



LREC 2026

Identity-Aware AI (IAAI) @ LREC 2026

Workshop Proceedings

Editors

Pranav A

Valerio Basile

Neele Falk

David Jurgens

Gabriella Lapesa

Anne Lauscher

Soda Marem Lo

16 May 2026

©ELRA Language Resources Association (ELRA), 2026

These proceedings are licensed under a Creative Commons Attribution-NonCommercial 4.0 International License (CC BY-NC 4.0)

ISBN 978-2-493814-88-3

Preface

Welcome to the Workshop on Identity-Aware AI (IAAI), held on 16 May 2026 as part of the 15th Language Resources and Evaluation Conference (LREC 2026) in Palma de Mallorca, Spain.

This workshop brings together researchers working at the intersection of natural language processing, AI ethics, and social science to examine how identity is represented, misrepresented, and operationalized in language technologies. A central concern of the workshop is the impact of these systems on marginalized communities, and the development of more equitable and identity-aware approaches to NLP.

We received submissions addressing a broad range of topics, including bias detection in clinical text, fairness in classification, cross-cultural value misalignment, social bias benchmarking, identity representation in NLP, honorifics and linguistic identity, and community-based auditing of AI harms. From these submissions, we selected 6 archival papers and 3 non-archival contributions for presentation. The non-archival papers appear in the program but not in these proceedings.

The workshop featured two keynote talks. The paired keynote and panel discussion was delivered by Rossana Damiano (University of Turin) and Samuel Goree (Stonehill College), followed by an invited keynote by Debora Nozza (Bocconi University). We are grateful to all three speakers for their contributions to the day's discussions.

We thank all authors for their submissions, all reviewers for their careful and constructive feedback, and the LREC 2026 organizing committee and ELRA for their support in making this workshop possible.

The Organizers
Identity-Aware AI (IAAI) @ LREC 2026

Organizing Committee

Organizers:

- Pranav A, University of Hamburg, Germany
- Valerio Basile, University of Turin, Italy
- Neele Falk, University of Stuttgart, Germany
- David Jurgens, University of Michigan, USA
- Gabriella Lapesa, GESIS, Leibniz Institute for the Social Sciences & Heinrich-Heine University of Düsseldorf, Germany
- Anne Lauscher, University of Hamburg, Germany
- Soda Marem Lo, University of Turin, Italy

Program Committee:

Gavin Abercrombie, Valerio Basile, Patrizio Bellan, Shree Harsha Bokkahalli Satish, Minh Duc Bui, Filipa Calado, Hongyu Chen, Luis Chiruzzo, Samuele D'Avenia, Lia Draetta, Esra Dönmez, Agnieszka Falenska, Neele Falk, Darja Fišer, Maryam Fooladi, Saba Ghanbari Haez, Fatemeh Ghavidel, Sara Goggi, Henning Hoffmann, Carolin Holtermann, David Jurgens, Os Keyes, Gabriella Lapesa, Anne Lauscher, Beckett LeClair, Soda Marem Lo, Anne-Marie Lutgen, Marta Marchiori Manerba, Alex Markham, Marcin Moskalewicz, Arianna Muti, Renáta Németh, Amir H. Payberah, Giorgio Pedrazzi, Alistair Plum, A Pranav, Ella Rabinovich, Yujie Ren, Julia Romberg, Laurel Stvan, Menno van Zaanen, Karina Vida, Urs Zaberer

Invited Speakers:

- Rossana Damiano, University of Turin, Italy
- Samuel Goree, Stonehill College, USA
- Debora Nozza, Bocconi University, Italy

Table of Contents

<i>The Point of View of a Sentiment: Towards Clinician Bias Detection in Psychiatric Notes</i> Alissa A. Valentine, Lauren Lepow, Lili Chan, Alexander Charney and Isotta Landi	1
<i>Speak Your Mind: The Speech Continuation Task as a Probe of Voice-Based Model Bias</i> Shree Harsha Bokkahalli Satish, Harm Lameris, Olivier Perrotin, Gustav Eje Henter and Eva Szekely	12
<i>Investigating the Automatic Translation of Korean Honorifics</i> Luis Cihlar, Minh Duc Bui, Kyung eun Park, Manuel Mager, Walter Bisang and Katharina von der Wense	20
<i>Balancing the Scales: Reinforcement Learning for Fair Classification</i> Leon Eshuijs, Shihan Wang and Antske Fokkens	32
<i>Evaluating LLMs for Detecting Demographic-Targeted Social Bias: A Comprehensive Bench- mark Study</i> Ayan Majumdar, Feihao Chen, Jinghui Li and Xiaozhen Wang	47
<i>Queering the Audits: Community-Based Auditing of AI Harms to Queer Communities</i> Organizers of Queer In AI, A Pranav, Alissa A. Valentine, Alex Markham, Beckett LeClair, Tereza Blazkova, Ekaterina Kornilitsina, Sofie H. Bruun, Gerasimos Spanakis and Anne Lauscher	66

Workshop Program

Friday, May 16, 2026

9:00–9:10 *Opening Remarks*
Organizers

9:10–10:15 *Paired Keynote & Panel Discussion*
Rossana Damiano and Samuel Goree

10:15–11:15 Session: Extended Poster Session and Coffee Break

The Point of View of a Sentiment: Towards Clinician Bias Detection in Psychiatric Notes

Alissa A. Valentine, Lauren Lepow, Lili Chan, Alexander Charney and Isotta Landi

Speak Your Mind: The Speech Continuation Task as a Probe of Voice-Based Model Bias

Shree Harsha Bokkahalli Satish, Harm Lameris, Olivier Perrotin, Gustav Eje Henter and Eva Szekely

Investigating the Automatic Translation of Korean Honorifics

Luis Cihlar, Minh Duc Bui, Kyung eun Park, Manuel Mager, Walter Bisang and Katharina von der Wense

Balancing the Scales: Reinforcement Learning for Fair Classification

Leon Eshuijs, Shihan Wang and Antske Fokkens

Evaluating LLMs for Detecting Demographic-Targeted Social Bias: A Comprehensive Benchmark Study

Ayan Majumdar, Feihao Chen, Jinghui Li and Xiaozhen Wang

Queering the Audits: Community-Based Auditing of AI Harms to Queer Communities

Organizers of Queer In AI, A Pranav, Alissa A. Valentine, Alex Markham, Beckett LeClair, Tereza Blazkova, Ekaterina Kornilitsina, Sofie H. Bruun, Gerasimos Spanakis and Anne Lauscher

11:15–12:00 *Invited Keynote*
Debora Nozza

12:00–13:00 *Virtual Presentations & Conclusion*
Organizers