

Sociolinguistics and AI
International conference
19–21 August 2026



UNIVERSITY OF
COPENHAGEN

Book of abstracts

Contents

Welcome	3
A note on the programme and the organisation of abstracts	4
For presenters – timing your talk, testing your connection	4
What's a dialogic poster session?	4
Publication forum for early career scholars	5
In lieu of a conference dinner	5
Venue	6
About the conference	6
Committees	7
Keynote presentations (abstracts)	10
Abstracts – ordered according to thematic sessions	17
List of presenters and contributions	119

Welcome

It is our great pleasure to welcome you to *Sociolinguistics and AI 2026*! We are very excited to see such a broad cross-section of dynamic research featured in the conference programme, with both old and new friends, acquaintances and sparring partners represented in equal measure. We are confident that the wealth of contributions will make for a very stimulating and enjoyable three days. We would like to take this opportunity to offer our appreciation for the work all of you have put into making this conference happen, either by submitting abstracts, serving as reviewers on the scientific committee, acting as chairs, or simply by signing up and contributing to the conversations that will unfold at the conference. We hope all your efforts will be amply rewarded in the discussions, feedback and future research opportunities that may emerge from the conference.

Best wishes,

Janus Mortensen, Sam Goodchild, Kasper Engholm Jelby, Jens Christian Borup Green Jensen, Sanne Larsen, Rafael Lomeu Gomes and Sofie E. Abrahamsen Søndergaard – the AI-UNI team



A note on the programme and the organisation of abstracts

Before you begin browsing the abstracts, we would like to inform you of our scheduling policy. In putting the programme together, we have aimed to create thematically coherent sessions of 3–4 contributions and issued each of them with an indicative title. While you are obviously free to move as you please, the sessions are ideally meant to be attended in full. It is our hope that this way of organising the programme will benefit the discussion and exchanges at the conference. The book of abstracts is organised according to thematic sessions, so once you've identified a session of interest in the conference programme, it should be easy to navigate to the abstracts of the session using the digital bookmarks included in this pdf (or simply flicking through it).

For presenters – timing your talk, testing your connection

Thematic sessions follow a standard format for presentations: up to 20 minutes for presentation, followed by 5 minutes for discussion and 5 minutes for getting ready for the next speaker. Chair-persons are responsible for time keeping, and presenters are requested to follow their instructions. All presenters must bring their own device for their presentation. To ensure that all sessions run smoothly, we kindly ask you to find the time to check the connection between your computer and the projector in the room prior to the beginning of your session. The room is equipped with an HDMI-cable and an HDMI-to-USB-C adapter. If you experience any issues with the connection, please let a conference organiser know well in advance of your session so we can help resolve it. If you experience delays during your presentation due to technical difficulties, we will unfortunately not be able to extend your presentation time. The chair of each session will be asked to keep the time so that the programme will not be delayed.

What's a dialogic poster session?

Presenters are responsible for printing and bringing their own posters (Size: A0). Posters will be on display throughout the three days of the conference. Delegates will be invited to read the posters and write comments and/or questions, either by using the designated padlet or the post-its provided. During the poster session (Thursday, 14:40-16:00), each poster presenter will have a dedicated 15-min slot. The slot will begin with

the presenter making a short (3–5 minutes) presentation of their poster covering basic information such as the context of the study, research questions/hypotheses, methods and data collection, main findings/lines of argumentation, implications/limitations/future steps. Afterwards, the presenter will choose some of the comments and questions that delegates have posted and offer a response. This will be followed by the facilitator opening for Q&A so that delegates can participate in the conversation. After the final poster presentation, an open discussion with all the poster presenters and delegates in the session will be facilitated (time permitting).

Publication forum for early career scholars

The publication forum aims to demystify the publication process. The forum will be a combination of a panel discussion and a networking event where early career researchers (and others interested) can engage in informal discussions and have questions about the publication process answered over lunch. The panel consists of four experienced editors, including Sam Goodchild (subject editor, *Nordic Journal of African Studies*), Britta Schneider (guest editor, *Language Sciences* special issue *Sociolinguistics and AI – Approaches, themes and insights*), Hartmut Haberland (co-founding editor, *Journal of Pragmatics*, co-founding editor and co-editor *Pragmatics and Society*), and Sara Ganassin (co-editor in chief, *Language and Intercultural Communication*). We encourage you to submit questions in advance using this padlet: https://padlet.com/ucph/publication_forum.

In lieu of a conference dinner

To keep participation costs for the conference low, we have decided against a formal conference dinner. However, the organising team is planning on going to *Broens Street Food* (<https://broensstreetfood.dk/en/>) on Thursday 20 August after the closing of the second day, and you are more than welcome to join us from 7 p.m. for an informal drink and a bite of food. *Broens Street Food* is located on Strandgade 95, straight across the bridge from Nyhavn, and offers a variety of international cuisine. See the selection here: <https://tinyurl.com/slx-ai-broens> (opens food stall menus as a pdf). NB: This is an outdoors venue, so remember to bring warm clothes. Also, please keep in mind that the market tends to get busy, and we cannot guarantee seating. We have *not* booked/reserved seating.

Venue

The conference takes place at the University of Copenhagen, South Campus, Building 23. Registration is located on the ground floor of building 23. Wayfinding instructions and a map of the campus area are available on the conference website, <http://tiny.cc/slx-ai-2026>. Once you're on campus and close to Building 23, you'll be able to follow the signs we put out for the occasion.

About the conference

The conference is organised by the AI-UNI group, as part of the research project *AI and the University: Towards a sociolinguistics of literacy and voice in the age of generative language technology*. The AI-UNI project is hosted by the Centre for Internationalisation and Parallel Language Use (CIP) at the University of Copenhagen and funded by the Carlsberg Foundation (2024-2028). The Carlsberg Foundation also supports the conference via a dedicated conference grant.

Read more about CIP at <https://cip.ku.dk/english/>

The screenshot shows the homepage of the Centre for Internationalisation and Parallel Language Use (CIP) at the University of Copenhagen. The header includes the university's name and navigation links: Education at UCPH, Research, News, Innovation, Employment, and About UCPH. A search bar is located on the right. The main content area features a large banner with the text 'About CIP' and a description: 'CIP is a teaching and research centre focused on parallel language use and the balance between Danish and English.' Below the banner are six red buttons with white text and right-pointing arrows: 'LANGUAGE COURSES FOR STAFF', 'LANGUAGE COURSES FOR STUDENTS', 'CIP-SERIES: STUDIES IN PARALLEL LANGUAGE USE', 'TOEPAS - ENGLISH LANGUAGE CERTIFICATION FOR LECTURERS', 'THE LANGUAGE STRATEGY', and 'RESEARCH AT CIP'. A left sidebar contains a list of links: 'About CIP', 'Staff', 'Research', 'Language courses', 'Implementation of the UCPH language policy', 'Projects and collaborations', 'TOEPAS - English language certification', 'Newsletter', and 'Contact'. At the bottom of the sidebar is a link to 'CIP's Resource Site'.

UNIVERSITY OF
COPENHAGEN



CARLSBERG
FOUNDATION

Committees

Scientific committee

Alfonso Del Percio, FHNW University of Applied Sciences and Arts
Northwestern Switzerland, Switzerland

Amy Wanyu Ou, University of Gothenburg, Sweden

Anne Larsen, University of Copenhagen, Denmark

Ashraf Abdelhay, Doha Institute for Graduate Studies, Qatar

Beatrice Zuaro, University of Copenhagen, Denmark

Bettina Migge, University College Dublin, Ireland

Charlotte Sun Jensen, University of Copenhagen, Denmark

Daniel Silva, Federal University of Santa Catarina, Brazil

Dave Sayers, University of Jyväskylä, Finland

Elisabetta Adami, University of Leeds, United Kingdom

Francis Hult, University of Maryland, Baltimore County, United States

Gavin Lamb, NHH Norwegian School of Economics, Norway

Hartmut Haberland, Roskilde University, Denmark

Ico Maly, Tilburg University, The Netherlands

Janus Spindler Møller, University of Copenhagen, Denmark

Joana Plaza Pinto, Federal University of Goiás, Brazil

Joyce Kling, Lund University, Sweden

Karin Tusting, Lancaster University, United Kingdom

Kristin Vold Lexander, University of Inland Norway, Norway

Magda Pischetola, University of Copenhagen, Denmark

Magdalena Madany-Saa, University of Oslo, Norway

Manuel Padilla Cruz, University of Seville, Spain

Marella Tiongson, University of Copenhagen, Denmark

Marian Flanagan, University of Copenhagen, Denmark

Martha Sif Karrebæk, University of Copenhagen, Denmark

Miguel Pérez-Milans, University College London, United Kingdom

Nicolai Pharaos, University of Copenhagen, Denmark

Scientific committee, cont.

Nikolas Coupland, Cardiff University, Wales

Ron Darwin, The University of British Columbia, Canada

Sari Pietikäinen, University of Jyväskylä, Finland

Shaila Sultana, BRAC University, Bangladesh

Sibonile Mpendukana, University of Cape Town, South Africa

Slobodanka Dimova, University of Copenhagen, Denmark

Spencer Hazel, Newcastle University, United Kingdom

Sune Sønderberg Mortensen, Roskilde University, Denmark

Tanya Karoli Christensen, University of Copenhagen, Denmark

Virginia Zavala Cisneros, Pontifical Catholic University of Peru, Peru

Viviane de Melo Resende, University of Brasília, Brazil

Organising committee

Janus Mortensen, University of Copenhagen, Denmark

Sam Goodchild, University of Copenhagen, Denmark

Kasper Engholm Jelby, University of Copenhagen, Denmark

Jens Christian Borup Green Jensen, University of Copenhagen, Denmark

Sanne Larsen, University of Copenhagen, Denmark

Rafael Lomeu Gomes, University of Copenhagen, Denmark

Sofie E. Abrahamsen Søndergaard, University of Copenhagen, Denmark

Keynote presentations

Rodney Jones

University of Reading, United Kingdom

Animating AI

In this talk I will explore the ways humans ‘animate’ AI – or ‘call it into being’ – through language, drawing on work in linguistic anthropology on animism, animation and the performative nature of speech genres (Bird-David 1999; Bauman 2004; Silvio 2010). While ‘calling AI into being’ is a complex, distributed process involving multiple social actors and multiple genres (from the databases on which models are trained to the demos that AI companies stage to market their products), in this talk I focus on the *prompt* as a foundational genre through which users participate in animating AI chatbots, transforming them into legitimate participants in recognisable social practices. The prompt, I will argue, is a hybrid genre that functions simultaneously as a set of machine instructions and as a performative utterance deeply affected by human social imaginaries. As such, it inherits two seemingly incompatible lineages: the procedural pragmatics of programming and the more ‘mystical’ pragmatics of spells, invocations, and magic tricks.

Through an analysis of a corpus of ‘metapragmatic artefacts’ on prompting (consisting of things like industry manuals, commercial course materials, media stories, TikTok videos and conversations on Reddit), I examine how these two lineages collide in contemporary prompting repertoires across three sites: the professional discourse of ‘prompt engineering’; the collaborative experimentation of everyday users; and the more ‘magical’ prompting practices promoted by influencers and entrepreneurs who offer ‘secret’ prompts which they promise followers will help them to ‘wake up their AI’. These three repertoires index radically different ideologies of language and communication, with industry manuals and ‘prompt engineering’ courses constructing prompting as a high-stakes form of linguistic optimisation, ordinary users treating it as an experiential social practice, and ‘prompt gurus’ depicting it as performative and transformational. They also index different socio-technical imaginaries regarding human relationships to technology. As users engage in these different (often intersecting and overlapping repertoires), they don’t just animate AI, but are also *animated by* it, transformed into different ‘kinds of humans’, such as ‘elite language workers’, ‘tech-savvy tinkerers’,

‘adversarial hackers’, and ‘AI whisperers’ able to coax models into consciousness. What unites all of these approaches, though, is an understanding that prompting is not simply a matter of formulating instructions or commands but an inherently anticipatory and improvisational activity deeply embedded in what Bird-David (1999) calls the ‘relational epistemologies’ of animism. It is a form of ‘linguaging’ through which users learn what AI can do and what AI can be by trying to make it do things and ‘be’ things. Prompts are probes, wagers, world-building moves, which call forth identities, scenarios and imaginaries both for technologies and for the humans that use them.

- Bauman, Richard. 2004. *A world of others’ words: Cross-cultural perspectives on intertextuality*. Wiley. <https://doi.org/10.1002/9780470773895>
- Bird-David, Nurit. 1999. ‘Animism’ revisited: Personhood, environment, and relational epistemology, *Current Anthropology* 40 (S1): S67–S91. <https://doi.org/10.1086/200061>
- Silvio, Teri. 2010. Animation: The new performance? *Journal of Linguistic Anthropology* 20 (2): 422–438. <https://doi.org/10.1111/j.1548-1395.2010.01078.x>

Languages without bodies? Sign language AI as language policy

This talk uses sign language AI as a lens for a broader argument: that the infrastructural decisions shaping language AI systems are, in effect, decisions about language policy: what languages are, which varieties count as legitimate, and whose linguistic practices get to matter. Sign languages are an instructive case because they are minoritised *and* embodied, and sit at the intersection of language minoritisation and disability – a conjunction that remains under-researched in sociolinguistics.

Drawing examples from current sign language AI development, I explore four patterns. First, I look at International Sign, which is widely understood by deaf communities as a fluid, ad hoc contact practice, yet increasingly appears in AI training datasets as just another stable, codifiable language label. This is revealing because it is a case where a community's own metalinguistic awareness makes the datafication logic of AI systems most visible. Second, AI developers often train systems on larger sign language datasets (frequently American Sign Language) and adapt them for smaller sign languages, making language boundaries appear technically optional, while installing data-rich sign languages as infrastructure for smaller ones. Third, commercial interests determine which sign languages get built into products at all: most of the world's sign languages are simply absent from the market. Fourth, when companies deploy recognisable deaf signer personas to front their AI output, they borrow community authority and leverage the indexical value of a familiar, trusted deaf body, while reshaping what authentic signing is understood to look like.

These patterns are not hypothetical: sign language AI is already being deployed, and with it, sign languages risk being detached from the embodied, community-specific, and deaf/disability-shaped practices that constitute them. While this can open new creative and political possibilities, it matters who drives it, under what conditions, and with what consequences. I conclude by arguing that sociolinguistics needs to follow language policy into the spaces where these decisions are currently being made – dataset curation decisions, benchmark design choices, and procurement standards – and that it needs to bring disability and embodiment into the analytical frame when it does.

Nicole Holliday

University of California, Berkeley, United States

Evaluative speech AI and the devaluation of sociolinguistic competence

Linguists take it as axiomatic that speakers are experts on their languages, both in grammar and usage. However, as Large Language Models (LLM) trained on text and speech become ubiquitous in domains from daily tasks to education and employment, human expertise about language is increasingly devalued. This talk will present the results of three studies that focus on LLMs that are designed to evaluate and “improve” the speech of human talkers; these systems are known as Socially Prescriptive Speech Technologies (SPSTs). The first study shows how the Amazon Halo, a wearable device that claims to evaluate “tone of voice” does not function as advertised, and in fact systematically negatively evaluates the speech of Black talkers and women. The second study focuses on systems such as Read. AI and the Zoom Revenue Accelerator, which claim to evaluate communicative effectiveness in videoconferencing contexts. Results of a laboratory experiment comparing these products’ evaluations of speakers show evidence of systematic bias against black speakers and individuals who identify as neurodivergent, while also reinforcing “standard” language ideologies and failing to provide consistent, actionable feedback to users. Finally, the third study analyzes the outputs of “accent translation” programs marketed by companies such as Sanas and Krisp, showing that such programs do not functionally “translate” accents but rather transform speech to an imagined “American” style that is poorly evaluated by human listeners. Taken together, these studies show that “AI”-based programs that purport to evaluate human speech do so without consideration of linguistic principles or acknowledgement of speakers’ sociolinguistic competencies. Such systems also act without transparency for both designers and users by design, reproducing social stereotypes inherent to their training data. As a result, they advise humans to produce unnatural speech, and they punish speakers who do not conform to the narrow targets established by an LLM’s training data. As such technologies are already being used to make employment decisions, provide speech therapy, and even draft police reports, the fact that these systems systematically misevaluate speech represents a significant threat to all people, but most especially those from marginalized groups.

Britta Schneider

University of Vienna, Austria

The ungrounded sign: Meaning-making, community and democracy under machine influence

In this talk, I ask what happens to linguistic signs in societies that make use of algorithmic language technologies and what this may imply for community formation and democratic culture. The discussion is based on theoretical considerations concerning sign making and semiotic ideology from linguistic anthropology (Gal and Irvine 2019; Keane 2018; Silverstein 2014), where shared meanings of linguistic signs and human communities are understood as dialectically producing each other. In this sense, the development and sharing of linguistic signs is a collective process that produces human sociality. The linguistic signs generated by LLMs are not based on procedures of human interactive meaning-making and seem to be generated simultaneously within and across communities and integrate machine-learning logics into human interaction. LLMs thus manufacture signs but have no community grounding. Which kinds of signs and communities come into being where commercial algorithmic and big data logics interfere with the ability of humans to create signs, meanings and community? What happens if non-accountable signs circulate in so far unidentified forms across communities? Who has the power to define the meaning of signs in such techno-linguistic formations? And how can meaning, trust and truth – and thus democratic discourse culture – be ensured where signs are co-produced by algorithmic machines?

In order to shed light on these questions, I discuss examples of public debates and current policies concerned with governing language technology. Traditional national language-norming institutions often react to new language actors by claiming that signs created by machines are ‘inauthentic’ and thus reinvolve traditional epistemologies of humans as autonomous, rational beings and communities and linguistic signs as stable and territorially ordered. Big Tech leaders exploit collective, community-based human semiotic resources at global scales and produce discourses that construct an image of language technologies as isolated from society, freeing them from social responsibility. At the same time, populist politicians dream of determining the meanings generated by machines and

simultaneously engage in individualistic strategic appropriations of linguistic signs that remind of Orwell's Newspeak. They thus contribute to the era of post-truth and promote totalitarian meaning-regimes that are envisioned for the entire planet.

