

NALOMA 2026

**The 6th Workshop on
Natural Language Meets Logic and Machine Learning
(NALOMA)**

Proceedings of the Workshop

August 3 - 7, 2026
Prague, Czechia

©2026 Association for Computational Linguistics

Order copies of this and other ACL proceedings from:

Association for Computational Linguistics (ACL)
317 Sidney Baker St. S
Suite 400 - 134
Kerrville, TX 78028
USA
Tel: +1-855-225-1962
acl@aclweb.org

ISBN 979-8-89176-389-0

Introduction

Welcome to the 6th edition of the Natural Language Meets Logic and Machine Learning workshop (NALOMA).

NALOMA continues to serve as a venue dedicated to bridging the gap between machine/deep learning approaches on the one hand, and symbolic/logic-based approaches to natural language understanding and reasoning on the other. A central focus of the workshop remains the development of hybrid approaches and the exploration of theoretical insights that shape and guide computational models of reasoning.

NALOMA took place in August 3-7 during ESSLLI 2026, hosted at Czech Technical University Prague, Czechia. We are deeply grateful to the ESSLLI organizers for their support. The workshop was held over a period of five days, with time slots of about one and a half hours. This year's program featured three inspiring keynotes, seven regular talks with accompanying archival papers included in this proceedings, and two contributed talks based on non-archival submissions.

We would like to thank all authors of archival or non-archival submissions, as well as the dedicated members of the program committee whose careful reviews ensured the quality of the workshop. Our thanks also go to our keynote speakers for sharing their expertise and vision.

As in prior years, NALOMA serves as a platform connecting the symbolic AI and logic communities with the machine learning community, with the dual purpose of promoting discussion and fostering joint research initiatives. We look forward to the collaborations and insights that will arise from this year's event.

NALOMA is endorsed by the Special Interest Group on Computational Semantics (SIGSEM), for which we are grateful.

Hitomi Yanaka and Lasha Abzianidze, Program Co-Chairs

Organization

Program Co-Chairs

Hitomi Yanaka, the University of Tokyo and Riken

Lasha Abzianidze, Utrecht University

Program Committee

Stergios Chatzikyriakidis, University of Crete

Katrin Erk, University of Texas at Austin

Hai Hu, Shanghai Jiao Tong University

Thomas Icard, Stanford University

Aikaterini-Lida Kalouli, Bundesdruckerei GmbH

Lawrence S. Moss, Indiana University

Valeria de Paiva, Topos Institute

Hitomi Yanaka, University of Tokyo

Robin Cooper, University of Gothenburg

Daisuke Bekki, Ochanomizu University

Staffan Larsson, University of Gothenburg

Bill Noble, Massachusetts Institute of Technology

Masaya Taniguchi, Riken

Gijs Wijnholds, Leiden University

Keynote Talk

Reasoning in Language Models: Insights from Cognitive Science

Raffaella Bernardi

Free University of Bozen-Bolzano

Abstract: Coming from the ESSLLI Logic & Language tradition, I view reasoning as a core component of language: language is the vehicle through which we express our otherwise hidden reasoning processes. Large Language Models have reached remarkable levels of linguistic proficiency. The question now is not only what they say, but what lies behind their linguistic expressions.

Reasoning is multifaceted. We deduce conclusions from premises, induce general rules from examples, ask questions to acquire information that is essential for solving a problem, and revise our beliefs when new evidence becomes available. Consequently, the question we face is equally multifaceted: which reasoning abilities have Large Language Models actually acquired, if any, and how can we determine whether they genuinely reason?

In this talk, I will present a series of recent studies from my research group investigating deductive, inductive, and information-seeking reasoning in LLMs. I will discuss the complementary evaluation methodologies we adopted, combining behavioral analyses with investigations of the models' internal reasoning processes. Finally, I will argue that understanding LLM reasoning requires moving beyond static reasoning benchmarks toward dynamic, interactive, multi-turn settings. To support this claim, I will draw on experimental paradigms that have long been used in cognitive science to study human reasoning, showing how they provide powerful tools for investigating the interplay between language and reasoning in modern language models.

