted 5 papers (acceptance rate: 63%). The accepted papers cover human-centered benchmarking of agentic app generation, multi-agent orchestration for instruction visualization, agentic AI for safety-critical multi-drone systems, a natural language specification language for framework-agnostic agent development, and latent pragmatic misalignment detection.

Program Committee

- Karthik Dinakar, Pienso Inc., USA (*chair*)
- Justin D. Weisz, Snorkle Labs, USA (*chair*)
- Henry Lieberman, MIT CSAIL, USA (*chair*)
- Werner Geyer, Oracle, USA (*chair*)
- Victor Dibia, Microsoft Research, USA
- Thariq Shihpar, Anthropic, USA
- Simret Gebreegziabher, University of Notre Dame, USA
- Sebastian Gehrmann, Bloomberg, USA
- Harang Ju, Johns Hopkins University, USA

- Steve Ross, IBM Research, USA
- Daniel Karl Weidele, IBM Research, USA
- Patrick Xian, Northeastern University, USA
- Zhiping Zhang, Northeastern University, USA
- Brian Cort, Pienso Inc., USA
- Arnav Nagpal, Pienso Inc., USA
- Meng-Hsin Wu, Carnegie Mellon University (now Pienso Inc.), USA
- Mathew Maradin, Pienso Inc., USA

3. Workshop 2: CURE 2026

Scope and Topics

The CURE2026 workshop provided a venue for research on how interactive and intelligent interfaces can be leveraged to ensure effective uncertainty communication. With the increased capabilities of black-box AI models, particularly large language models (LLMs), there is a growing urgency to understand exactly what these models can and cannot do. Effective uncertainty communication can help users calibrate their trust and reliance on the AI system. However, effectively conveying uncertainty, be it through conversation or by visualizing their confidence levels, remains a major challenge, as users often get confused or misinterpret uncertainty information due to poorly designed interfaces. CURE2026 brought together researchers from HCI, AI, design, and psychology to investigate how we can design, develop, and evaluate intelligent and interactive interfaces for uncertainty communication. The workshop included a mini-conference section aimed at inspiring discussion and investigating the role of intelligent and interactive interfaces in uncertainty communication. The workshop also featured keynotes by Professors Jessica Hullman and Katrien Verbert. Finally, a part of the workshop was dedicated to an interactive activity to discuss and develop guidelines for communicating uncertainty across different domains and tasks.

Organizing Committee

- Jasmina Gajcin, IBM Research, Ireland
- Jovan Jeromela, Trinity College Dublin, Ireland
- Joel Wester, Aalborg University, Denmark
- Sarah Schömbs, University of Melbourne, Australia
- Styliani Kleanthous, Open University of Cyprus, Cyprus
- Karthikeyan Natesan Ramamurthy, IBM Research, USA
- Hanna Hauptmann, Utrecht University, Netherlands
- Rifat Mehreen Amin, LMU Munich, Germany

Paper Selection

CURE 2026 received 6 submissions. Each paper was evaluated by at least two programme committee members using a double-blind review process, prioritising technical quality, theoretical soundness, novelty, and clarity. The programme committee accepted 4 papers (acceptance rate: 67%), all of which were full five-page submissions. The accepted papers address conceptual alignment in human-centered design teams using LLMs, visual cues in wrong-path conversational AI interactions, longitudinal philosophers' judgments of generative AI interfaces, and uncertainty perception by humans and LLMs in multi-agent social simulations.

Program Committee

- Abhraneel Sarma, Northwestern University, USA
- Alex C. Williams, AWS, USA
- Alok Debnath, Trinity College Dublin, Ireland
- Danielle Albers Szafir, University of North Carolina – Chapel Hill, USA
- Eoin Delaney, Trinity College Dublin, Ireland
- Eelco Herder, Utrecht University, Netherlands
- Francesca Naretto, University of Pisa, Italy
- Hariharan Subramonyam, Stanford University, USA
- Hyo Jin Do, IBM Research, USA
- Ivana Dusparic, Trinity College Dublin, Ireland
- Jacy Reese Anthis, University of Chicago, USA
- Jan Leusmann, LMU Munich, Germany
- Jeroen Ooge, Utrecht University, Netherlands
- Justin Edwards, University of Oulu, Finland
- Leon Fröhling, GESIS Leibniz Institute for Social Sciences, Germany
- Maxwell Szymanski, KU Leuven, Belgium
- Michael Hedderich, LMU Munich, Germany
- Owen Conlan, Trinity College Dublin, Ireland
- Prakash Jamakatel, University of the Bundeswehr Munich, Germany
- Rune Møberg Jacobsen, Aalborg University, Denmark
- Samuel Rhys Cox, Aalborg University, Denmark
- Selene Baez Santamaria, University of Zurich, Switzerland
- Siddharth Swaroop, Harvard University, USA
- Simret Araya Gebreegziabher, University of Notre Dame, USA
- Siya Kunde, IBM Research, USA
- Srishti Yadav, University of Copenhagen, Denmark
- Sunnie S.Y. Kim, Apple, USA
- Timo Speith, University of Bayreuth, Germany
- Yan Zhang, University of Melbourne, Australia
- Yue Fu, University of Washington, USA