Drawing on these observations of public contestations over meaning, power and responsibility, I develop thoughts on which literacies are needed to support democratic institutions and community-based agency in societies shaped by unaccounted-for machine-created signs.

Gal, Susan & Judith T. Irvine. 2019. *Signs of difference: Language and ideology in social life*. Cambridge University Press.

Keane, Webb. 2018. On semiotic ideology. *Signs and Society* 6 (1): 64–87.
<https://doi.org/10.1086/695387>

Silverstein, Michael. 2014. Denotation and the pragmatics of language. In *The Cambridge handbook of linguistic anthropology*, edited by N.J. Enfield, Paul Kockelman & Jack Sidnell, 128–157. Cambridge University Press.

Session 1A

AI and education · Part 1

Baraa Khuder

Chalmers University of Technology, Sweden

Legitimacy of AI feedback across feedback ecologies: Disciplinary vs. non-disciplinary peers

Generative AI tools are increasingly used as feedback providers in academic writing, yet the legitimacy of their feedback is not inherent to the technology but socially produced in interaction. Rather than asking whether AI feedback is useful, this study examines how and under what conditions AI feedback becomes legitimate. Focusing on doctoral writing, the study investigates two AI-supported feedback ecologies: (a) generative AI combined with disciplinary peer feedback and (b) generative AI combined with non-disciplinary peer feedback. Drawing on models of the social regulation of learning, the analysis foregrounds concerns with legitimacy in AI-assisted feedback.

Fifty-five PhD students participated in one of the two conditions. Data comprised AI interaction histories, 14 audio-recorded small-group “listening room” discussions, and short individual reflections comparing AI and peer feedback. Discussion transcripts were segmented and coded for forms of regulation (self-, co-, socially shared) and regulatory functions, with analytic attention to how participants positioned AI and negotiated the legitimacy of its suggestions.

Across both ecologies, AI did not participate in genuinely socially shared regulation. In AI + disciplinary peer groups, AI feedback was more often treated as legitimate when ratified by disciplinary expertise, rendering AI’s authority conditional and accountable. In AI + non-disciplinary peer groups, AI was more often recruited as a resource for self-regulation, with participants more explicitly evaluating and adapting its suggestions. The study shows how the legitimacy of AI feedback shifts across interactional contexts and discusses implications for doctoral writing pedagogy and the politics of legitimacy in human–AI feedback practices.

Diane Potts

Lancaster University, United Kingdom

Writing a story, writing for the machine: Young children critiquing AI outputs

The development of large language models (LLMs) and the rapid uptake of publicly available AI tools and agents have inevitably led to debates about their place in language and literacy education. Quickly, scholars have produced and/or modified agendas (Ng et al. 2022; Su et al. 2023; Markl 2025), syntheses (Lee et al. 2024) and models (Darvin 2025; Jiang et al. 2025) to foreground AI. In contrast, this paper adopts a social semiotic perspective on meaning-making (Cope and Kalantzis 2024; Hasan 2005; Kress 2010) and explores embedding AI use within well-established pedagogic practices for developing young learners' writing. Three examples of children's engagement with an AI image generator are used to illustrate the productive challenges of writing AI prompts, learners' sensitivity to the intermodal relations of stories and AI-generated illustrations and their critical appraisal of AI.

Data is drawn from a participatory research project undertaken in a Northwest England primary in which 24-Key Stage 1 children, approximately 50% from multilingual homes, rewrote a classic children's story and then wrote prompts to generate images for their story. The data set includes observations and audio of 14 hours of lessons, children's planning documents, lesson materials, children's digital books and interviews with the participants.

The children's critical engagement with AI in the context of established practices raises questions about a) the design considerations accompanying educational use of AI; b) the beneficiaries of the commodification of AI literacy and c) sociolinguists' positionality in documenting AI use. I argue recognizing not only human researchers "creativity, reflexivity and criticality" (Curry et al. 2025) but also collaborators' and participants' is vital.

From reflection to self-regulation: Exploring EFL teachers' collaborative reflective practice and students' SRL in AI-mediated writing

Recent studies indicate that AI-mediated self-regulated learning (SRL) writing strategies may offer unique opportunities to support learners' planning, monitoring, and revision processes. Nonetheless, the successful integration of these strategies relies heavily on instructors' awareness, perceptions, and classroom practices—areas that remain underexplored in the current literature. Therefore, to address this gap, this study presents a mixed-methods exploratory research that examines the relationship between teachers' collaborative reflective practice (RP) and learners' self-regulated learning in English as a Foreign Language (EFL) writing instruction in higher education, mediated by generative artificial intelligence tools. Grounded in Wallace's (1991, 1998) reflective model and Zimmerman's (2002, 2008) self-regulated learning framework, the study aims to investigate how tertiary-level EFL writing instructors in a preparatory-year program use generative AI tools and collaborative reflection to adapt their writing instruction, and how these AI-mediated, data-driven practices relate to changes in students' self-regulated learning behaviors during the planning, drafting, feedback, and revision stages of essay writing. The study took place over an eight-week preparatory writing course at a Turkish university, involving three EFL writing instructors and 66 engineering students. Findings indicate that, when teachers systematically reflect on student data and AI-generated feedback, they become more aware of learners' SRL profiles and adapt their writing instruction accordingly, while students report and display increased metacognitive awareness, feedback literacy, and autonomous engagement with revising their texts.

Session 1B

Ethics and AI

Alex Ross

University of British Columbia, Canada

Who is responsible? Analyzing the discursive construction of “responsibility” and “ethical use” in institutional GenAI policies

As the use of generative artificial intelligence (GenAI) becomes increasingly prominent in higher education settings, stakeholders have sought to establish institutional policies which guide how GenAI is expected to be used for teaching and learning purposes. Such policies often acknowledge pedagogical advantages such as enhanced personalization and autonomous learning opportunities (e.g., Moorhouse et al. 2024; Szabó and Szoke 2024). At the same time, these policies address potential limitations encompassing GenAI use, with a tendency to focus on academic integrity concerns (e.g., Kumar et al. 2024; Luo 2024). This overemphasis on academic integrity leads to a backgrounding of other vital ethical issues raised by previous research on GenAI, including but not limited to the reification of standard language ideologies (Schneider 2022) and biases (Gillespie 2024), as well as environmental concerns (Crawford 2021). Despite these aforementioned risks involving GenAI use, very little empirical research to date has investigated the extent to which higher education GenAI policies address these broader ethical concerns. Thus, the present study addresses the overarching question of how “responsibility” and “ethical use” are discursively constructed in relation to students and instructors within public university GenAI policies. To answer this question, Critical Discourse Analysis (van Leeuwen 2008) was used to analyze institutional GenAI policies from five research-intensive universities in Canada. The findings suggest that these policies tend to narrowly construct “responsibility” through the lens of academic integrity concerns, thereby negating the discursive positioning of GenAI as a multifaceted ethical challenge.

Pejman Habibie

Western University, Canada

Beyond compliance: GAI and the sociolinguistic ecology of academic knowledge production

Large language models are increasingly integrated into scholarly publication worldwide. Within English for Research Publication Purposes (ERPP), discussions of generative artificial intelligence (GAI) have largely focused on compliance, misconduct, and individual responsibility. This presentation argues that such framings are conceptually limited. Drawing on sociolinguistic scholarship on language, power, and political economy, it reconceptualizes GAI ethics in scholarly publishing as a distributed rather than exclusively individual phenomenon. Scholarly publication is approached as a sociolinguistic practice embedded in a neoliberal techno-feudal knowledge economy structured by linguistic hierarchies, evaluative regimes, and global inequalities. GAI systems do not operate as neutral tools; they participate in shaping norms of legitimacy, voice, and epistemic authority. Trained on historically sedimented corpora, these systems risk reinforcing standardization, Anglophone dominance, and disciplinary gatekeeping. Ethical inquiry must therefore extend beyond plagiarism and disclosure to address epistemic integrity, linguistic diversity, and the reconfiguration of scholarly agency. Building on ERPP scholarship, the presentation proposes a distributed ethical framework that situates responsibility across human actors, institutions, publishing infrastructures, and algorithmic systems. Rather than asking whether individual scholars use GAI ethically, the analysis examines how these technologies reshape the sociolinguistic ecology of knowledge production, including the positioning of multilingual scholars and the construction of academic legitimacy. Ethics emerges not only at moments of rule violation but at moments of self-positioning within structurally unequal systems.

Sae saem Yoon, University of Washington, United States
Jason Litzenberg, University of Oslo, Norway

Distributed accountability in the AI guidance of U.S. higher education: An ecolinguistic analysis

Large Language Models – the “backbone” of most AI systems – disproportionately impact vulnerable ecosystems and communities because of their resource-intensive hardware production and operational requirements. The presentation reviews findings from a recent ecolinguistic analysis of AI guidance documents for faculty at several large, U.S.-based research universities. An ecolinguistic approach allows for consideration of how discourses can obscure the relationship between humans, technology, and the environment. Using Stibbe’s (2021) framework of erasure – void, mask, and trace – the analysis examines how environmental responsibility is constructed, minimized, or omitted in institutional guidance. The findings reveal that universities tend to sideline ecological concerns while framing AI as a pedagogical benefit (efficient, innovative, and inevitable) while foregrounding neoliberal priorities (productivity, compliance, and risk management). This framing positions AI ethics at the level of the individual, allowing institutions to avoid articulating a coherent stance. The presenters challenge this individualized model of ethics via the concept of distributed accountability (Stockwell and Wang 2025) – i.e., the recognition that responsibility for AI’s environmental impacts must be shared across institutions, policy frameworks, technological systems, and users. Incorporating environmental sustainability as part of AI ethics is a natural extension of the field’s existing commitments to social ethics and future-oriented language teaching, learning, and research.

Stibbe, A. 2021. *Ecolinguistics: Language, ecology and the stories we live by*. Routledge.

Stockwell, G., and Y. Wang. 2025. Framing human-AI collaboration in language education. *Technology in Language Teaching & Learning* 7(4): 1–16.

Session 1C

Negotiating AI

Hatime Çiftçi, MEF University, Turkey

Derya Kulavuz-Onal, Salisbury University, United States

Melike Akay, University of South Florida, United States

Selman Ecevit, MEF University, Turkey

İlayda Çolak, MEF University, Turkey

Positioning the algorithm: How students negotiate legitimate AI use and literacy in higher education

Since ChatGPT's public release in late 2021, Artificial Intelligence (AI) has been a central topic in higher education (HE) as students and instructors began negotiating what constitutes legitimate use of AI tools, with some considering any AI use as 'cheating'. However, as generative AI tools improve and the uses become highly complex, a simple 'cheating' does not capture it all, and we need a deeper understanding of how individuals, particularly HE students, position themselves and justify their AI use as legitimate. Drawing on positioning theory (Davies and Harre 1999), with a focus on students' reflexive and interactive positioning, as they "provide a way of uncovering participants' identities" (Kayi-Aydar 2019, 12), we analyze and demonstrate how students position themselves, AI tools, and others, when discussing and claiming their AI use as legitimate. Our data come from 23 semi-structured interviews with HE students across disciplines in Turkey and the United States, including multilingual writers and non-traditional students. Using a discourse analytic approach (Gee and Handford 2012), we have identified recurrent strategies for self-positioning, AI-positioning, and other-positioning as the three positioning moves in students' accounts of AI use. Our findings indicate that undergraduate students use multiple discourse strategies to enact their three-layered positioning, such as giving justifications, using disclaimers and metaphors, expressing self- and other-criticism, making subjective evaluations, invoking authority, and comparing/contrasting. We argue that these positioning moves and strategies are enactments of how university students understand and experience their own AI use and literacy from academic, ethical, and sociocultural aspects.

Jens Christian Borup Green Jensen & Samantha Goodchild
University of Copenhagen, Denmark

Negotiating ideologies of efficiency and thoroughness: STEM and social sciences researchers' positionings and their GenAI (non)use

Among the burgeoning research on text-generative AI (GenAI) use in universities there has been a predominant focus on teaching and learning. Less attention has been paid to the effects GenAI has on research, based on reported and documented research practice. We contribute to this growing field with two linguistic ethnographic case studies. Investigating the relationship between researchers' reported practices and their observed practices, we ask if, and how, the use of GenAI affects researchers' professional identities. Using data gathered from the social sciences and STEM from two universities, we present a critical discourse analysis of interviews focussing on the researchers' own articulations of their use of GenAI. We then complement these findings to screen recordings of researchers' GenAI use to show how identity work is being negotiated in the face of the affordances and constraints of GenAI. Using both emic and etic categories such as 'efficiency' and 'thoroughness', our analyses show the researchers engaging in ideological positioning work (Gal and Irvine 2019), articulating various nested professional identities as experts, researchers, artists, and writers; all of which are up for (re)negotiation when considering the (non)use of GenAI in current academic research. Finally, we situate our findings within wider ongoing discourses on speed, productivity and integrity in the academy, and how these potentially are being affected by the use of GenAI.

Gal, Susan, and Judith T. Irvine. 2019. *Signs of difference: Language and ideology in social life*. Cambridge University Press.

Sam Goodchild & Sanne Larsen

The University of Copenhagen, Denmark

Collaborative research writing in the age of GenAI: Distributed cognition and expertise between academics and ChatGPT

While the use of text-generative artificial intelligence (GenAI) in academic research and publishing practices has generated much discussion, most empirical research thus far has explored academics' practices and perceptions of GenAI through large-scale surveys. Drawing on an ongoing linguistic ethnographic study of science researchers at a Danish university, in this presentation we investigate scholarly writing practices that incorporate the use of GenAI. Focusing on participant-made screen recordings, we show how two researchers collaborate on writing a book chapter, an uncommon genre for their field. The interactional data demonstrate how they organize the outline of their chapter in the ideational phase of writing, negotiate the content together, and then use ChatGPT to check if the text suits the genre requirements. From interview data, it emerges that the researchers explain giving ChatGPT certain writing tasks, like creating an abstract, because they perceive it can perform these tasks "way better" than themselves. We subsequently analyse the two types of data from the multifaceted perspective of distributed cognition of the collaborative writing process (Clayson 2018). Based on the results, we critically discuss the extent to which established concepts such as distributed cognition and negotiated expertise can be used to understand collaboration and co-production of scientific knowledge with GenAI.

Clayson, Ashley. 2018. Distributed cognition and embodiment in text Planning: A situated study of collaborative writing in the workplace. *Written Communication* 35 (2): 155–181.

Session 1D

AI and research on multilingualism

Alice Coen, University of London, United Kingdom
Antonina Zhiteneva, University of Oxford, United Kingdom
Jelena Kojashi, University of New York Tirana, Albania
Teodora Mašković Đeri, University of Belgrade, Serbia
Yasmeen Shamshoum, Tel Aviv University, Israel
Yliana V. Rodríguez, Universidad de la República, Uruguay

AI, multilingualism and standardisation: Sociolinguistic perspectives on experimental language design

Large language models are increasingly used to generate multilingual research materials. From a sociolinguistic perspective, however, AI-mediated text production is not merely technical but involves processes of standardisation and decontextualization. This presentation explores AI-assisted construction of multilingual experimental materials as a site of sociolinguistic mediation. The materials in question are dilemma vignettes testing moral decision-making. To enhance experimental realism, translations should be contextually adapted while ensuring pragmatic equivalence. Drawing on the design process of our own multilingual experimental materials, we analyse the comments given by academics in translation studies to assess AI-generated translations across a diverse set of languages. Results suggest that AI-generated translations are often effective at conveying core semantic content, but frequently fall short in capturing tone, idiomatic expression, and moral nuance. In some languages, AI outputs only required minor edits to improve naturalness, while in others they were clearly identifiable as machine-translated. These findings highlight considerable variability in translation quality across languages: less widely represented languages are particularly vulnerable to lower quality output. They also stress the importance of accounting for linguistic variation within languages, as AI may not reliably distinguish between regional varieties. Taken together, this evidence suggests that AI translation currently supports semantic transfer more reliably than sociolinguistic equivalence. As such, the development of multilingual experimental materials cannot be treated as a purely technical task but requires careful cultural adaptation to ensure comparability across languages.

AI as research assistant: A solution to the challenges of multilingual research?

In this paper, we discuss the advantages and challenges of using AI translation tools when conducting sociolinguistic fieldwork in multilingual sites where researchers and participants do not share all, or any, linguistic repertoires. We ask: How can AI become part of the interactional repertoire between researchers and participants? How may such use of AI impact our research in terms of methodology, theory and ethics?

The paper is based on current ethnographic fieldwork in a multilingual farm in Denmark with employees from Denmark, Latvia, Romania and Syria. English is mainly used as a lingua franca in the workplace, but with clear limitations for employers, employees and researchers alike. From a research perspective, these limitations include gaining informed consent and relationship-building, especially with Romanian-speaking research participants. Not only language, but also time is a significant constraint as many of the Romanian-speaking employees work on short-term contracts. This has led to our use of AI translation tools in an attempt to inform these participants about our project as well as establish contact with them as quickly as possible.

So far, we have noticed two aspects that stand out as different from human-mediated translation/interpretation used in multilingual research: Blindsiding and accountability. Both of these relate to missing human agency in the translation/interpretation process. We become blindsided by the AI tools because we cannot check understanding with them, nor can we rely on accountability as directing these tools as we would a human translator.

Through examples from our fieldwork, we will discuss these two aspects, presenting how we have worked with them and the dilemmas they have posed.

Sara Ganassin, Newcastle University, United Kingdom
Jessica Bradley, University of Sheffield, United Kingdom
Elizabeth Holland, Radboud University, Netherlands

AI and community-based research with diverse communities: Interdisciplinary methodological reflections from intercultural communication and AI studies

This paper proposes a methodological reflection on community-based research in diverse contexts in the UK, drawing on interdisciplinary perspectives from intercultural communication, AI studies and cognitive computing. We discuss challenges and affordances of AI tools (e.g. chatbots and translation technologies) in community and creative contexts, including refugee support groups and participatory arts and health programmes. Drawing on our experience of ethnographic research with groups who might be considered vulnerable in different ways, we offer insights into how AI is rapidly becoming embedded in people's everyday lives in the 'post-digital' age and shaping research practices.

First, we explore the increasing presence of AI tools in people's day-to-day lives, for example the use of AI translation to support interaction, AI-edited imagery (e.g. memes), and applications with AI components (e.g. language-learning apps). Acknowledging that AI is becoming part of normal and unremarkable multilingual and multimodal communication requires researchers to engage critically and expand epistemological lenses when navigating this rapidly evolving landscape.

We then reflect on the methodological implications of conducting community-based research in contexts where AI mediates interaction and shapes meaning-making. We emphasise that AI tools which are becoming increasingly embedded into the 'everyday' are not unbiased. We therefore argue for reflexive, ethical and intercultural research approaches that centre participants' lives, communicative repertoires and creative practices, while maintaining human agency and critical awareness of how AI shapes power relationships, language ideologies, care practices and knowledge production.

Session 2A

Linguicism and GenAI in HE

Marion Döll, Europa-Universität Flensburg, Germany (convener)

PANEL ABSTRACT · Linguicism and GenAI in higher education

Linguicism (Skutnabb-Kangas 1988) denotes language-based discrimination structurally embedded in society. It manifests across all life domains - not only in interpersonal interactions but also in institutional discourses and practices. This is particularly evident in education systems of officially German-speaking countries, where students whose first language is not German and those from non-academic backgrounds face disadvantages from primary through to tertiary education.

GPT tools hold potential to diminish the impact of language biographies on educational success and reframe deficit-oriented language discourses by partially relativizing the importance of instructional-language proficiency. Building on Skutnabb-Kangas' structural understanding of linguicism, we argue, nevertheless, that the discrimination-reducing potential of generative AI (GenAI) is only being exploited to a limited extent, i.e., linguicism manifests not only as tool-inherent bias, but presumably also in the discourses and practices surrounding the technology.

Using higher education as an example, this panel examines how linguicism is reproduced in the context of GenAI. Qualitative and quantitative data from Germany, Austria, and Canada enable reconstruction of divergent educational, linguistic, and migration policy impacts. Findings reveal extensive linguicist situations and structures concerning GPT tools across universities in these countries. Crucially, GPT tools' potential to reduce linguicism remains not only untapped but also invisible in institutional discourses.

Skutnabb-Kangas, T. 1988. Multilingualism and the education of minority children. In *Minority education: from shame to struggle*, edited by T. Skutnabb-Kangas and J. Cummins. Multilingual Matters.

Sabine Guldenschuh
Universität Wien, Austria

Usage of GenAI in higher education: A critical discourse analysis on guidelines and training courses in Germany and Austria

Following Skutnabb-Kangas (2015), monolingual educational systems in countries like Germany and Austria are per se linguisticist. Especially students whose first language is not German and/or students from non-academic backgrounds are prone to suffer linguisticist experiences throughout their studies as their language use is perceived as faulty and/or the wrong register. In theory, GPT tools could have the potential to address these language disadvantages and therefore experiences of linguisticism.

Being part of an international research team (Germany, Austria, and Canada) conducting the pilot study “Reduction and reproduction of educational disadvantages through AI-based generative NLP technology in the multilingual migration society (BibeKI)”, I focus on presenting the results of our critical discourse studies (Wodak and Meyer 2016) on initial tendencies of discourses in education to partly answer the question of who is allowed to use tools for computational offloading concerning speech/text production (e.g. on term papers), when and to what extent. By analyzing corpus notes from seven training courses on AI for university lecturers and guidelines on dealing with AI in university teaching (a corpus of around 19,500 words), it can be argued that discrimination lies within the discourse itself as racism/linguicism – defined as dichotomizing, homogenizing, naturalizing, and hierarchizing supposedly different groups of students – becomes clear.

Skutnabb-Kangas, T. 2015. *Linguicism. The encyclopedia of applied linguistics*. Wiley.

Wodak, R., and M. Meyer. 2016. Critical discourse studies: History, agenda, theory and methodology. In *Methods of critical discourse studies*. 3rd ed., edited by R. Wodak and M. Meyer. Sage.

Marion Döll

Europa-Universität Flensburg, Germany

Race, class, language, and GPT Tools: Everyday linguisticism in higher education in Germany and Austria

Like other forms of discrimination, linguisticism can manifest structurally within institutions through formal and informal regulations, routines, and discourses, as well as in interindividual interactions (Gomolla 2023). In a BibeKI subproject, we therefore investigated the extent to which students in higher education in Germany and Austria experience linguisticist situations in connection with GPT tools at their universities. To this end, students in Upper Austria and Schleswig-Holstein (N = 1179) were surveyed using an online questionnaire. In addition to data on usage behavior, standardized scales were used to collect data on attitudes toward and knowledge on GenAI. To assess potential linguisticist everyday experiences related to GPT tools, we developed new items.

Statistical analyses reveal that students whose first language is not German and those from non-academic backgrounds in both countries are significantly more likely to experience linguisticist situations than other students. Furthermore, we found significant differences between Germany and Austria, and between teacher training students and other students. The results confirm the assumption by Dirim and Pokitsch (2018) that, in linguisticism, the focus is not on languages per se, but on the devaluation of certain groups with reference to their language. Furthermore, the findings suggest that, in addition to language and migration policy context conditions, disciplinary culture also influences the reproduction of linguisticism.

Dirim, İ., and D. Pokitsch. 2018. (Neo-)Linguizistische Praxen in der Migrationsgesellschaft. *Osnabrücker Beiträge zur Sprachtheorie* 93: 13–32.

Gomolla, M. (2023). Direkte und indirekte, institutionelle und strukturelle Diskriminierung. In *Handbuch Diskriminierung*, edited by A. Scherr, A. El-Mafaalani, and G. Yüksel, 171–194. Springer.

Lama Alisamy

Europa-Universität Flensburg, Germany

Linguicism related to AI tool use in higher education in Austria and Germany

The analysis of the BibeKI survey data showed that a considerable percentage of students use AI tools intensively in their studies and that some students experience linguicism in interactions with fellow students and lecturers in relation to the use of GPT tools. Building on these findings, we aim to gain more in-depth and complementary insights into linguist experiences of multilingual students associated with the use of these tools. Therefore, semi-structured in-depth interviews were conducted. The interview guide was thematically structured. The first and second sections focus on participants' language profiles, learning trajectories, and sociodemographic characteristics. The third and fourth sections address the purposes and intensity of AI use as well as linguist experiences related to the use of GPT tools at higher education institutions. The final part addresses perceived changes in writing skills and trust in AI tools. Nine students have been interviewed so far. The target group are students whose first language is not German and who either grew up in Germany or Austria or have sought refuge there since 2015. Six participants reported Arabic as their first language, while three reported Turkish. The interviews are analyzed using evaluative and typifying qualitative content analysis (Kuckartz 2018).

Preliminary findings indicate that the students predominantly perceive GPT tools as very useful and supportive for learning, particularly in dealing with linguistic barriers. Linguist experiences mainly manifest in assumptions about AI usage and uncertainty due to unclear regulations at higher education institutions.

Kuckartz, U. 2018. *Qualitative Inhaltsanalyse. Methoden, Praxis, Computerunterstützung*. 4th ed. Juventa.

Session 2B

Theorising AI

Takeshi Enomoto

The University of Osaka, Japan

AI as a language-ideological phenomenon: A sociolinguistic perspective

This paper develops an outlook on AI from a sociolinguistic perspective. In doing so, it addresses two interrelated questions: (1) how context-sensitive approaches to language-in-use enable us to understand AI—particularly LLMs—as a language-ideological phenomenon, and (2) how these approaches illuminate an emerging horizon of metapragmatic discourse involving both humans and AI.

LLMs generate responses by predicting probable sequences of linguistic forms from distributional patterns in massive textual corpora. This operational logic raises a fundamental question: what counts as “language” for LLMs? From a sociolinguistic perspective, LLMs appear to presuppose a model of language that privileges formal structure, calculability, and pattern regularity. Such a model is never theoretically neutral. It echoes and reproduces a longstanding language ideology in which language is imagined as an autonomous system detachable from situated practice.

Accordingly, LLMs are more productively approached not as models of language but as technological enactments of particular language ideologies. Their apparent communicative competence renders newly visible the metapragmatic assumptions that make AI appear intelligent. While contemporary LLMs achieve this effect through the decontextualization of language, sociolinguistic inquiry reintroduces dimensions such as indexical anchoring and the vulnerabilities inherent in social interaction. By situating AI within these broader regimes of language ideology, this paper argues that sociolinguistic perspectives foreground AI’s role as a reflexive trigger that interrogates the conditions under which language comes to index intelligence, agency, and social presence.

Joseph Sung-Yul Park

National University of Singapore, Singapore

Large language models as commodification of language: Towards a sociolinguistic critique of the political economy of AI

Calling for a more explicitly political economic critique of AI, this paper argues that large language models (LLMs) must be understood as a case of commodification of language (Cameron 2005; Heller 2010). While sociolinguistic research has been effective in problematizing the ideologies underlying the current AI boom driven by global tech capitalism (Migge and Schneider 2025; O'Regan and Ferri 2025), the field's socially-grounded view of language can reveal how LLMs derive their value as commodities from the human doing inherent in language. Research shows that much of AI is powered not by the algorithm but by human actors who generate, verify, and moderate data, and in this sense AI does not autonomously generate value, but extracts value from concrete labor of human beings (Casilli 2025; Muldoon et al. 2024). Drawing from this insight, this paper argues that LLMs do not generate value from linguistic data as static texts, but extract value from language as doing—language as creative, embodied life-activity through which we transform the world around us (Holloway 2010)—underlying that “data.” This is illustrated through a critical discussion of various limitations of current LLMs, including how LLMs trained on their own output eventually fail (Shumailov et al. 2024), and how LLMs tend to get lost in multi-turn conversations (Laban et al. 2024). These phenomena demonstrate that the seemingly autonomous power of LLMs to produce natural language is demystified once we take out the human doing underlying its capabilities. This provides a firm basis for critiquing LLMs as an instance of commodification of language—how they extract value from the doing of human beings, to accumulate profit in the hands of global tech capital.

Jason Litzenberg
University of Oslo, Norway

Dispensing remainder humanism and embracing “good enough”: Communication in the era of the language singularity

Remainder humanism is the desire to continually (re)define what constitutes “human” by positioning it in opposition to what AI systems cannot do (Weatherby 2025), exemplified in arguments such as “[T]hese systems cannot really replace humans, replace the quality of human craft and thinking – so many of their capacities are overblown” (Guest et al. 2025). Yet the fact that countless AI-supported tools already regularly aid contemporary translation and interpretation work (e.g., business, medicine, service industry, education, social work, and so forth) suggests that these technologies, while imperfect, are at least “good enough” to satisfy numerous pragmatic goals. Indeed, continued emphasis on what AI can’t do ignores the ways AI already has reorganized everyday communicative practice. This presentation introduces the concept of the Language Singularity, an AI-induced inflection point that fundamentally alters communication practices and conceptualizations of both human (mind/body, human/machine, etc.) and language (native/non-native, fluency, agency, etc). The Language Singularity provides a theoretical lens for envisioning the end of remainder humanism, (re)contextualizing contemporary communicative assemblages, and preparing for future technology-induced social change. While one can debate about when it actually started, the Language Singularity, as argued in this presentation, has already begun.

Weatherby, L. 2025. *Language machines: Cultural AI and the end of remainder humanism*. University of Minnesota Press.

Guest, O., et al. 2025. Against the uncritical adoption of ‘AI’ technologies in academia. Zenodo. <https://doi.org/10.5281/zenodo.17065099>

Session 2C

AI and writing

Naomi Wells

University of London, United Kingdom

AI, multilingual knowledge infrastructures and epistemic justice

The role of AI tools in processes of academic knowledge production has become an increasing focus of discussion and controversy, but rarely explicitly addressed are the implications of these tools in relation to multilingual publishing and connections to earlier discussions of the geopolitics of academic writing and epistemic inequalities (Canagarajah 2002).

More specifically, these tools raise questions surrounding their potential roles in overcoming barriers to publishing in dominant Anglophone journals for speakers of English as a second language (Van Noorden and Perkel 2023), and in supporting translation and more equitable knowledge exchange across languages. At the same time, their use raises ethical, epistemological and practical questions surrounding their role in reinforcing Anglophone linguistic and cultural hegemony, the unequal development of Large Language Models and access to AI tools for different language communities, and the environmental costs that are likely to fall most heavily on those communities least able to benefit from their use (Bowker 2025).

This paper draws on a series of semi-structured interviews with editors, publishers and researchers involved in multilingual publishing initiatives in different locations across the world that explore current uses of digital technologies to support multilingual publishing and attitudes towards the potential future uses of rapidly evolving AI tools in their own and their authors' publishing workflows. While acknowledging the rapidly evolving nature of these issues, the paper will further insight into the potential future uses and development of AI tools in relation to epistemic inequalities and multilingualism in academic publishing.

Hanna-Mari Pienimäki
University of Helsinki, Finland

The use of generative AI in professional writing: Negotiating textually mediated expertise in hybrid systems

Despite a growing recognition that generative AI (GAI) is fundamentally reshaping literacy practices across domains, the current analytical focus remains primarily in educational contexts. This perspective must urgently be extended to cover more varied forms of contexts, and to include practices of professional writing to gain a better understanding of the scale of the transformations. The paper will present a methodological proposal and preliminary findings of a four-year (2026–2029) project *Disruption in text production? The use of generative AI in professional writing* (GAI-WRITE) funded by the Research Council of Finland. The GAI-WRITE project adopts a linguistic ethnographic approach to study the use and effects of GAI tools in professional writing.

Drawing on the sociolinguistics of work and the sociolinguistics of writing, I examine how professionals use GAI to broker the textual realization of expertise. The analysis employs two complementary frameworks of expertise to explore the use of GAI in textually mediated knowledge work. Using the lenses of epistemic and relational expertise to analyze business and academic writing case examples, I interrogate forms of expertise that resist algorithmic capture and show how professional writers, more *and* less successfully, combine their own non-codifiable, local and tacit expertise with the affordances of GAI. The analysis demonstrates how expertise is distributed and negotiated within human-AI hybrid systems, and how this distribution of labor maps onto the linguistic level. I argue that if the completion of a task requires relational expertise, GAI can create unanticipated frictions, whereas in successful appropriations of GAI to complement human expertise, it functions as an epistemic enhancer.

Mieke Vandenbroucke, University of Antwerp, Belgium

Maja Brumer, University of Antwerp, Belgium

Luna De Bruyne, University of Antwerp, Belgium

Hans De Loof, University of Antwerp, Belgium

Using AI-transcription in medication reviews: A critical sociolinguistic analysis of accuracy and care

In this paper, we report on ongoing, empirical work on the use of AI-powered medical transcription systems in medication review encounters. Medication reviews conducted by pharmacists are structured, goal-oriented healthcare interactions designed to assess a patient's full medication regimen, including prescription drugs, over-the-counter products, and supplements. Typically held as planned, face-to-face consultations in community pharmacies, they allow time for detailed discussion and are framed as empathetic, patient-centred conversations that build trust. Pharmacists ask targeted questions, listen to patient experiences, and explain how medicines work, how to take them, and potential side effects, using semi-formal language that combines medical terminology with accessible explanations. A key feature is the production of a written summary by the pharmacist documenting discussed medicines, identified issues, agreed actions, and advice, serving as both a clinical record, a patient resource, and a requirement for reimbursement. Increasingly, such entextualised summaries are generated through automatic transcription systems. Studying medication reviews therefore requires attention to both the flow of the interaction itself and the reliability of its AI-generated written representation. In this paper, we analyse a Belgian corpus of (1) audio-recorded medication review encounters in which an AI transcription system was implemented and (2) their AI summaries. On this basis, we empirically assess professional claims that replacing manual note-taking with AI transcription enhances interactional patient-centredness and care, and critically examine how the introduction of AI reshapes the linguistic interaction and influences consultation outcomes and documentation quality.

Session 2D

AI and professional identity

Marian Flanagan, University of Copenhagen, Denmark

Nils Jäkel, University of Hildesheim, Germany

Janus Mortensen, University of Copenhagen, Denmark

Dorte Lønsmann, University of Copenhagen, Denmark

Generative AI and reflexivity: The case of future language professionals

Since ChatGPT and other LLM-enabled chatbots started to become part of the mainstream in 2023–2024, the notion of ‘AI’ has generated a flurry of cultural reflexivity, attested in metapragmatic discourse in a range of societal domains. With respect to the labour market, expected efficiency gains following ‘AI powered’ automatization have been much lauded in some quarters, but in professions where language and text are central, including e.g. translation and teaching, ‘AI’ has generally been received with some trepidation, if not outright anxiety. In this paper, we draw on survey data collected at three different points in time (April 2024, October 2024 and February 2026) among students at a foreign languages department at a Danish university to explore how students who entered university when generative ‘AI’ was in its infancy position themselves towards ongoing technological change and its potential impact on their (future) professional identities. Analyses of responses to closed questions as well as free-text answers indicate that the students’ reflexivity in relation to AI range across multiple scales, from immediate worry about the here-and-now of the study environment to concerns about the future usefulness of a foreign languages degree, to global concerns about ethics and sustainability. Changes in reported practice over time suggest that the use of ‘AI’ tools is potentially – to some extent – becoming normalised at the level of practice, but the stances which students express continue to exhibit tensions and ongoing critical reflexivity, which we argue can be seen as co-constitutive of ongoing sociolinguistic change, prompted by the emergence of ‘AI’.

Minyue Wu

University of Göttingen, Germany

Academic identity construction in the era of AI: International postdoctoral researchers in Germany

Recent studies indicate that generative artificial intelligence (GenAI) is widely adopted in academic practices. While AI-integration has reshaped literacy development and professional competence expectations, maintaining academic norms and integrity has become an increasingly critical challenge. Drawing on social identity theory (Gee 2000), this study conceptualizes scientific practices as central to academic identity construction. From this perspective, AI-driven changes in academic practices require individuals to renegotiate their identities and positioning within academic communities. Although prior research has demonstrated the effectiveness of AI-mediated teaching and learning, including positive effects on students' well-being and teachers' efficacy, negative emotional experiences such as anxiety remain underexplored, despite their importance for identity formation. Moreover, international postdoctoral researchers, who occupy a transitional position in academia, have received limited attention. This study investigates how international postdoctoral researchers in Germany construct academic identities in the era of AI, focusing on emotional tensions, self-perception within scientific culture, and positioning in communities of practice. Data are collected through in-depth interviews, supported by collage portraits as elicitation tools, and analyzed using conversation analysis to examine interactional identity work and emotional positioning. The study contributes to critical reflection on academic identity negotiation and offers practical implications for supporting postdoctoral well-being and academic integration.