4. Workshop 3: HealthIUI 2026

Scope and Topics

As Artificial Intelligence (AI) continues to transform health and care, the integration of Intelligent User Interfaces (IUI) in health and wellness applications presents both significant opportunities and challenges. This workshop aimed to bring together researchers and practitioners from HCI, AI, healthcare, and related fields to explore how IUIs can impact long-term user engagement, personalization, and trust in health-oriented interactive systems. We focused on interdisciplinary approaches to design systems that are technically advanced but also responsive to user needs, demands of context of use, values and ethical requirements, and privacy. Through presentations, discussions, and collaborative sessions, participants identified key challenges, shared emerging solutions, and outlined pathways for responsible and impactful innovation in health IUI.

Organizing Committee

- Peter Brusilovsky, University of Pittsburgh, USA
- Behnam Rahdari, Stanford University, USA
- Shriti Raj, Stanford University, USA
- Helma Torkamaan, TU Delft, Netherlands

Paper Selection

HealthUI 2026 received 30 submissions. All submissions underwent a peer-review process with each submission receiving at least three reviews. The programme committee accepted 14 papers (acceptance rate: 47%). The accepted papers span personalized health behavior support, clinical decision systems, conversational and generative AI for health, digital biomarkers for mental health, elder care, physiotherapy support, AI-based wearable assistance, and intelligent system design for chronic conditions.

Program Committee

- Bart Knijnenburg, Clemson University, USA
- Federica Cena, University of Torino, Italy
- Maxwell Szymanski, KU Leuven, Belgium
- Alain D. Starke, University of Amsterdam, Netherlands
- Hanna Hauptmann, Utrecht University, Netherlands
- Shatha Degachi, TU Delft, Netherlands
- Min Lee, Singapore Management University, Singapore
- Bereket Yilma, University of Luxembourg, Luxembourg
- Werner Geyer, Oracle, USA
- Marko Tkalcic, University of Primorska, Slovenia
- Adir Solomon, University of Haifa, Israel
- Niels van Berkel, Aalborg University, Denmark
- Clara-Maria Barth, University of Zurich, Switzerland
- Tobias Schreck, Graz University of Technology, Austria
- Stefan Hillmann, Technische Universität Berlin, Germany
- Stefan Lengauer, Graz University of Technology, Austria
- Francisco Nunes, Fraunhofer Portugal AICOS, Portugal
- Tariq Osman Andersen, University of Copenhagen, Denmark
- Olga C. Santos, UNED, Spain
- Martin Wiesner, Heilbronn University, Germany
- Ludovico Boratto, University of Cagliari, Italy
- Noemi Mauro, University of Torino, Italy
- Ine Coppens, Ghent University, Belgium
- Xuhai Xu, Columbia University, USA
- Hyunggu Jung, Seoul National University, South Korea
- Soohwan Lee, UNIST, South Korea
- Edward Choi, KAIST, South Korea
- Hubert Dariusz Zajac, University of Copenhagen, Denmark

5. Workshop 4: SHAPEXR 2026

Scope and Topics

This workshop explored the relationship between eXtended Reality (XR), Artificial Intelligence, and Human Factors, focusing on the fundamental principles of Human-Computer Interaction within the Intelligent User Interfaces topic. With the increasing integration of large language models, conversational agents, and autonomous AI services into everyday interactions, XR offers embodied environments that blend language, gesture, gaze, touch, and spatial context. The workshop examined how AI can enable real-time personalization of immersive interfaces, the design of intuitive multimodal interactions with

AI agents, and methods for assessing and optimizing human factors such as cognitive load, emotional engagement, and trust. Additional topics included accessibility by design and the social dynamics of AI-enhanced multi-user XR spaces. Positioning XR as both an application domain and a gateway to AI, the workshop fostered research on intelligent, user-centric, and inclusive immersive interfaces.

Organizing Committee

- Giuseppe Caggianese, National Research Council of Italy, Institute for High-Performance Computing and Networking (CNR-ICAR), Naples, Italy
- Marta Mondellini, National Research Council of Italy, Institute of Intelligent Industrial Systems and Technologies for Advanced Manufacturing (CNR-STIIMA), Lecco, Italy
- Nicola Capece, University of Basilicata, Italy
- Mario Covarrubias, Politecnico di Milano, Italy
- Gilda Manfredi, University of Basilicata, Italy

Paper Selection

SHAPEXR 2026 received 11 submissions. Each submission received two to three reviews. The programme committee accepted 7 papers (5 full papers and 2 short papers; acceptance rate: 64%), representing contributions from Germany, Luxembourg, Italy, and Israel. The accepted papers cover cognitive failure detection in VR via eye-tracking, LLM-powered humanoid robotics, multimodal adaptive XR for de-escalation training, AI wearable assistance for clinical ward visits, hand-based inclusive region-of-interest selection, cultural heritage in adaptive virtual reality, and occlusion-robust multimodal emotion recognition.