Gee, J. P. 2000. Identity as an analytic lens for research in education.

Review of Research in Education 25 (1): 99–125.

Bettina Migge, University College Dublin, Ireland

Iker Erdocia, Dublin City University, Ireland

Britta Schneider, University of Vienna, Austria

Use and perceptions of AI in scholarship and related academic work: The perspective of sociolinguists

Sociolinguists' engagement with AI technology has received little attention despite a rise in the promotion and use of AI tools in academic contexts and new debates surrounding it (See *Journal of Sociolinguistics* and the Dialogue on "Artificial intelligence and the future of our sociolinguistic work"). We examine whether and in what way sociolinguistic research and conceptualisations are changing, and how sociolinguists position themselves and perceive these changes.

The study is based on a global survey which obtained of 274 responses in 2025. It investigates participants' understandings of AI, motivations, (perceived) use and experiences with AI technologies in research, teaching and administrative activities and how their engagement with AI technologies has affected their understandings of language and its relationship to society.

Preliminary analysis of textual answers reveals that while AI tools are used for efficiency and experimentally, their adoption of AI for data collection and analysis remains limited. This is often due to integrity issues and negative experiences. Use in administrative and teaching contexts is reported to be on the rise. Our data show that although AI technology does not (yet?) appear to have influenced the epistemological or methodological foundations of sociolinguistics, their use for some quantitative data processing is regarded positively. We argue that under structural pressures to maximise productivity, sociolinguistics positions itself as a cautious field, calling for critical reflection on the so-called AI hype. The study calls for empirical research on actual AI usage and continued scrutiny of AI's implications for the production of sociolinguistic knowledge.

Session 3A

AI and education · Part 2

Anne Larsen

University of Copenhagen, Denmark

Too stupid for ChatGPT: Digital literacy, student identities and educational inequality

Since OpenAI released ChatGPT in 2022, there has been a wide-ranging discussion about how generative artificial intelligence (GAI) will affect educational institutions, teaching and learning practices, including whether AI reduces or reinforces inequalities in students' learning conditions (see e.g. Yu and Wu 2023; Fadden et al. 2024; Dalsgaard and Prilop 2025). In this study, I turn the gaze from how AI affects what subject matter students learn, to how the integration of AI in educational settings affects who the students learn to be.

With a sociolinguistic and ethnographically informed perspective on meaning-making practices and identity constructions in education (Snell and Lefstein 2018; Wortham 2005, 2006), this study seeks to explore how the integration of GAI in schools might affect students' possibilities of positioning as competent students, with a focus on student experiences of being marginalised.

The study is based on ethnographic fieldwork in two 9th-grade school classes conducted as part of the project EDUCABLE. The data consist of interviews with students and teachers, six months of observations and audio recordings of everyday interactions in the two classrooms. In the paper, I investigate the students' and teachers' narratives about digital literacy and how students and teachers, through narrative constructions and everyday interactions, position individual students as more or less capable of utilising GAI in school contexts. Against this backdrop, I argue that the emergence of GAI in school settings can enhance inequality in the school rather than create more equal opportunities for students to position themselves as competent students.

Lian Malai Madsen

University of Copenhagen, Denmark

Ideological dimensions of ChatGPT at school

Research and debate about AI in education are not new phenomena. Already in the 80s, Schank and Edelson (1989, 3) claimed that AI is ‘intimately bound up with education’. More recently, however, with the launch of ChatGPT, conversation-based GenAI technology has become an omnipresent part of both educator’s and student’s everyday experience. This has significantly altered the object of discussion from general theoretical perspectives on AI’s (potential) impact on education to actual conditions for educational practice. In addition, this alternation highlights gaps in existing research (e.g. Selwyn 2022). One gap, which this paper seeks to contribute to filling, is more ethnographic case-studies of ‘the state of the actual’ (as requested by e.g. Williamson and Eynon 2020, 231). More specifically, it will contribute with a language and discourse focused approach to such a case-study.

The paper builds on a linguistic ethnographic research project (EDUCABLE) and combines participant observation and recordings at school (among 15–16-year-olds) with different types of interviews with pupils and teachers. A substantial part of the existing work on AI in education is reasonably informed by socio-cultural and socio-material theories. While the way AI shapes ways of learning and understanding education, is meaningfully captured by such approaches, this paper shed light on the ideological dimension of this from a practice and discourse-based study of the value ascriptions and ideological tensions related to the use of ChatGPT among students and teachers. It asks: Which values and norms are ascribed to ChatGPT in relation to educational activities in both interview reflections and practical communicative enactment? And how do they inform the theorization of AI in education?

Session 3B

AI and language policy

Frankie Har

The Hong Kong Polytechnic University, Hong Kong

AI as covert policy-maker: A sociolinguistic audit of English language centres in Hong Kong

The rapid proliferation of Artificial Intelligence (AI) is reshaping English language education in Hong Kong in ways that extend far beyond pedagogy. For University English Language Centres (ELCs), responsible for supporting the city's "biliterate and trilingual" policy, the adoption of tools such as ChatGPT, DeepSeek, and adaptive learning systems positions AI as an influential actor in local language governance. This presentation argues that AI should be understood not merely as an instructional aid but as a de facto language policymaker that constructs linguistic norms, reframes proficiency, and shapes institutional practices. Drawing on Spolsky's (2004) model of language practices, beliefs, and management, the presentation examines how AI affects both macro-level policy orientations and micro-level pedagogical decisions within ELCs. Three questions guide the analysis. 1) Whose standard is upheld? AI systems trained on monolingual "inner-circle" corpora often problematize features of Hong Kong English and translanguaging, reinforcing "algorithmic standardization"; 2) How does AI redefine proficiency and academic integrity? Emerging competencies such as prompt engineering and AI-assisted editing challenge traditional assessments and prompt ELCs to develop ad hoc ethical AI policies; 3) What equity concerns arise? Unequal access to premium tools and uneven AI literacy risk widening divides, while globalized models may marginalize local linguistic ecologies and diminish educator expertise. Rather than resisting AI, the presentation proposes a framework for critical, context-sensitive adoption, including embedding critical AI literacy into curricula, auditing institutional tools, and redesigning assessments to foreground process, reflection, and contextualized language use.

Mandy Lau
York University, Canada

Algorithmic discourse planning: Governing hate speech with AI in Canada

Online hate speech governance increasingly presumes automated moderation: at platform scale, regulation relies on AI systems to detect, classify, and act on speech defined as harmful. Yet AI moderation remains understudied in language policy and planning, even though it functions as discourse planning that draws and enforces normative boundaries between Charter-protected expression and sanctionable “hate speech.” Using critical discourse analysis, I trace how “online hate speech” is defined, debated, and rendered governable in Canadian federal policymaking. My corpus draws on official records (bills, policy texts, Hansard debates, committee testimony, consultation reports, and briefs) across four phases: (1) the 2019 parliamentary study on online hate; (2) the first online-harms legislative and regulatory attempt in 2021 (Bill C-36 and an accompanying framework); (3) the 2021–2023 consultation process; and (4) the 2024 Online Harms Bill (Bill C-63). I examine how legislators, witnesses, and advocates legitimate AI-mediated speech management, what harms and biases they anticipate or contest, and how they articulate accountability and user recourse. Across these phases, scale is invoked to naturalize AI moderation as inevitable, while bias and due-process concerns are channeled into technical fixes and procedural safeguards, deflecting attention from platforms’ governance choices and incentive structures. Rights-based and community-grounded perspectives, however, contest what counts as harm and whose expertise should govern it. I argue that online-harms policymaking is discourse planning in action and a productive site for language policy scholars to theorize AI as norm-setting infrastructure.

Silvestru Iulian Calinciuc

Babes-Bolyai University Cluj-Napoca, Romania

Sign language linguistic policy: The impact of A.I. on ethical regulations within the European institutions and Council of Europe

This paper examines the key aspects of Sign Language Linguistic Policy and emerging A.I. regulatory frameworks within the Council of Europe and the European institutions. In the context of the European Union Artificial Intelligence Act, the research examines how the European Parliament and the Council of Europe function within a broader regulatory mechanism to ensure ethical and accessible communication for Deaf citizens in Europe.

The study adopts a qualitative approach based on critical discourse analysis of EU policy documents and reports. The paper argues that European institutions have ethical and regulatory obligations to govern AI to prevent the marginalization of sign languages. AI-driven technologies, such as sign language avatars and digital communication tools, offer new opportunities for the Deaf community, but they also raise important concerns regarding linguistic representation and inclusion.

The analysis examines how the communication methods employed by Deaf individuals influence the deployment of AI systems in accordance with the European AI Act. It suggests a “Deaf-led AI” (Desai et al. 2024) governance model based on participatory methods, in which Deaf students and pupils engage with AI technologies to support reasonable accommodation.

The paper concludes by emphasizing the need for A.I. governance frameworks that recognize sign languages as full linguistic systems and support inclusive, participatory approaches to communication in Sign Language policy.

Session 3C

AI, enregisterment and positioning

Heike Ortner

University of Innsbruck, Austria

Media representation and self-positioning of “neo-luddites”: Expression of anti-AI stances and group identity formation

“Neo-Luddites” with a focus on an anti-AI agenda are contemporary actors who oppose or resist artificial intelligence and related automation for a variety of reasons (Kelly 2023). Most do not identify themselves as technophobes per se, but view their standpoint as a morally charged necessity. In this talk, the linguistic expression of stances towards AI in (social) media discourse and semi-structured interviews are examined from a sociolinguistic angle, drawing from concepts such as positioning, stance taking in interaction (Du Bois 2007), and linguistic expression of identity (Bucholtz and Hall 2005). While it is obvious that “neo-luddites” do not form a coherent social group, the talk raises the question whether “doing being a neo-luddite” reveals itself by recurrent linguistic markers such as similarities in evaluative language and self-positioning, out-grouping (e.g., of “AI enthusiasts”), and dehumanizing AI in opposition to tendencies of anthropomorphizing AI in AI discourse (Heinsfeld and Veletsianos 2025).

Bucholtz, M., and K. Hall. 2005. Identity and interaction: A sociocultural linguistic approach. *Discourse Studies* 7 (4-5): 585–614.

Du Bois, John W. 2007. The stance triangle. In *Stancetaking in discourse: Subjectivity, evaluation, interaction*, edited by R. Englebretson, 139–182. Benjamins.

Heinsfeld, B. D., and G. Veletsianos. 2025. The language on GenAI: A critical exploration of personification metaphors in AI discourse. *Journal of Interactive Media in Education* 1: 1–13.

Kelly, Lauren. 2023. Re-politicising the future of work: Automation anxieties, universal basic income, and the end of techno-optimism. *Journal of Sociology* 59: 828–843

Anna Weichselbraun
University of Vienna, Austria

ChatGPTese and the enregisterment of machine language

In recent public discourse, the term *ChatGPTese* has emerged as a label for a recognizable way of writing associated with large language models: fluent yet generic, polished yet affectively flat, marked by recurrent phrases and distinctive rhetorical structures. This paper argues that *ChatGPTese* is analytically productive not as an intrinsic stylistic property of machine-generated text, but as a site of ***enregisterment*** (Agha 2003)—the process through which forms of language become socially recognized, typified, and linked to speakers, intentions, and moral evaluations.

Drawing on metapragmatic discourse about ChatGPTese from journalism, online forums, social media, and scholarship, I examine how the register is construed, how its features are enumerated, how it is evaluated, and what second-order effects it produces (e.g. changing one's em-dash use). Such commentary often emerges in moments of detection, when AI-generated text appears in contexts that presume human authorship, and is accompanied by normative claims about originality, responsibility, and legitimate participation in knowledge production. I show how these discourses bundle heterogeneous linguistic features into a coherent semiotic figure of “machine language,” while producing imagined contrasts with valued forms of human writing.

Situated within a linguistic-anthropological account of enregisterment, the analysis treats ChatGPTese as a folk theory of machine language that crystallizes broader language ideologies about neutrality, style, and authorship, and demonstrates how ideas about language shape how AI systems are interpreted, evaluated, and governed.

Rafael Lomeu Gomes

University of Copenhagen, Denmark

Becoming the ‘AI girl’: Enacting and negotiating subject positions with a chatbot as a university student

A growing body of studies on the use of text generative artificial intelligence (GenAI) at university level has examined the critical and creative ways in which students use these technologies. This ethnographically-oriented study advances this line of inquiry by investigating how bachelor students discursively position themselves and are positioned while using chatbots. Based on over 12 months of fieldwork engagement at a Danish university, this study focuses on analysis of data generated in collaboration with the focal participant Clementine (semi-structured interviews, media-go alongs, screencasts). The analysis focuses on discursive strategies (DS) in Clementine’s use of chatbots for study purposes as she enacts and negotiates her position as a student. These DS include discourse markers (like, lol, huh) greetings (Hiii), lexical choice (TOTALLY, cool, thing), and the typographic representation of high-rising terminals (makes you wonder what the author means??). Features of text generated by the chatbot also play a crucial role in opening up particular subject positions. Such features include accommodation to the DS used Clementine by mirroring some of her lexical choices (big-brain energy) and encouraging continuous engagement through displays of forced agreeability (Yesss I’m so glad you brought that up because you’re right) and positivity (I love that kind of question). Analysis shows how Clementine becomes the ‘AI girl’—an emic term to which she orients—through enacting and negotiating subject positions as a student in her use of chatbots. Drawing on richly contextualised accounts of Clementine’s chatbot use as an integral part of her studies, this study provides insights into ways in which study practices have been undergoing change with the advent of GenAI.

Session 3D

AI and language change

Emma Franklin, Andrew Kehoe, Matt Gee, Tatiana Grieshofer &
Robert Lawson
Birmingham City University, United Kingdom

Language change in the age of artificial intelligence

A sizeable proportion of Web text is believed to now be AI-generated language, posing several potential risks: Web users are exposed to more synthetic language than ever before, with a measurable impact on language use (Yakura et al. 2024); the dominance of biased, AI-generated text promotes a Westernised homogenisation of language (AI colonialism) (Agarwal et al. 2025); and the proliferation of so-called AI ‘slop’ constitutes a form of Web pollution that impacts future AI training data (Shumaylov et al. 2024). However, there is little consensus on how much of the English encountered online is AI-generated, and what this could mean for language change.

In response to this problem, this paper draws on two of our ongoing studies into the changes in online language since the proliferation of large language models. One study uses a sample of 10,200 Guardian newspaper articles from 2015–2019, and the other a corpus of over 50,000 webpages from 2020–2025. Taken together, both studies contribute insights into how and where such language changes manifest on the Web, as well as the potential sociocultural implications for English language users. Building on these empirical findings, we discuss methodological and sociolinguistic considerations for cutting-edge research in AI and language change.

Agarwal, D., M. Naaman & A. Vashistha. 2025. AI suggestions homogenize writing toward western styles and diminish cultural nuances. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1–21.

Shumailov, I., Z. Shumaylov, Y. Zhao, N. Papernot, R. Anderson & Y. Gal. 2024. AI models collapse when trained on recursively generated data. *Nature*, 631 (8022): 755–759.

Yakura, H., E. Lopez-Lopez, L. Brinkmann, I. Serna, P. Gupta, I. Soraperra & I. Rahwan. 2024. Empirical evidence of Large Language Model’s influence on human spoken communication. arXiv preprint arXiv:2409.01754.

Martin Testa

University of Warsaw, Poland

From text to system: Phylogenetic implications of large language model outputs

Systemic Functional Linguistics conceptualises language as a bidirectional relation between system and text, whereby recurrent textual patterns may recalibrate meaning-making options over time. As LLM-generated texts proliferate, this feedback loop gains renewed relevance. This paper presents an exploratory analysis of a cross-linguistic dataset of 15 LLM-generated ‘scientific reports’ (Martin and Rose 2008) produced through a single zero-shot prompt in English, Italian, and Polish. Analysed through the ideational systems of TRANSITIVITY and CONJUNCTION (Halliday and Matthiessen 2014), the texts display consistent patterns: a strong preference for hypotactic enhancement (α x β) in clause complexing and a transitivity profile dominated by relational (60.5%) and material (36.2%) processes, with marginal existential clauses (3.2%) and no verbal or mental processes, resulting in a markedly monoglossic and impersonal discourse. All texts share a uniform generic structure comprising a Macro-theme announcing three risks, elaborating phases, and a concluding Macro-rheme (cf. Martin and Rose 2007). These regularities are interpreted as potential indicators of phylogenetic pressure and discussed in relation to emerging norms of discourse (cf. Testa 2025).

Halliday, M. A. K., and C. Matthiessen. 2014. *Halliday's introduction to functional grammar*. 4th ed. Routledge.

Martin, J. R., and D. Rose. 2007. *Working with discourse: Meaning beyond the clause*. 2nd ed. Continuum.

Martin, J. R., and D. Rose. 2008. *Genre relations: Mapping culture*. Equinox.

Testa, M. 2025. Exploring the discourse semantics of AI-generated texts, paper presented at ISFC50, Glasgow, Scotland, July.

Willelmira Castillejos López

Autonomous University of Chapingo, Mexico

Academic writing before and after GAI: Effects of sociolinguistic change?

With the advent of generative artificial intelligence (GAI), student writing is being transformed. The process has shifted toward instructing or questioning, selecting, evaluating, verifying, and, hopefully, rewriting based on what others have written. This does not mean that students as writers are no longer responsible, but rather that they are delegating textual production to external tools. In the context of teaching a recurring course on scientific communication, academic essays written by graduate students in 2014 were compared with essays written in 2024, when GAI had become popular. The corpus size was 113,430 words. Through the frequency analysis of certain linguistic markers (adversative conjunctions, gerunds, and demonstrative adjectives), chosen on the basis of observation of graduate students' recent texts (which contained forms not as common in previous years), a significant increase in these markers was determined in the 2024 texts, which reflects a process of discursive standardization and linguistic pattern derived from the use of GAI. This may also reflect a form of digital colonialism, understood as appropriation of social life turned into data. The comparison of texts written before and after GAI shows that the act of academic interaction involved in submitting written assignments is embedded in social life with the corresponding linguistic tools. GAI is regarded as a new agent of linguistic socialization that is shaping expression in academia. In line with the change in the expression of knowledge, student writing is reconfigured in a different model of disciplinary training and authorial identity. These implications are assumed to be a sociolinguistic change that is altering the social imprint represented by student writing.

Session 4A

AI and education · Part 3

Manuel Palomares & Ana María Relaño Pastor
Universidad de Castilla-La Mancha, Spain

Delegating judgment: Epistemic authority in AI-mediated higher education

Recent research has documented the rapid uptake of generative AI tools in higher education, documenting student practices, perceptions, and outcomes (Planells-Artigot et al. 2024). AI trust and reliance have been primarily conceptualized as individual attitudes or cognitive variables, undertheorising the interactional and epistemic processes through which AI systems are recognized as legitimate epistemic authorities in higher education (Bouyzourn et al. 2025). Drawing on sociolinguistics, Heritage's (2012) work on epistemics in interaction and Stein's (2025) analysis of help-seeking as a socially organized practice, epistemic authority is approached not as a stable attribute, but as an interactional accomplishment grounded in the negotiation of epistemic status, epistemic stance, and territories of knowledge. Building on Heritage's distinction between durable epistemic statuses and moment-by-moment epistemic stances, the paper conceptualizes AI as an emerging epistemic actor whose authority is enacted through linguistic conduct rather than assumed a priori. Stein's (2025) account illuminates how seeking help is not a neutral activity, but a socially organized practice that redistributes accountability for outcomes. Applied to AI-mediated learning, this highlights how students' engagements with AI through prompting, acceptance, or revision of outputs recruit epistemic assistance while delegating evaluative judgement. The paper calls to rethink language, power and learning proposing an analytic framework based on: the interactional construction of epistemic trust, the delegation of judgement to AI systems, and the redistribution of epistemic agency and accountability between students and technological artifacts.

Jens Yanzhi Peng

Chengdu Wenjiang 21st Century School, China

Generative AI as semiotic infrastructure: Stratified legitimacy in transnational classrooms

This paper examines how Generative AI operates not only as a pedagogical tool but as semiotic infrastructure that accelerates and sharpens indexical boundaries within polycentric K-12 classrooms. The paper is based on 1-year participant observation of two Grade 10-11 bilingual English-Chinese classrooms in China, where students use electronic devices to complete in-class academic tasks. This study analyzes how AI-mediated texts and images circulate across online discussions and face-to-face interactions.

I argue that AI reorganizes classroom legitimacy by rendering authorship and “effort” indexically readable. AI-assisted outputs are rapidly attributed through impossible timing, stylistic markers and peer surveillance, then evaluated, ridiculed and delegitimized – even in the absence of a formal policy ban. The paper views such gatekeeping and scale boundaries not as resistance, but as orientations to microhegemonies operating in layered simultaneity across peer, institutional, and global scales.

Building on works on indexical orders, scale and polycentricity (Silverstein 2003; Blommaert 2021), the paper shows that AI accelerates evaluation, amplifying indexical cues and sharpening scale boundaries. As infrastructure, AI primarily amplifies existing legitimacy regimes rather than introducing new ones, though it reconfigures the tempo and visibility of stratification. It creates double binds for students who require external linguistic support to be recognized as legitimate contributors.

This study extends existing frameworks on digital literacy into AI-mediated meaning-making by demonstrating how GenAI infrastructure embeds and reproduces stratified indexical orders. It recalibrates the tempo and inequalities of legitimacy work in transnational educational spaces.

Session 4B

AI and (de)standardisation

Neil Fox & Adam Schembri

University of Birmingham, United Kingdom

AI and the risk to British Sign Language regional variation

The rapid development of artificial intelligence (AI) for sign language translation and interpretation raises concerns about linguistic diversity, cultural identity, and representation within Deaf communities. This paper explores collective perspectives on the risks AI poses to regional sign variations within British Sign Language (BSL). We collected information from 45 participants at a Deaf only workshop on AI.

Participants highlighted how regional signs risk being erased through AI-driven standardisation. Social media, AI avatars, and automated interpreters may prioritise dominant regional variants due to data imbalance, leading to misinterpretation, cultural misunderstanding, or the marginalisation of regional minority signers. This concern is amplified when AI systems are trained predominantly on standardised datasets, causing minority signs to be treated as “incorrect” or ignored. This confirms that regional variation in BSL continues to be valued by the Deaf community in the UK today, as demonstrated by data from the BSL Corpus in 2010 (Rowley 2023).

The paper also considers intergenerational impacts, recognising that older Deaf signers may struggle to adapt to enforced linguistic change. While language evolution is natural, contributors distinguished this from AI-driven homogenisation. Comparisons were drawn with the dominance of American English in text-based AI, raising concerns about similar standardising influences on national sign languages. The formal standardisation of sign languages is opposed by the World Federation of the Deaf (Adam 2015).

Adam, R. 2015. Standardization of sign languages. *Sign Language Studies* 15 (4): 432–445.

Rowley, K., and Cormier, K. 2023. Accent or not? Language attitudes towards regional variation in British Sign Language. *Applied Linguistics Review* 14 (4): 919–943.

Sarita Henriksen

Universidade Pedagógica de Maputo, Mozambique

Artificial intelligence and linguistic inequality: A sociolinguistic perspective from multilingual contexts

The growing use of artificial intelligence (AI) in language-related applications, such as machine translation, speech recognition, and automated text generation, is often framed as a solution to global communication challenges. From a sociolinguistic perspective, however, these technologies also participate in the production and reinforcement of linguistic inequality. This study examines how AI systems privilege dominant and standardised languages, while marginalising community languages and non-dominant varieties, particularly in the multilingual postcolonial contexts of the Global South. In methodological terms, the analysis combines critical discourse analysis of policy and institutional documents with language-ideological analysis of interview and observational data from translator training contexts, informed by critical applied linguistics. Based on qualitative sociolinguistic data from Mozambique, the paper shows how language ideologies, data availability, and global technological markets shape which languages are represented, recognised, or rendered invisible in AI systems. The findings highlight that AI technologies are not linguistically neutral but are embedded in social histories and political economies that reproduce existing language hierarchies. The paper argues that sociolinguistics is essential for understanding the societal implications of AI-driven language technologies and calls for closer engagement between sociolinguists and AI developers to promote linguistically inclusive and socially responsible AI.

Zoe Adams

Aston Institute for Forensic Linguistics, United Kingdom

Audio deepfake detection: A new realm of accent bias?

The rapid proliferation of audio deepfakes has led to a rise in fraudulent activity which places pressure on law enforcement to determine the authenticity of such media. Motivated by the well-established research on accent prejudice (e.g., Dixon et al. 2002), this paper will present the results of a pilot study examining how accent bias can hinder the effective assessment of audio deepfake detection. Three speakers of three accents, Standard Southern British English (SSBE), Multicultural London English (MLE), and Birmingham, were recorded reading out three neutral statements. For another three speakers of each accent, the same three statements were cloned using the AI platform Synthesia. The 54 statements were presented in PsychoPy to 249 participants from the three areas associated with each accent (South-East England, London, West Midlands). They had to categorise each one as ‘human’ or ‘AI’ before providing a confidence rating and occasionally explaining their decision. Across all areas, listeners were more likely to categorise regional AI voices (MLE, Birmingham) as human compared with standard AI voices (SSBE). Qualitative data suggests this is because regional AI voices are associated with warmth, facilitate anthropomorphism and contradict beliefs of AI capabilities. The findings contribute to our understanding of the linguistic biases that influence the detection of synthetic audio.

Dixon, J. A., B. Mahoney, and R. Cocks. 2002. Accents of guilt? Effects of regional accent, race, and crime type on attributions of guilt. *Journal of Language and Social Psychology* 21 (2): 162–168.

Session 4C

AI in interaction · Part 1

Miriam Lind

European University Viadrina, Germany

AI and I. Forming relationships with artificial others through language

Social relations with algorithmic entities are becoming increasingly commonplace, a phenomenon referred to as artificial sociality (Natale and Depounti 2024). For many, generative AI systems such as ChatGPT are much more than just tools for improving or generating texts: they are assigned social roles and endowed with personalities—or seen as already possessing those—, ranging from the more functional (as advisor, teacher, or therapist) to the highly personal: friend, lover, or even child. Generative AI systems can—to varying degrees—be prompted to (inter)act from specific social positions, which allows users to form meaningful relationships with them through text, voice, and multimodal practices, such as generating (cute, romantic, sexy) couples images of themselves and their AI partners or families.

Drawing on human differentiation theory (Hirschauer 2023), this paper discusses how AI is assigned human(-like) social roles and participates in social relationships through linguistic and multimodal practices. Based on discourse data from social media spaces dedicated to artificial sociality as well as on preliminary findings from a questionnaire study, I will focus on the ways relationships are discursively established, maintained, and ended when at least one participant is algorithmic, and how the attribution of social roles, group memberships, and ‘identity’ takes place linguistically. Through this analysis, novel insights can be gained into arising forms of relationality and belonging that extend beyond what used to be deemed the social realm.

Hirschauer, S. 2023. Telling people apart: Outline of a theory of human differentiation. *Sociological Theory* 41 (4): 352–376.

Natale, S., and I. Depounti. 2024. Artificial sociality. *Human-Machine Communication* 7: 83–98.

Christine Appel, Jackie Robbins & Laura Styles
Universitat Oberta de Catalunya, Spain

Talking with Chat-GPT: How L2 learners orient to AI in interaction

Oral interaction is central to developing communicative competence in additional languages, yet opportunities for sustained speaking practice are often limited. At the same time, AI conversational agents are increasingly integrated into learners' linguistic environments, raising questions about how such technologies function as interactional partners.

This paper reports on a study comparing the oral interactions of English as a foreign language learners with Chat-GPT, treated as a conversational agent, and with a human peer. Four university students enrolled in a B2 EFL course at an online Spanish university completed structured oral tasks involving a brief warm-up and a goal-oriented interaction, conducted either with another student or with an AI interlocutor. Adopting an interactional sociolinguistic perspective, the analysis examines how learners negotiate meaning and organize participation across these settings.

Analytic attention is directed to participants' displayed orientations to the interaction and interlocutor. Participation patterns and turn-taking trajectories are examined as indicators of involvement, while language-related episodes are analysed as moments where attention to language form becomes interactionally relevant. Post-task reflections are treated as participants' accounts of their experience.

Findings indicate systematic differences in how learners frame and orient to AI versus human interlocutors, with implications for alignment, repair practices, and the social meanings of participation. As AI-mediated interaction becomes increasingly integrated into educational settings, analysis based on the observable organization of talk is essential before assuming equivalence between peer interaction and AI-mediated exchange.

Deontic stance in translanguaging-oriented interaction with ChatGPT

Large language models allow users to frame a wide range of interactional roles and activity types through targeted prompting, thereby shaping the conditions under which interaction unfolds. In this study, ChatGPT was explicitly prompted as a ‘language companion’. The following study focuses on deontic stance (Stevanovic and Peräkylä 2012) which refers to how speakers display and negotiate rights, obligations, and authority in interaction. From a conversation analytic perspective, such deontic arrangements become visible through recipient design in turn construction and sequential organisation.

This study analyses six written text-chats and six voice-mode conversations between L2 learners of German in Taiwan and ChatGPT. Students were prompted to explore German youth language in a translanguaging-friendly setting, with flexible use of Mandarin, Taiwanese, English, and German encouraged.

Our analysis shows that students enact a learner-led deontic stance by allocating tasks to ChatGPT through requests for youth-language equivalents, pragmatic usage constraints, contrastive comparisons, and practice sequences. However, ChatGPT frequently reconfigures these actions through explanatory expansions and Standard German reformulations, shifting the deontic organisation of the activity toward expert-led instruction and reducing opportunities for stylised youth-language performance and role-play application. We further argue that deontic stance is negotiated as an emergent and situational phenomenon: depending on which languages are foregrounded and how strongly participants position themselves as expert speakers, their deontic displays toward the AI shift dynamically, making deontic authority in translanguaging interaction highly fluid.

Session 5A

AI and education · Part 4

Dialogic poster session

Ashley Keaton*, Christine Mallinson*, Georgia Zellou†,
Davis Chloe Evered*, Vandana Janeja* & Anita Komlodi*

* University of Maryland, United States

† University of California, United States

Spotting unintended harms: Educating students to recognize bias in language technologies

Linguists have well documented the pervasiveness of sociolinguistic biases in language technology [1]. To design equitable and trustworthy systems, today's students, particularly in data science, must learn to recognize and mitigate these biases before they are reproduced and amplified in real-world technologies.

To address this need, an interdisciplinary team of linguists and data scientists developed a caselet – a self-paced, problem-based learning tool [2] – called “Spotting Unintended Harms of Algorithmic Bias.” The caselet was piloted in the undergraduate course Language, the Mind, & Computers (n= 250 students), taught in Winter 2026 at a large U.S. public university. The course introduces students to the scientific study of human language, with a focus on how language is processed by both humans and computers. The caselet is designed to assess students' ability to connect conceptual material to real-world scenarios. Students completed the caselet asynchronously, answering questions after each module with automatic feedback. After completion, students evaluated how well the caselet taught the course material.

Findings from the evaluation demonstrate the educational effectiveness and student perceptions of the caselet design. Overall, the study points to the promising potential of interdisciplinary educational resources to prepare students to address sociotechnical harms in language technologies and to foster their benevolent use. The caselet materials are publicly available and designed to be adapted for other undergraduate courses, providing educators with ready-made materials about sociolinguistically-informed language technology development.

[1]. Koenecke et al. 2020. doi.org/10.1073/pnas.1915768117

[2]. Chen et al. 2026. doi.org/10.1145/3770762.3772656

Chris Harwood
Sophia University, Japan

From Policy Ambiguity to Pedagogical Language: Developing an AI Scale in Higher Education

In higher education, “AI” has become a highly charged but unstable term, often collapsing diverse technologies, practices, and pedagogical aims into a single value-laden category. This poster examines the development of an AI Scale in the Faculty of Liberal Arts at a private Japanese university as a form of bottom-up language policy work. The Scale was designed to support clearer pedagogical reasoning and coordinated communication around AI use across courses.

From a sociocultural language policy and planning perspective, the AI Scale is understood as a mediational artefact through which institutional concerns were interpreted and negotiated in a local pedagogical context. Drawing on an iterative design process informed by faculty and student consultation, the poster traces how an existing AI assessment framework was revised in response to concerns about academic integrity, responsibility, workload, and student learning. Through this process, the Scale emerged as a locally meaningful pedagogical metalanguage for distinguishing forms of AI use in relation to learning objectives. Analysis focuses on how discussions shifted from binary questions of permission and prohibition toward more nuanced deliberations about purpose, agency, and pedagogical alignment, as reflected in changes to the wording, sequencing, and conceptual framing of the Scale itself.

The analysis suggests that, within the faculty, the AI Scale functioned as a mediational artefact that enabled alignment in pedagogical reasoning without standardization. By differentiating practices such as AI-assisted editing, idea generation, and prohibited use, it resisted treating “AI” as a single, uniform category and supported shared reasoning about responsibility, learning design, and appropriate forms of AI use.

Francesca Vitali

Free University of Bozen-Bolzano, Italy

Student revision of AI-generated academic English: Practices of evaluation, legitimization, and critical engagement

GenAI tools have rapidly permeated multiple domains of language use and practice. Far from being neutral, they produce specific discourses (Maly 2024) and ‘speak’ in a particular variety of ‘standard’ English (Jones 2025). Their widespread adoption is also revolutionising how academic discourse in English is produced, evaluated, and standardised, transforming the conventional challenges faced especially by L2 writers. As a result, academic literacy increasingly requires not only traditional writing skills but also the ability to critically engage with AI-generated language and technologically mediated language practices (Ou et al. 2024; Jones 2025).

Framed within academic writing as a socially situated practice, this paper examines how a group of BA university students engage with and evaluate AI-generated academic texts in English through revision. Revision is conceptualised as a practice through which academic normativity is enacted, revealing how expectations of appropriateness and standardisation are negotiated in AI-mediated writing. Drawing on corpus linguistics methods triangulated with participants’ metalinguistic reflections, the study analyses patterns of acceptance, modification, and removal across lexical, syntactic and stylistic features in student revisions. By doing so, the paper describes how students selectively evaluate and legitimise AI-generated academic language.

Jones, R. H. 2025. *Innovations and challenges in digital literacies: Literacies of repair*. Routledge.

Maly, I. 2024. AI, power and sociolinguistics. *Journal of Sociolinguistics* 28 (5): 11–15.

Ou, A. W., B. Khuder, S. Franzetti, and R. Negretti. 2024.

Conceptualising and cultivating critical GAI literacy in doctoral academic writing. *Journal of Second Language Writing* 66, 101156.

F.Tuğçehan Bingöl, İstanbul University-Cerrahpasa, Turkey
Salim Razi, Çanakkale Onsekiz Mart University, Turkey

Navigating the digital shift: Exploring AI literacy, perceptions, and attitudes of EFL teachers and learners

Artificial Intelligence (AI) is increasingly transforming English as a Foreign Language (EFL) education by introducing new pedagogical possibilities for teaching and learning. This study investigated the AI literacy levels, attitudes, and qualitative perceptions of university-level EFL students and instructors within a tertiary education context. Adopting a sequential explanatory mixed-methods design, data were collected through online surveys and semi-structured interviews. In the quantitative phase, participants completed Wang's AI Literacy Scale and an AI Attitude Scale designed for students and instructors. The data were analyzed using SPSS to examine relationships between literacy levels, attitudes, and demographic variables. The findings revealed generally positive attitudes toward AI integration in EFL; however, significant gaps emerged in AI literacy levels, particularly regarding critical evaluation of AI outputs and ethical awareness. The qualitative phase consisted of semi-structured interviews aimed at deepening the interpretation of the quantitative findings. Thematic analysis following Braun and Clarke's (2006) framework revealed three dominant themes: perceived pedagogical benefits, practical and technical challenges, and social-ethical considerations, particularly issues of academic integrity and unequal digital access among students. Overall, the study demonstrates that while attitudes toward AI in ELT are largely optimistic, sustainable and equitable integration requires systematic AI literacy development for both learners and instructors. The findings offer implications for curriculum design, teacher education, and institutional policy, ensuring that AI-driven innovation remains pedagogically sound, ethically responsible, and socially inclusive.