Program Committee

- Agnese Augello, CNR-ICAR, Italy
- Sabrina Bartolotta, Università Cattolica del Sacro Cuore, Italy
- Sara Buonocore, University of Naples Federico II, Italy
- Luigi Casoria, CNR-ICAR, Italy
- Valerio Colonnese, University of Basilicata, Italy
- Sabatina Criscuolo, CNR-STIIMA, Italy
- Valerio De Luca, Pegaso University, Italy
- Giovanni D'Errico, University of Salento, Italy
- Luigi Duraccio, University of Naples Federico II, Italy
- Ugo Erra, University of Basilicata, Italy
- Antonio Esposito, University of Naples Federico II, Italy
- Silvia Franchini, CNR-ICAR, Italy
- Luigi Gallo, Pegaso University, Italy
- Ludovica Gargiulo, CNR-STIIMA, Italy
- Boriana Koleva, University of Nottingham, United Kingdom
- Katia Lupinetti, CNR-IMATI, Italy
- Federico Manuri, Politecnico di Torino, Italy
- Alfonso Matropietro, CNR-STIIMA, Italy
- Pietro Neroni, CNR-ICAR, Italy
- Marta Pizzolante, Università Cattolica del Sacro Cuore, Italy
- Anran Qi, INRIA, France
- Luca Sabatucci, CNR-ICAR, Italy
- Lorenzo Stacchio, University of Macerata, Italy
- Veronica Sundstedt, Blekinge Institute of Technology, Sweden
- Walter Terkaj, CNR-STIIMA, Italy

6. Workshop 5: AI CHAOS! 2026

Scope and Topics

As AI systems are increasingly adopted in high-stakes domains such as healthcare, autonomous driving, and criminal justice, their failures may threaten human safety and rights. Human oversight of AI systems is therefore critically important, as a potential safeguard to prevent harmful consequences in high-risk AI applications. Although regulations like the European AI Act mandate human oversight for high-risk AI, we lack methodologies and conceptual clarity to implement it effectively. Independent of policy and regulation, poorly designed oversight can create dangerous illusions of safety while obscuring accountability. This interdisciplinary workshop brought together researchers from various disciplines, including AI, HCI, psychology, law, and policy, to address this critical gap. We explored the following questions — How can we design AI systems that enable meaningful human oversight? What methods effectively communicate system states and risks to human overseers? How do we ensure scalable and effective interventions? Through papers, talks, and interactive group discussions, participants identified oversight challenges, examined stakeholder roles, discussed supporting tools, methods, regulatory frameworks, and established a collaborative research agenda. Our central goal was to further a roadmap that enables effective human oversight for the responsible deployment of AI in society.

Organizing Committee

- Tim Schrills, University of Lübeck, Germany
- Patricia Kahr, University of Zurich, Switzerland
- Markus Langer, University of Freiburg, Germany
- Harmanpreet Kaur, University of Minnesota, USA
- Ujwal Gadiraju, Delft University of Technology, Netherlands

Paper Selection

AI CHAOS! 2026 received 14 submissions. Each submission was reviewed by at least two programme committee members. The programme committee accepted 12 papers (acceptance rate: 86%). The accepted papers address moral responsibility under intelligent support, oversight-by-design for accessible generative interfaces, benchmarks for superintelligence, distributed oversight via community notes, gaze-based reinforcement learning for personalized highlighting, neurodivergent learning states, data degradation in recursive generative AI training, perceived trustworthiness, compliance assessment under the Digital Services Act, disagreement in LLM-as-a-judge systems, the human oversight fallacy, and risk-aware AI design for employee incentive planning.

Program Committee

- Agathe Balayn, Microsoft Research, USA
- Margaret Burnett, Oregon State University, USA
- Wen Duan, Clemson University, USA
- Anna Maria Feit, Saarland University, Germany
- Hanna Hauptmann, Utrecht University, Netherlands
- Tim Miller, University of Queensland, Australia
- Niels van Berkel, Aalborg University, Denmark
- Martijn Willemsen, Eindhoven University of Technology, Netherlands
- Hanwei Zhang, Saarland University, Germany

Acknowledgments

We thank the authors, programme committee members, and session chairs of all five workshops for their contributions to this volume. We are grateful to the IUI 2026 organizing committee for hosting these workshops in Paphos, Cyprus, and to CEUR-WS.org for publishing these proceedings under a Creative Commons CC BY 4.0 license.

Declaration on Generative AI

The authors have used Generative AI tools for improving grammar and spelling during the preparation of this preface.