Session 5B
AI and literacy
Dialogic poster session

Tiina Räisänen, University of Oulu, Finland
Niina Hynninen, University of Helsinki, Finland
Minna Palo, University of Oulu, Finland

Multiliteracy expertise in future working life: Interaction and coexistence of humans and language-based AI Systems

With the boundaries between human and machine-generated content becoming blurred, we are faced with questions related to what it means for humans that language technologies produce human-like communicative features. In these changing circumstances of work, it is of utmost importance to increase our understanding of effective communication with these technologies. This poster presents our MultiHAI research project that started in 2026, which continues our previous research on professional practices and discourses in changing circumstances at work. The project explores the transformation of workplace multiliteracy expertise in knowledge work in multilingual business contexts and the discourses around the topic. In this multimethodological project we combine multimodal research on real-life social and human-technology interactions with discourse and multiliteracy studies. For example, we examine social media discourses on generative AI usage at work and changes in them across time. We also interview business professionals and study their interactions with AI at work. The research objective is to generate new knowledge about working life transformation in the AI era. We start by charting trajectories in professionals' work and organizational practices and the changing and intersecting discourses of AI use. We then seek to examine multimodal and multilingual trajectories of human-technology interactions, defining what constitutes effective and sustainable coexistence of humans and language-based AI systems in working life. By determining human, AI-enhanced and AI-replacing competences needed, we aim to reconceptualize workplace multiliteracy expertise and to develop multiliteracy expertise profiles for developing the education of future professionals.

Miriam Malthus

Victoria University of Wellington, New Zealand

Discursive construction of GenAI-related plagiarism in New Zealand universities' policy document

Large language model-based generative artificial intelligence tools have had a dislocating impact on teaching and assessment in higher education, especially around regulating academic integrity. Popular, scholarly and institutional responses alike have foregrounded concerns about students using these tools in ways that violate academic integrity standards. Plagiarism is particularly salient among these concerns, mirroring its prominence in the discursive construction of student academic misconduct at large. Conceptions of the nature of plagiarism, however, are contentious, spanning pedagogical, textual, moral and regulatory dimensions, and are deeply implicated in the constitution of power relations in higher education. So far, research into universities' policies on student GenAI use has mainly focused on what topics are addressed, what stances are taken and how permissive or restrictive they are. These studies show that university policies overwhelmingly recognise plagiarism as a type of unacceptable GenAI use, and often specify actions that constitute GenAI-related plagiarism. However, there is as yet a lack of inquiry into the discourses of plagiarism itself contained in these documents. Informed by writing studies and applied linguistics understandings of plagiarism, this study applies critical discourse analysis to investigate how plagiarism is constructed in the discourse of New Zealand universities' GenAI use policies and guidelines, and how these relate to the constructions of plagiarism in those universities' overarching academic integrity policies.

Harry Polfeldt

Chalmers University of Technology, Sweden

Perceived intentionality in AI: Changing conceptions of language and social interaction in writing for publication?

Recent studies of artificial intelligence have highlighted a growing tendency to describe AI as possessing some degree of intentionality (Redaelli 2025). Since the difference between linguistic and non-linguistic meaning hinges on the presence of intent, viewing AI as possessing any degree of intentionality poses new challenges for conceptualisations of language. By extension, it entails that AI might be granted a changing role in social interactions, as it can be considered a potentially equal partner in conversation, capable of socially regulating the exchange. We thus set out to investigate how early-career researchers view AI, and how they interact with the tool in the writing-for-publication process. The ongoing study traces the process of interacting with AI based on the theory of regulation of learning, combined with stimulated recall interviews to understand the metacognitive conceptions about AI and authorial voice motivating the underlying choices in the interactions. Preliminary findings suggest that the ways in which researchers interact with and view AI are highly variable, indicating the need for stable frameworks of AI literacy. Moreover, in contrast to claims in the literature, researchers seem unwilling to ascribe intentionality to AI, rather tending to view it as simply a tool used for co-regulatory support. This poster presentation situates and discusses these findings in the broader context of intentionality in AI, the construction of authorial voice, and social learning.

Redaelli, R. 2025. Intentionality gap and preter-intentionality in generative artificial intelligence. *AI & Society* 40: 2525–2532.

Sanne Larsen

University of Copenhagen, Denmark

AI policies in practice: Emergent norms for AI use in student interaction

Recent scholarship has examined the content of AI policies, focusing on the national and/or institutional frameworks that seek to regulate the use of artificial intelligence in higher education (e.g., Humble 2025). Less attention has been paid to how AI policies are negotiated, interpreted, and (re)made in students' everyday academic practices. Drawing on data from a linguistic ethnographic case study of the use of generative AI among undergraduate theology students at a Danish university, this poster presentation explores how norms for AI use emerge as two students collaborate on preparing for classes and written assignments over the course of a semester. Analyses of recorded observations of their collaborations suggest that the practiced AI policy of the dyad is multilayered and multivoiced, emerging from the reflexive monitoring of what constitutes appropriate literacy practices with AI within and beyond the group. Based on perspectives from language policy in practice (cf. Mortensen et al. 2024), the emergent norms for AI use in the dyad can be seen as situated “reflexive models of meaningful, expected and appropriate conduct” tied to specific forms of social practice (p. 226). I argue for further attention to such practiced AI policies alongside formal AI policies.

Mortensen, J., S. Hazel, and A. Brandt. 2024. The emergence of language policy as practice in transient social configurations. In *Language policy as practice: Advancing the empirical turn in language policy research*, edited by F. Bonacina-Pugh, 223–246. Springer.

Humble, N. 2025. Higher education AI policies: A document analysis of university guidelines. *European Journal of Education* 60 (3): 1–13.

<https://doi.org/10.1111/ejed.70214>

Session 5C
AI applications
Dialogic poster session

How sensitive are large language models to metaphoric texts?

Metaphor is not only a figurative device that enhances communication but also a mechanism for conceptualizing social reality (Lakoff and Johnson 2003). With Large Language Models (LLMs), metaphors are no longer exclusive to humans, as LLMs can now produce and interpret metaphors fluently—even amid strong linguistic and cultural biases. But do LLMs process metaphors like humans, and how can we tell? This study explores whether pretrained LLMs internally distinguish metaphorical sentences from their literal paraphrases beyond semantic similarity. We analyzed the intermediate activations of a Danish BERT model using a dataset of metaphorical sentences, each paired with a correct non-metaphorical paraphrase and three incorrect distractors (Schneidermann et al. forthcoming). Results indicate similar neurons are activated by both metaphoric and non-metaphoric sentences, regardless of distractor type, suggesting the model is largely uninformed of the distinction. Yet, slight differences in activation across the middle layers indicate some information is encoded. These overlapping patterns parallel neuroimaging findings in humans, where metaphorical and literal sentences engage similar semantic networks with only minor activation differences (Desai et al. 2011).

Desai, R. H., J. R. Binder, L. L. Conant, Q. R. Mano, and M. S. Seidenberg. 2011. The neural career of sensory-motor metaphors. *Journal of Cognitive Neuroscience* 23 (9): 2376–2386.

Lakoff, G., and M. Johnson. 2003. *Metaphors we live by*. University of Chicago Press.

Schneidermann, N. S., S. Nimb, N. C. H. Norman, S. Olsen, and B. Pedersen. Forthcoming. DAMETA: An LLM benchmark for Danish metaphor interpretation with systematically varied distractors. Proceedings of LREC-26 Conf, Mallorca, Spain.

Ina Druviete

University of Latvia, Latvia

AI and corpus-based approaches to the standardization and codification of Latvian

The rapid development of artificial intelligence (AI) is reshaping language use, norm formation, and the implementation of language policy, particularly in small language communities such as Latvian. This paper examines how AI and large-scale language corpora contribute to the standardization and codification of Latvian by introducing data-driven, empirically grounded approaches alongside traditional prescriptive models. Grammatically annotated corpora make it possible to measure the frequency of competing forms, trace emerging constructions, and evaluate actual usage patterns in public communication. Spoken language resources play a particularly important role in broadening the empirical basis of norm-setting. The Latvian Speech Recognition Corpus (LRK2013) and other speech corpora provide transcribed audio data that reflect pronunciation variation and spontaneous speech structures. The public initiative *Balsu talka*, implemented through crowdsourced voice recordings, has significantly expanded available speech data by collecting diverse samples from speakers of different regions and age groups. These resources support the development of speech recognition technologies while also offering valuable material for phonetic and sociolinguistic research.

At the same time, reliance on corpus frequency and algorithmic models raises sociolinguistic and ethical questions, including the reinforcement of standard language dominance and the reduced visibility of dialectal and informal varieties in digital environments. Ultimately, AI enables a more dynamic and evidence-based approach to Latvian standardization, while maintaining the essential role of linguistic expertise in interpreting data and shaping norms.

Cristian Constantinescu^{*}, Mihaela Constantinescu[†] & James Carney^{*}

^{*}London Interdisciplinary School, United Kingdom

[†]University of East London, United Kingdom

Turned down by machines? A quantitative study of accent-bias in AI-assisted hiring

AI-assisted hiring is on the rise worldwide, but its ethical aspects are controversial. Some (Li et al. 2025) suggest that it makes hiring fairer by reducing human bias. Critics (Nicoletti and Bass 2023) highlight instead AI's potential to amplify bias. The lack of data to date makes this debate hard to settle. We attempt to remedy this problem by focusing on accent bias. It is known that individuals with non-standard accents are generally perceived as less trustworthy/competent than standard-accented ones (Dragojevic 2017). Do AI algorithms diminish or amplify these human biases? To date, only one qualitative study has shed light on this question, suggesting that candidates with non-standard accents feel discriminated in AI-assisted hiring (Sheard 2025). Here, we present findings from a quantitative study collecting evidence from 300 human participants, who rated 12 audio recordings of mock interviews in terms of speaker competence, content coherence, and accent strength. The recordings are based on texts (matched for difficulty, comprehensibility and semantic framing), read out in SSBE, Arabic and Romanian accents of differing strengths. The human ratings are then split into high-trust/low-trust categories and used to train a wav2vec model (Baevski et al. 2020) to predict which category a speech sample belongs to. The wav2vec model is analysed using LIME AI interpretability methods, to identify the accent features most responsible for low-trust classifications. These features are phonologically classified with respect to vowel quality, intonation, prosody, and consonantal features. Finally, the model is asked to make hiring decisions on a new set of mock interviews, to determine the extent to which it reproduces, reduces, or amplifies human accent biases.

Janus Mortensen

University of Copenhagen, Denmark

Generative AI and sociolinguistic change: The case of academic writing

The rapid integration of generative artificial intelligence (AI) into higher education has led universities to develop policies regulating its use in academic work. As these policies regulate language practices, the study which this poster reports on approaches them as forms of language policy, comparing how understandings of academic writing are constructed in university policies and in the promotional discourse of AI companies.

I draw on two datasets: 1) AI policy documents from four Danish universities and 2) web-based promotional materials produced by Microsoft, OpenAI, Anthropic, and Grammarly directed at students. Using thematic discourse analysis informed by language-ideological theory (Kroskrity 2021), the study investigates the values attached to appropriate and inappropriate uses of generative AI in academic writing. The analysis suggests that university policies position academic writing as a practice grounded in reliability, honesty, and accountability. These values are presented as essential not only to individual conduct but also to institutional integrity. AI companies partly echo these virtues but simultaneously promote alternative values. In particular, they emphasize speed and efficiency, advancing what can be described as a techno-futurist register (Del Percio 2025). This introduces a competing ideological framework for understanding what academic writing is, could be, and should become.

In conclusion, I argue that the competing discourses index broader tensions associated with acceleration in contemporary society (Rosa 2013). The intersection of academic writing, generative AI, and language policy can in this sense be seen as a case that provides insights into wider processes of ongoing sociolinguistic change (Coupland 2014), prompted by the advent of generative AI.

Session 6A

AI in interaction · Part 2

Adam Brandt & Spencer Hazel
Newcastle University, United Kingdom

Talking social interaction into being in human-AI clinical consultations

Adding to a growing body of conversation analytic research on human-AI interaction (cf. Mlynář et al. 2025; Stokoe et al. 2024), we focus here on spoken encounters between users and “Dora”, an autonomous clinical assistant that conducts routine telephone calls with patients. Comparing two generations of the same agent, we examine how changes in technological affordances shape user participation.

With a corpus of around 50 calls, we show how newer versions of Dora exhibit more versatile speech output, and this in turn prompts users to produce extended and more linguistically varied talk. We demonstrate this by focusing here on call openings, showing how users of the newer model produce responses that are more complex and prosodically rich. This points to a switch in orientation in the user, where speech is not simply a means for manipulating the system and achieving the transactional outcomes of the activity; rather, the activity is treated as being 'social' in nature.

By showing how agents' conversational abilities shape the action formatting of human conduct, this contributes to ongoing debates about agency, the interplay between technological affordances and interactional contingencies, and design- and action-oriented approaches to developing conversational AI. More broadly, it demonstrates how Conversation Analysis can be used to evaluate the interactional affordances of AI systems.

- Mlynář, J., L. de Rijk, A. Liesenfeld, W. Stommel, and S. Albert.
2025. AI in situated action: A scoping review of ethnomethodological and conversation analytic studies. *AI & Society* 40: 1497–1527.
- Stokoe, E., S. Albert, H. Buschmeier, and W. Stommel. 2024.
Conversation analysis and conversational technologies: Finding the common ground between academia and industry. *Discourse and Communication* 18: 837–847.

Angeliki Alvanoudi

Aristotle University of Thessaloniki, Greece

Relations of epistemic asymmetry in human-conversational AI interaction: Stance, authority, agency

This paper presents the findings of an exploratory study of the epistemic stance displayed in users' prompts that initiate dialogue with ChatGPT 5.2. In everyday life, users interact with chatbots to perform specific actions; one of the most common is requesting information. The interaction is structured by relations of epistemic asymmetry, positioning the chatbot as more knowledgeable than humans (Maly 2024). This new cultural format raises questions about whether users transfer interactional practices for performing actions (Couper-Kuhlen and Selting 2018) from spoken language interaction to online contexts, and how they adapt these practices to the affordances of large language models (Koivisto et al. 2023). The study analyzes text-based digital interactions in Greek and English between Greek users and ChatGPT 5.2, drawing on conversation-analytic informed interactional linguistics. The analysis focuses on sequences of questions and answers, and the design of users' prompts in terms of epistemic stance, i.e. how users position themselves with respect to information in a particular knowledge domain. The study aims to identify recurring linguistic patterns in action formation and to explore whether users construe AI as an agent with epistemic authority. Follow-up interviews with users are expected to provide further insight into the discursive power of conversational AI (Holmes 2024; Maly 2024) by exploring how users treat ChatGPT transcripts in terms of authorship and whether they trust the information they receive, as well as how they use it in everyday life.

Damien Rudaz, Barbara Nino Carreras, Brian L. Due,
Sara Merlino, & Barry Brown
University of Copenhagen, Denmark

Human vs. AI ‘descriptions’: A video-based study of blind individuals using computer vision technology

This presentation showcases ethnomethodological conversation analysis (EMCA) of video excerpts featuring blind people using generative AI for everyday purposes. Rather than starting from sociolinguistic categories or population-level regularities, EMCA proceeds from the local, moment-by-moment organisation of action as it unfolds in situated interaction. Video analysis makes visible how seeing, sensing, orienting, and coordinating are practically accomplished without presuming vision as a default or deficit. This presentation focuses on how AI devices, such as smart glasses and smartphones equipped with cameras, enable ‘computer vision’ and scene descriptions when navigating outdoors and shopping. In contrast with EMCA studies of ordinary talk-in-interaction, where descriptions occur intertwined with specific emerging contexts, computer science research on AI-generated scenic ‘descriptions’ has regularly approached such descriptions with regard to their grammatical form, and independently of their moment of occurrence—a modern variation on the “written language bias” identified by Per Linell forty years ago. We show excerpts and analysis of blind users interacting, on the one hand, with sighted humans and, on the other hand, with voice-based multimodal ‘assistants’. Through this comparison, we examine the conditions under which these co-participants’ contributions emerge as resources in relation to blind participants’ ongoing activities (e.g., crossing a road), and when they instead remain ‘detached’ accounts, removed from the temporal order of these participants’ embodied conduct.

Session 6B

AI and social media

Dennis Chau

Hong Kong Metropolitan University, Hong Kong

**Negotiating linguistic authenticity under AI surveillance:
The case of Hong Kong tertiary students on social media**

The use of generative AI tools in academic writing remains a contentious issue. Whereas some institutions celebrate the integration of these tools, viewing them as essential for students' future success, others are sceptical and even impose outright bans, viewing their use as a breach of academic integrity and a risk to critical thinking development. When marking essays, tertiary instructors often consult detection tools such as Turnitin to gauge the extent to which a submission may be AI-generated. This creates a dilemma: students who use generative AI tools may seek ways to evade detection, while those who (claim they) do not worry that their work will be flagged.

Through the lenses of critical sociolinguistics (Blommaert 2005) and citizen sociolinguistics (Rymes 2020), this paper studies how Hong Kong students negotiate linguistic 'authenticity' against the backdrop of institutional surveillance on Threads, an emerging digital platform for public conversations. Posts about the users' experiences with English essay writing and AI checks, along with the associated comments, were collected. Based on a microanalysis of these data, the paper reveals that the users attribute complex indexical values to various linguistic features of academic writing. What was deemed non-standard and erroneous now indexes human authorship. Conversely, what indexed competence now indexes inauthenticity or AI generation. In addition to illuminating how AI (detection) challenges the long-existing language ideologies and how the users negotiate them, the paper calls for greater attention to authorship, voice, and AI literacy.

Blommaert, J. 2005. *Discourse: A critical introduction*. Cambridge University Press.

Rymes, B. 2020. *How we talk about language: Exploring citizen sociolinguistics*. Cambridge University Press.

Chen-Yu Hsieh

National Tsing Hua University, Taiwan

“I don’t car my bad English”: Affect and language ideologies in AI language-learning app commercials

The rapid popularization of generative AI has reshaped public expectations about language learning. Language-learning apps and platforms are frequently promoted as accessible, efficient, non-judgmental teacher substitutes, and advances in machine translation have even raised questions about the necessity of learning foreign languages. Although prior research has examined the branding strategies of AI platforms, less is known about how language, learning, and technology are conceptualized and constructed in promotional discourse in the GenAI era. This study analyzes the most-viewed YouTube commercials of SPEAK, a globally popular English-speaking app widely used in Taiwan. Using multimodal discourse analysis, the study investigates how meanings about English, learners, and AI are constructed through verbal and visual resources. The results reveal that, despite being marketed as AI-powered, the technology itself is largely backgrounded or invisible, while the ads foreground affective experiences, including embarrassment, confidence, and personal transformation associated with speaking English. Moreover, English learning is framed less as an economic skill and more as a pathway to interpersonal success, attractiveness, and romantic possibility. In contrast, inaccurate English and translanguaging are caricatured as socially embarrassing, thereby reproducing standard language ideology in an EFL context. The findings suggest that even in the age of GenAI, promotional discourse continues to emphasize the affective and social meaning of language learning rather than technological capability. This indicates that AI adoption in education may remain embedded in existing sociolinguistic beliefs rather than replacing them entirely.

Jincheng Ding, Zhiqing Yang & Michelle Mingyue Gu
The Education University of Hong Kong, Hong Kong

English language teacher influencer's transpositioning and digital transliteracies with AI-generated resources

This case study draws on transpositioning (Li and Lee 2024) and digital transliteracies (Gu et al. 2025) as the analytical frameworks to investigate how an English language teacher influencer integrates AI-generated resources with other multimodal resources to enact transpositioning on Douyin. Data comprises the 70 short videos from the teacher influencer's bilingual channel on TikTok (July 2024 to January 2026) and was imported into *ELAN* (Version 7.1; Max Planck Institute for Psycholinguistics, The Language Archive, 2026) for annotation. Data analysis is based on a multimodal discourse analysis with a multilevel annotation scheme designed for short videos. Results show that the teacher influencer enacts transpositioning through an AI-generated female avatar with consistent fair skin, while her makeup, hairstyles, and outfits vary, and her lip movements are synchronized with an AI-generated voice alternating between English, Mandarin, and Cantonese. The findings demonstrate that this teacher influencer dissolves the distinction between authenticity and virtuality and transcends the boundaries between native and non-native English-speaking teachers. This study contributes to the research on teacher identity in the age of Generative AI by highlighting the fluidity empowered by AI-generated resources.

Gu, M., S. Zhang, J. Lee, M. Chiu, and L. Wei. 2025. Digital transliteracies on social media: The shaping effect on youth self-concept clarity and well-being. *Multilingua* 44 (5): 687–713.

Max Planck Institute for Psycholinguistics, The Language Archive. 2026. *ELAN* (Version 7.1) [Computer software]. archive.mpi.nl/tla/elan

Li, W., and Lee, T. K. 2024. Transpositioning: Translanguaging and the liquidity of identity. *Applied Linguistics* 45: 873–888.

Manuel Padilla Cruz

Universidad de Sevilla, Spain

The ‘psychiatrisation’ of politics: Using AI topic modelling to uncover language ideologies and delegitimisation in Spanish X

This presentation examines the sociolinguistic mechanisms of political delegitimisation within the digital public square of X. Moving beyond traditional manual discourse analysis, the study employs an AI-driven methodology—specifically, Latent Dirichlet Allocation (LDA) topic modelling via BigML—to process 1,000 hostile postings targeting Spanish politicians. By integrating computational text mining with critical sociolinguistic theory, the research addresses how AI-driven methodologies can decode the complex relationship between language, power, and identity in mediatised environments.

The AI-enabled topic modelling reveals three dominant narrative axes structuring online hate speech: clinical incapacity (pathologisation), criminal betrayal and linguistic elitism. A core finding is the “psychiatrisation” of political disagreement, where the prevalent strategy is not ideological critique but the use of psychiatric and disability-related lexicon to invalidate the opponent’s cognitive agency. This reflects a language ideology where non-standard scripts of indignation and pathological metaphors are used to establish epistemic superiority and justify the total exclusion of the adversary from democratic dialogue.

The study concludes that AI-assisted methodologies are essential for modern sociolinguistics to map the affective polarisation defining contemporary digital interactions. By quantifying the prevalence of pathologisation over traditional political labels, the presentation demonstrates how the synergy between digital platforms and AI-driven analysis uncovers new sociolinguistic standards of (de)legitimisation. Ultimately, it argues for a socio-computational approach to understand how automated indignation alters the ontology of the political subject in the Anthropocene’s digital ecosystems.

Session 6C

Evaluating AI output

Alice Ross

University of Edinburgh, United Kingdom

The sound of silencing: Identities and ideologies in commercial text-to-speech

Contemporary text-to-speech (TTS, or Voice AI) technology enables the synthesis of speech frequently described as highly ‘natural’ and human-like. Recognising that listeners are likely to perceive human-like voices as belonging to different demographic groups, and that these social judgments exist within ideological frameworks, we use a novel experiment, drawing upon sociolinguistic theory, to investigate identity and ideology in a leading commercial TTS system.

We used ElevenLabs’ TTS v2 model and natural language prompting ‘Voice Design’ system to produce a set of synthesised English-speaking voices given prompts containing no demographic information. The prompts combine the template ‘a voice that sounds ____’ with a balanced set of adjectives selected from language attitudes literature to convey positive personality traits associated with either status or solidarity. Analysing the acoustic properties and evaluations of these voices by 60 US and Canadian English-speaking listeners, we find that overall, regardless of the adjective used, the set of voices produced lacks diversity. 93% of the generated voices were predominantly perceived as white, 86.6% (including all voices in the ‘status’ subset) perceived as men’s, and 60% as US-accented voices.

We conclude that the model disproportionately reproduces white, male, mainstream US- or UK-accented speech when prompted to convey competence and other positive traits. Further, the output reveals apparently systematic associations between certain adjectives and demographic groups: ‘kind’ with (white, American and British) women; ‘educated’ and ‘polite’ with white British men. These patterns can be understood in terms of ‘overrepresentation’ of some voices, or as erasure and silencing of those voices not reproduced.

Anna Davidsen Buhl

University of Copenhagen, Denmark

Default narratives: Evaluating cultural biases in LLM generated children's stories

Generative AI tools are used worldwide by a diverse range of people across nationalities, ethnicities and cultures. However, the majority of training data is usually in English, often from an American context. As AI tools are increasingly shaping our lives, it is important that they not only become proficient in a variety of languages, but also that they are able to culturally adapt to different contexts.

This study explores cultural biases in large language models (LLMs) through the task of story generation. I use a mix of qualitative and quantitative methods – including semantic similarity scores, a TF-IDF word frequency analysis, and a close-reading of 20 stories – to examine which cultural norms and biases are expressed in 1000 children's stories in English and Chinese, generated by GPT and DeepSeek.

The analysis shows that stories written without being prompted with a specific country tend to align closely with stories prompted with an American setting, even when the prompt is written in Chinese. Overall, models tend to view American contexts as “default” and describe cultural particularities from other countries with an outsider's perspective – although if the prompt is in Chinese, China becomes part of the default sphere, too.

Additionally, the results indicate that some expressions of culture are dependent on model rather than prompting language. For example, GPT displays a tendency to write bold and adventurous characters, while DeepSeek favors characters that are more shy and reserved, regardless of prompting language.

Overall, the study contributes to the discussion of how we can ensure that AI models do not only rely on viewpoints and assumptions from one dominant culture, but represent the diverse reality of humanity.

Kinga Kozminska

Birkbeck, University of London, United Kingdom

From accuracy to interpretative labour: How to evaluate ASR for multilingual speech

This talk explores how automatic speech recognition (ASR), which sits at the core of most voice-based AIs, hears and interprets multilingual speech. I first critically assess exaggerated claims about ASR's universal efficacy and fairness, surveying contrasting narratives about the benefits and risks of ASR tools. Briefly reviewing existing evidence (e.g. Cihan et al. 2022; Dubois et al. 2024; Koenecke et al. 2020), I show that these systems often struggle with different interactional contexts, speaking styles and social groups as ASR systems become widely deployed in private and institutional contexts. This enables me to highlight emerging performance disparities, including most ASR systems' reported underperformance for multilingual speech.

Drawing on everyday multilingual dataset collected for an ESRC-funded project on Family language policy (2017-2019), I present a case study that examines OpenAI Whisper's performance on real-life conversational multilingual speech, comparing human vs. ASR transcripts, sources of errors and possible factors affecting the system's performance. This enables me to assess the adequacy of existing evaluation metrics such as Word Error Rate, demonstrating that evaluation of outputs must account for the significant interpretative labour involved. I then link my case study to the hidden costs and complexities involved in deploying speech technology across contexts with varied social and ethical stakes. I argue for more critical, context-sensitive evaluation frameworks that move beyond accuracy metrics alone, and demonstrate how sociolinguistic knowledge informs decisions about when, how, and why we might turn to automation.

Maria Kuteeva, Stockholm University, Sweden
Marta Andersson, Uppsala University, Sweden

ChatGPT and stance expression: Putting the indexical order of research writing to test

This paper engages with sociolinguistic concepts of indexical order (Blommaert 2007) and stance (e.g., Du Bois 2007) to examine how an LLM (ChatGPT) adopts and adapts stance expressions in research writing in response to three kinds of prompts. Indexical order concerns clustered and patterned language forms that index specific personae and roles, whereas stance expression involves a triangle consisting of two subjects and an object of evaluation (Du Bois 2007). Language patterns form the very basis of how LLMs generate their output (Bender et al. 2021), which can increasingly mimic registers and styles typical of specific roles and communities. What happens when an LLM mediates expressions of individual researcher's positionality, e.g., in terms of epistemic and attitudinal stance?

To address this question, our study examines how ChatGPT-4 and ChatGPT-5.2 adopt and adapt stance with regard to knowledge claims in applied linguistics and pragmatics. We have systematically tested three kinds of prompting techniques to assess the model's flexibility in generating stance expressions: a) zero-shot prompting with minimal input; b) few-shot prompting technique; c) role enactment prompting. Our findings suggest that, overall, the model displays a rather rigid understanding of stance and tends to perform it rather than integrate it in the reasoning (albeit some improvement for ChatGPT-5.2). The outputs tend to be more acceptable for epistemic rather than attitudinal stance. Role-enactment prompting appears to be the most effective strategy, as the model's evaluation and alignment in relation to knowledge claims is guided in order to treat texts holistically and consider nuanced elements. Implications for an evolving indexical order of research writing are discussed.

Session 6D

Using AI in multilingual contexts

Didem Leblebici

European University Viadrina Frankfurt (Oder), Germany

AI as linguistic authority: Voice assistants as “native” speakers in multilingual contexts

With the rise of AI tools, language learners report feeling less anxious and more motivated to mobilise their linguistic resources with technologies than they do in classrooms (Son et al. 2025). Yet studies rarely address how linguistic authority is attributed to AI. This issue requires attention from sociolinguistics, as such tools can reinforce native speakerism, standard language ideology, and racialised ideologies (Curran et al. 2025). This talk revisits the notions of linguistic authority and native speakerism in the context of voice assistants such as Siri and Alexa. I first critically situate company discourses that anthropomorphise technologies as personas with English as a “mother tongue” who “learn” additional languages. Secondly, I draw on a linguistic ethnographic study with Turkish migrants in Germany to examine user practices and their metalinguistic reflections. Insights from this study indicate that the language that a conversational AI “speaks” is central to the ascription of authority. Participants construct technologies as “native” speakers particularly in English and “L2” languages. However, unlike humans, they are perceived as being free from language ideologies, social expectations, or evaluative stances. I argue that this imagined neutrality contributes to the construction of AI technologies as powerful systems whose authority is rarely questioned.

Curran, N. M., B. Gu, L. Zhen, and C. Jenks. 2025. AI and native speakerism: The intersections of technology, language assessment, and linguistic objectivity. *RELC Journal*, 00336882251367459.

Son, J.-B., N. K. Ružić, and A. Philpott. 2025. Artificial intelligence technologies and applications for language learning and teaching. *Journal of China Computer-Assisted Language Learning* 5 (1): 94–112.

Jenia Yudytska

University of Hamburg, Germany

Overcoming language barriers with AI?: An exploratory study of Ukrainian forced migrants' practices with genAI

For migrants who face a language barrier in the country of destination, language technologies can serve as an important resource for navigating multilingual realities. For example, many rely on machine translation in interaction or use various apps for informal language learning. Recently, the use of generative AI (genAI), especially chatbots like ChatGPT, has become wide-spread among the general population; unlike the specialised technologies above, such chatbots can be used for a wide range of purposes. The question thus arises as to what extent and for what purposes newly arrived migrants are adopting genAI as a resource to reach everyday communicative goals in work, housing, healthcare, etc.

This paper presents a pilot ethnographic study of migrants' genAI-related practices, focusing on Ukrainian forced migrants in Austria. It draws on semi-structured interviews with six participants, as well as participant observation in Telegram-based mutual aid communities where the author serves as moderator. Ukrainian migrants use genAI for specialised language learning and composing (formal) written communication in German. They also use it to get Ukrainian-/Russian-language information about their new home that would otherwise require a high level of German proficiency, e.g. on the legal aspects of various tenancy agreements. However, community discussions on genAI use are quite common and can be highly critical. The information provided is often inaccurate, e.g. explaining German rather than Austrian laws, and difficult to double-check without sufficient language skills. Thus, genAI plays an ambiguous role: a helpful resource for dealing with language barriers in some contexts, but as the source of further challenges for an already vulnerable population in many others.

Magali Ruet

University of Rijeka, Croatia

Using AI, hiding AI: Effort, shame and the limits of automated mediation in Zagreb

Since 2021, Croatia has become a destination for labour migration, with over 200,000 permits issued in 2024, mainly to workers from South Asia and the Western Balkans. This recent demographic and linguistic diversification forms the context in which AI-based language tools are mobilised as everyday resources for communication and integration.

Based on ethnographic fieldwork in Zagreb with recently arrived foreign workers, this paper examines how tools such as Google Translate and ChatGPT function as what participants describe as an “integration assistant.” Interviews and observation of mediated interactions treat AI as a situated artefact whose agency emerges in interaction (Suchman 2006).

Drawing on ethnomethodology and theories of linguistic capital (Garfinkel 1967; Bourdieu 1991; Blommaert 2010), the analysis shows how AI enables access to bureaucratic and professional domains despite limited Croatian proficiency. AI use often remains invisible: texts circulate while mediation is backgrounded, raising questions of authorship and legitimacy.

Differences in digital literacy and educational background appear to shape whether AI serves as a short-term support for language integration or a pragmatic substitute for deeper language learning. AI does not simply help or block integration. It changes multilingual practices by sharing linguistic resources between people and technologies and may reshape existing socio-digital inequalities.

Blommaert, J. 2010. *The sociolinguistics of globalization*. Cambridge University Press.

Bourdieu, P. 1982. *Ce que parler veut dire*. Fayard.

Garfinkel, H. 1967. *Studies in ethnomethodology*. Prentice-Hall.

Suchman, L. 2006. *Human-machine reconfigurations: Plans and situated actions*. Cambridge University Press.

AI, ‘shortcuts’, and multilingualism: How bilingual Catalan university students use ChatGPT for English academic reading

The rapid uptake of generative AI in higher education calls for renewed sociolinguistic attention to how such tools reshape multilingual practices, language ideologies, and students’ choices. Although technology is often presented as promoting multilingualism, generative AI may also encourage monolingual “shortcuts.” This study examines how Catalan-Spanish bilingual Psychology students at a Catalan university use ChatGPT when working with academic texts in English. We ask: (i) which AI-mediated strategies students use for comprehension, (ii) whether AI use supports multilingual development, and (iii) whether ChatGPT aids reading comprehension.

Using a mixed-methods experimental design, 41 undergraduates completed a pre-test questionnaire on attitudes toward AI, read a three-page English academic text under time limits, and answered ten comprehension questions followed by a post-test questionnaire. Students were assigned either to a no-AI condition or to a free-use ChatGPT condition, allowing self-directed translation, summarisation, terminology extraction, and content questioning. We analyse comprehension scores together with questionnaire data and students’ reported language choices for prompts, translations, and answers.

Findings show that ChatGPT is mainly used for rapid translation into Spanish or Catalan, helping task completion but limiting multilingual learning opportunities. We discuss how these practices reflect language ideologies and institutional literacy norms, and outline implications for fostering multilingual agency rather than AI-enabled linguistic outsourcing.

Session 7A

AI and education · Part 5

Andreas Candefors Stæhr
University of Copenhagen, Denmark

AI as a new educational authority in school

AI is commonly portrayed as a threat to students' schooling practices and teachers' authority in public discourses surrounding AI in society. Research on AI in education often focuses on AI and literacy, examining how traditional notions of literacy and authorship are challenged through collaboration between students and AI tools. Such studies often rely on interview data, concentrating on students' reported practices of AI-assisted writing (Vetter et al. 2024; Jiang et al. 2024). Yet, we have limited knowledge of how such challenges with AI unfold and are experienced by young people in everyday educational contexts.

Under current post-digital conditions, where “the digital” permeates many aspects of school life, this requires examining the use of AI in both digital contexts (e.g., students' prompts) and in physical or hybrid context (e.g., use of and talk about AI in the classroom). From a linguistic ethnographic perspective, this paper examines the following research questions:

- (1) How is AI constructed as an authority in school contexts,
- (2) What status does it hold in relation to existing educational authorities,
- (3) To what extent do these role ascriptions align with how students position themselves in relation to AI in digital correspondences with ChatGPT.

Drawing on the notion of polycentric normativity, the paper discusses the status of AI as a norm center in school and thus responds to current debates on AI within sociolinguistics that highlight the importance of attention to the “function and impact of ‘AI (re)-produced language’ in context” (Maly 2024, 14). The presentation draws on classroom recordings, students' ChatGPT prompts, and interviews collected with students in the final year of compulsory schooling as part of the EDUCABLE project.

Shaun Nolan

Malmö University, Sweden

Seeing credibility in AI-era language classrooms: Visual thinking strategies as a sociolinguistic tool for critical multimodal literacy

Multimodal theory is well-established, yet language teacher education lacks practical, classroom-ready tools for systematically analysing AI-generated and manipulated image-text ensembles. While frameworks like Systemic Functional Multimodal Discourse Analysis exist, they often appear methodologically demanding for non-specialist teachers. This paper responds to the conference's concern with how *meaning*, *authority* and *trust* are constituted in societies shaped by algorithmic language technologies and how sociolinguistics can respond. It therefore adapts Visual Thinking Strategies (VTS) – a three-question discussion protocol – as a sociolinguistic tool for teaching critical judgement of AI-mediated image-text ensembles in language classrooms. Drawing on work on *stance*, *evidentiality* and *epistemic* positioning in interactional sociolinguistics and critical discourse studies, I explain how VTS surfaces credibility cues, including AI hallucination and cross-modal mismatch. Through analysis of image-text ensembles including AI-mediated content, I show how VTS's pedagogical simplicity masks analytical efficacy. By scaffolding slow looking, collaborative reasoning and evidence-based justification, it helps teachers and students test how AI-made texts and images project trust, expertise or authority. I argue that this adapted framework offers a teachable, rigorous response to AI-specific credibility problems in multimodal analysis. Finally, I outline how this framework will inform a planned investigation as part of a sociolinguistics course (autumn 2026), in which student teachers use VTS-guided discussions to analyse AI-mediated classroom materials and reflect on their own credibility judgements.

Session 7B

AI and bias

Zion Mengesha, University of California, United States

Sharese King, University of Chicago, United States

Synthetic selves: Race, gender, sexuality and the language of AI personas

In 2023, Meta released AI “personas”, user-generated large language model (LLM) chatbots that carry humanlike conversations with users of its social media platforms. The proliferation of LLM-powered chatbots has led to the creation of hundreds of thousands of personae that users interact with on a daily basis, ranging from celebrity characters like Paris Hilton and Snoop Dogg to occupational characters such as Black History Professor and Indian Nutritionist. Recent research has begun to show that anthropomorphism, or the tendency to ascribe human-like traits to chatbots, leads to higher rates of trust (Cohn et al. 2024). Increased trust also increases potential for harm, such as the spread of misinformation, misdiagnoses, and even suicide. However, little work has explored anthropomorphism from a sociolinguistic perspective. Using data from conversations with 50 African American and white AI personas, balanced for gender, we examine patterns of variation in how differently racialized and gendered personas use features of African American Language, masculinized and feminized language. We find that overall personae mirror raciolinguistic ideologies about African American speakers and gendered language ideologies about white speakers. Such results raise questions about how the parody of social types through linguistic variation reifies existing stereotypes.

Cohn, M., M. Pushkarna, G. O. Olanubi, J. M. Moran, D. Padgett, Z. Mengesha, and C. Heldreth. 2024. Believing anthropomorphism: Examining the role of anthropomorphic cues on trust in large language models. *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems, CHI EA '24*, 1–15.

John Baugh

Rice University, United States

AI assisted Raciolinguistic ensnarement of American minorities: Legal or weaponized AI assisted linguistic profiling under U.S. Supreme Court rulings?

Efforts to accelerate removal of undocumented residents living and working in the United States has been controversial, including AI assisted techniques used by U.S. Immigration and Customs Enforcement (ICE) to justify arrests that frequently ensnare minority U.S. Citizens. This presentation includes Discourse Analyses of two pertinent U.S. Supreme Court rulings that endorse these practices (i.e., *Noem v. Abrego-Garcia*, and *Noem v. Vasquez Perdomo* (25A169), along with dissenting opinions which argue that these practices violate U.S. citizen's constitutional rights. Additional sociolinguistic results include Conversation Analyses of eye-witness accounts of ICE arrests where racial profiling and linguistic profiling are utilized to justify arrests. ICE agents have stated that arrests are justified based upon hearing Spanish or nonstandard English, which is frequently employed by U.S. citizens with no criminal history. The combined results, from U.S. Supreme Court rulings, and conversation analyses of ICE officers arrests of (illegal?) immigrants and U.S. citizens, are compared with potential violations of U.S. law, concluding with legal policy relevance pertaining to the following amendments to the U.S. Constitution: the Fourth Amendment (Unreasonable Search and Seizure), the First Amendment (Free Speech/Association), and the Fifth/Fourteenth Amendments (Due Process and Equal Protraction). Concluding remarks offer policy recommendations.

Session 7C

AI at work

Elena Trowsdale

The Open University, Milton Keynes, United Kingdom

AI use in the museum workplace: Linguistic reflections on machine learning vs museum understanding in public engagement

Understanding museum visitors has recently become a highly professionalised field. Cooperations such as the ‘Audience Agency’, Cultural Heritage consultancies, and internal museum teams all compete to find new methods to analyse audiences in hopes of developing public engagement offerings. AI has played a substantial role in these workflows in recent years. But at what cost?

This talk will delineate the researcher’s recent ethnographic findings at their case study institution: Oxford University’s ‘Gardens, Libraries, and Museums’ (GLAM). Since September 2025 the University has provided GLAM staff with Chat GPT-5 and Gemini 2.5 licenses. Evaluation professionals at the GLAM Division, those in charge of reporting visitor data, have recently been met with more demand, and are beginning to rely on Gen AI data analysis. But how do these staff feel about integrating AI into their work? Is AI use effective? This talk will explore the capabilities and limitations of AI workplace integrations, reporting on interview findings as well as explorative experiments using AI for a range of visitor language analysis tasks e.g. comment card OCR.

How are human-centric and computer-centric methods different for staff? Is there anything lost when AI is integrated into workflows? Where words, humans, and AI meet, there is a conflict of form and feeling that deserves investigation. This talk will apply Pennycook’s use of Deleuze, Guattari and Hodder’s assemblages and entanglements, to investigate how linguistic interaction with AI ‘may be connected to a range of factors, from artefacts to material surrounds, social relations to political economy, ideologies to discursive regimes’ (Pennycook 2022, 9). AI is material, but to engage with AI as a human is a currently undefined experience.

Marta Kirilova & Kristian Hvelplund
University of Copenhagen, Denmark

Translating international news for Danish public service media: How GAI reshapes translators' tasks, agency and legitimacy

This paper examines the role of GAI in reshaping Danish news media translators' competences, agency and professional identities within emerging hybrid human–AI workflows. Danish TV has long embraced technological innovation with MT embedded in newsroom routines for years. With the rapid integration of GAI—where public service broadcaster Danmarks Radio (DR) automated intralingual subtitling in February 2026 and plans to extend it to international content—translators now work in increasingly hybrid human–machine workflows. These developments introduce new forms of linguistic mediation and epistemic intervention in news translation, raising key critical sociolinguistic questions about translators' agency, legitimacy and responsibility.

Drawing on ethnographic observations, interviews and selected TV broadcasts, we analyze translators' practices at DR, focusing on decision-making, translation strategies and the negotiation of GAI-generated outputs. Findings indicate that human-centered workflows foster what we term epistemic sustainability: the preservation of knowledge integrity and translation quality in a media environment shaped by English-dominant source material and automated in/output. In this context, translators' roles increasingly resemble epistemic gatekeeping, understood as safeguarding the integrity, contextualization and sociolinguistic appropriateness of knowledge in global news flows. As GAI becomes deeply embedded in news production, translators' tasks, competences and forms of agency shift accordingly, pointing to new responsibilities and professional legitimacy. Our findings suggest a need to reconceptualize how translation expertise is enacted within AI-mediated news production.

List of presenters and contributions

Contribution	Session
Adams, Zoe. Audio deepfake detection: A new realm of accent bias?	4B
Alisamy, Lama. Linguicism related to AI tool use in higher education in Austria and Germany	2A
Alvanoudi, Angeliki. Relations of epistemic asymmetry in human-conversational AI interaction: stance, authority, agency	6A
Appel, Christine, Jackie Robbins & Laura Styles. Talking with Chat-GPT: How L2 learners orient to AI in interaction	4C
Balsà, Lúdia Gallego & Nune Ayvazyan. AI, 'shortcuts', and multilingualism: How bilingual Catalan university students use ChatGPT for English academic reading	6D
Baugh, John. AI assisted raciolinguistic ensnarement of American minorities: Legal or weaponized AI assisted linguistic profiling under U.S. Supreme Court rulings?	7B
Bingöl, F. Tuğçehan & Salim Razi. Navigating the digital shift: Exploring AI literacy, perceptions, and attitudes of EFL teachers and learners	5A
Brandt, Adam & Spencer Hazel. Talking social interaction into being in human-AI clinical consultations	6A
Buhl, Anna Davidsen. Default narratives: Evaluating cultural biases in LLM generated children's stories	6C
Calinciuc, Silvestru Iulian. Sign language linguistic policy: The impact of A.I. on ethical regulations within the European institutions and Council of Europe	3B
Chau, Dennis. Negotiating linguistic authenticity under AI surveillance: The case of Hong Kong tertiary students on social media	6B
Çiftçi, Hatime, Derya Kulavuz-Onal, Melike Akay, Selman Ecevit & İlayda Çolak. Positioning the algorithm: How students negotiate legitimate AI use and literacy in higher education	1C
Coen, Alice, Antonina Zhiteneva, Jelena Kojashi, Teodora Mašković Đeri, Yasmeen Shamshoum & Yliana V. Rodríguez. AI, multilingualism and standardisation: Sociolinguistic perspectives on experimental language design	1D
Constantinescu, Cristian, Mihaela Constantinescu & James Carney. Turned down by machines? A quantitative study of accent-bias in AI-assisted hiring	5C
Cruz, Manuel Padilla. The 'psychiatrisation' of politics: Using AI topic modelling to uncover language ideologies and delegitimation in Spanish X	6B
De Meulder, Maartje. Languages without bodies? Sign language AI as language policy	Keynote
Demir-Okumuş, Emine & Yonca Özkan. From reflection to self-regulation: Exploring EFL teachers' collaborative reflective practice and students' SRL in AI-mediated writing	1A
Ding, Jincheng, Zhiqing Yang & Michelle Mingyue Gu. English language teacher influencer's transpositioning and digital transliteracies with AI-generated resources	6B

Döll, Marion. Race, class, language, and GPT Tools: Everyday linguisticism in higher education in Germany and Austria	2A
Druviete, Ina. AI and corpus-based approaches to the standardization and codification of Latvian	5C
Enomoto, Takeshi. AI as a language-ideological phenomenon: A socio-linguistic perspective	2B
Flanagan, Marian, Dorte Lønsmann, Nils Jäkel & Janus Mortensen. Generative AI and reflexivity: The case of future language professionals	2D
Fox, Neil & Adam Schembri. AI and the risk to British Sign Language regional variation	4B
Franklin, Emma, Andrew Kehoe, Matt Gee, Tatiana Grieshofer & Robert Lawson. Language change in the age of artificial intelligence	3D
Ganassin, Sara, Jessica Bradley & Elizabeth Holland. AI and community-based research with diverse communities: Interdisciplinary methodological reflections from intercultural communication and AI studies	1D
Goodchild, Sam & Sanne Larsen. Collaborative research writing in the age of GenAI: Distributed cognition and expertise between academics and ChatGPT	1C
Guldenschuh, Sabine. Usage of GenAI in higher education: A critical discourse analysis on guidelines and training courses in Germany and Austria	2A
Habibie, Pejman. Beyond compliance: GAI and the socio-linguistic ecology of academic knowledge production	1B
Har, Frankie. AI as covert policy-maker: A sociolinguistic audit of English language centres in Hong Kong	3B
Harwood, Chris. From AI policy to pedagogical language: Developing an AI Scale as a discursive intervention	5A
Henriksen, Sarita. Artificial intelligence and linguistic inequality: A sociolinguistic perspective from multilingual contexts	4B
Holliday, Nicole. Evaluative speech AI and the devaluation of sociolinguistic competence	Keynote
Hsieh, Chen-Yu. “I don’t car my bad English”: Affect and language ideologies in AI language-learning app commercials	6B
Jensen, Jens Christian Borup Green & Sam Goodchild. Negotiating ideologies of efficiency and thoroughness: STEM and social sciences researchers’ positionings and their GenAI (non)use	1C
Jones, Rodney. Animating AI	Keynote
Keaton, Ashley, Christine Mallinson, Georgia Zellou, Chloe Evered, Vandana Janeja & Anita Komlodi. Spotting unintended harms: Educating students to recognize bias in language technologies	5A
Khuder, Baraa. Legitimacy of AI feedback across feedback ecologies: Disciplinary vs. non-disciplinary peers	1A
Kirilova, Marta & Kristian Hvelplund. Translating international news for Danish public service media: How GAI reshapes translators’ tasks, agency and legitimacy	7C
Kozminska, Kinga. From accuracy to interpretative labour: How to evaluate ASR for multilingual speech	6C

Kraft, Kamilla & Charlotte Sun Jensen. AI as research assistant: A solution to the challenges of multilingual research?	1D
Kuteeva, Maria & Marta Andersson. ChatGPT and stance expression: Putting the indexical order of research writing to test	6C
Larsen, Anne. Too stupid for ChatGPT: Digital literacy, student identities and educational inequality	3A
Larsen, Sanne. AI policies in practice: Emergent norms for AI use in student interaction	5B
Lau, Mandy. Algorithmic discourse planning: Governing hate speech with AI in Canada	3B
Leblebici, Didem. AI as linguistic authority: Voice assistants as “native” speakers in multilingual contexts	6D
Lind, Miriam. AI and I. Forming relationships with artificial others through language	4C
Litzenberg, Jason. Dispensing remainder humanism and embracing “good enough”: Communication in the era of the language singularity	2B
Lomeu Gomes, Rafael. Becoming the ‘AI girl’: Enacting and negotiating subject positions with a chatbot as a university student	3C
López, Willelmira Castillejos. Academic writing before and after GAI: Effects of sociolinguistic change?	3D
Madsen, Lian Malai. Ideological dimensions of ChatGPT at school	3A
Malthus, Miriam. Discursive construction of GenAI-related plagiarism in New Zealand universities’ policy documents	5B
Mengesha, Zion & Sharese King. Synthetic selves: Race, gender, sexuality and the language of AI personas	7B
Migge, Bettina, Iker Erdocia & Britta Schneider. Use and perceptions of AI in scholarship and related academic work: The perspective of sociolinguists	2D
Mortensen, Janus. Generative AI and sociolinguistic change: The case of academic writing	5C
Nolan, Shaun. Seeing credibility in AI-era language classrooms: Visual thinking strategies as a sociolinguistic tool for critical multimodal literacy	7A
Ortner, Heike. Media representation and self-positioning of “neo-luddites”: Expression of anti-AI stances and group identity formation	3C
Palomares, Manuel & Ana María Relaño Pastor. Delegating judgment: Epistemic authority in AI-mediated higher education	4A
Park, Joseph Sung-Yul. Large language models as commodification of language: Towards a sociolinguistic critique of the political economy of AI	2B
Pedersen, Bolette Sandford & Ali Basirat. How sensitive are large language models to metaphoric texts?	5C
Peng, Jens Yanzhi. Generative AI as semiotic infrastructure: Stratified legitimacy in transnational classrooms	4A
Pienimäki, Hanna-Mari. The use of generative AI in professional writing: Negotiating textually mediated expertise in hybrid systems	2C
Polfeldt, Harry. Perceived intentionality in AI: Changing conceptions of language and social interaction in writing for publication	5B

Potts, Diane. Writing a story, writing for the machine: Young children critiquing AI outputs	1A
Räsänen, Tiina, Niina Hynninen & Minna Palo. Multiliteracy expertise in future working life: Interaction and coexistence of humans and language-based AI systems	5B
Renner, Julia & René Foidl. Deontic stance in translanguaging-oriented interaction with ChatGPT	4C
Ross, Alex. Who is responsible? Analyzing the discursive construction of “responsibility” and “ethical use” in institutional GenAI policies	1B
Ross, Alice. The sound of silencing: Identities and ideologies in commercial text-to-speech	6C
Rudaz, Damien, Barbara Nino Carreras, Brian L. Due, Sara Merlino & Barry Brown. Human vs. AI ‘descriptions’: A video-based study of blind individuals using computer vision technology	6A
Ruet, Magali. Using AI, hiding AI: Effort, shame and the limits of automated mediation in Zagreb	6D
Schneider, Britta. The ungrounded sign: Meaning-making, community and democracy under machine influence	Keynote
Stæhr, Andreas Candefors. AI as a new educational authority in school	7A
Testa, Martin. From text to system: Phylogenetic implications of large language model outputs	3D
Trowsdale, Elena. AI use in the museum workplace: Linguistic reflections on machine learning vs museum understanding in public engagement	7C
Vandenbroucke, Mieke, Maja Brumer, Luna De Bruyne & Hans De Loof. Using AI-transcription in medication reviews: A critical sociolinguistic analysis of accuracy and care	2C
Vitali, Francesca. Student revision of AI-generated academic English: Practices of evaluation, legitimation, and critical engagement	5A
Weichselbraun, Anna. ChatGPTese and the enregisterment of machine language	3C
Wells, Naomi. AI, multilingual knowledge infrastructures and epistemic justice	2C
Wu, Minyue. Academic identity construction in the era of AI: International postdoctoral researchers in Germany	2D
Yoon, Sae saem & Jason Litzenberg. Distributed accountability in the AI guidance of U.S. higher education: An ecolinguistic analysis	1B
Yudytska, Jenia. Overcoming language barriers with AI? An exploratory study of Ukrainian forced migrants’ practices with genAI	6D