Bio: Raffaella Bernardi is Associate Professor at the Faculty of Engineering, Free University of Bozen-Bolzano, Italy. She works on Computational Linguistics focusing on the interplay between Language and Reasoning as displayed during conversations and/or in multimodal contexts. She is part of the Association of Logic, Language and Information (FoLLI) for which she has served in multiple roles (as Secretary, as Logic and Language Area Chair (AC) of the ESSLLI Student Session 2001, local Program Chair (PC) of ESSLLI 2016, PC chair of ESSLLI 2021, etc.). She is an active member of the international Association of Computational Linguistics (ACL), serving as a Senior Chair, AC of *ACL conferences; she is also an active member of the Italian Association of Computational Linguistics (AILC) serving as PC chair, AC, member of the scientific committee of AILC Lectures. She has recently been the EU representative within the ACL Sponsorship Board. She has also been quite active in disseminating Computational Linguistics by means of teaching activities. She has long been the local coordinator, at unibz and then at unitn, of the Erasmus Mundus European Masters Programme in Language and Communication Technologies. She is an ELLIS (European Laboratory for Learning and Intelligent Systems) elected Fellow. Currently, she is serving as Director of the MSc in Computing for Data Science (Faculty of Engineering), member at-large of the ACL Management Board, and member of AILC Management Board.

Keynote Talk

LLMs as a Synthesis between Symbolic and Distributed Approaches to Language

Gemma Boleda
Universitat Pompeu Fabra

Abstract: Since the middle of the 20th century, a fierce battle is being fought between symbolic and distributed approaches to language and cognition. In this talk, I reflect on deep learning models of language in the context of this broader discussion, and consider the possibility that LLMs represent a synthesis between the two antagonistic traditions. This possibility is supported by the fact that deep learning architectures allow for both distributed/continuous/fuzzy and symbolic/discrete/categorical-like representations and processing; what is left to see is how modern language models make use of this flexibility. I will discuss recent research in interpretability that suggests that grammar is processed in a near-symbolic fashion in modern LLMs, while lexical semantics receives a much more distributed treatment; and I will also discuss work in progress in our group that is starting to probe this view. If time allows, I will end the talk by summarizing a recent systematic review of interpretability work on syntactic knowledge in LLMs.

Bio: Gemma Boleda is an ICREA Research Professor in the Department of Translation and Language Sciences of the Universitat Pompeu Fabra (Barcelona, Spain), where she co-leads the Computational Linguistics and Linguistic Theory (COLT) research group. She previously held post-doctoral positions at the University of Trento (Italy), The University of Texas at Austin (USA), and Universitat Politècnica de Catalunya (Spain). She was also a visiting researcher at the Computational Linguistics & Phonetics department (CoLi) of Saarland University and the Institute for Natural Language Processing (IMS) of the University of Stuttgart, both in Germany. In her research, Prof. Boleda uses quantitative and computational methods to investigate how languages work; in particular, how they are shaped by human cognitive constraints, on the one hand, and communicative needs, on the other.

Keynote Talk

Where Did the Symbol Go? Faithful Explanations in the Age of Large Language Models

Mateusz Lango
Charles University

Abstract: The fluency of modern large language models may lead to a perception that natural language generation is close to solved, yet a closer look reveals a recurring pattern of failure whenever a task requires exact, discrete fidelity to some underlying structure rather than fluent pattern completion. Reasoning that is straightforward in a single turn becomes surprisingly brittle across a dialogue; generation under explicit lexical or content constraints remains difficult to guarantee; and explanations produced by LLM-as-judge systems are frequently unfaithful to the actual decision process they claim to justify. I will argue that these three failures share a common shape: each is a case where an LLM must maintain or produce something exact and discrete, but is instead relying on an implicit, continuous representation not well suited to that job. This talk focuses on the third case - faithful explanation - and asks what happens when we make the intermediate representation explicit and symbolic, rather than leaving it to the model to reconstruct implicitly.

I will present two lines of work built on this idea. The first constructs faithful explanations for image classifiers via a pipeline that separates the extraction of the (symbolic, verifiable) evidence behind a decision from its realization into fluent natural language, so that plausibility and faithfulness can be optimized somewhat independently. The second uses LLMs not to generate text directly, but to author interpretable, rule-based generation systems from scratch, producing data-to-text generators that are faithful and inspectable by construction rather than by post-hoc checking. Together, these case studies suggest a general recipe for faithful explanation: let neural models handle fluency and search, and let symbolic structure carry the guarantees. I will close with some open questions about how far this division of labor can be pushed, and what it would take to test the broader thesis motivating the talk.

Bio: Mateusz Lango is a postdoctoral researcher at the Institute of Formal and Applied Linguistics, Charles University, and an assistant professor in the Machine Learning Division at Poznań University of Technology. He completed his Ph.D. in Machine Learning in 2022 under the supervision of Prof. Jerzy Stefanowski, and now works in the group of Prof. Ondřej Dušek. His research focuses on explainable methods for NLP and machine learning more broadly, particularly on producing explanations that are both plausible and faithful. He is also interested in improving natural language generation, making it more accurate, controllable, and faithful.

Table of Contents

<i>Revisiting the Systematicity in Negation in the Era of In-Context Learning</i>	
Hitomi Yanaka and Taisei Yamamoto	1
<i>Neural DTS: Integrating Hyperbolic Classifiers into Natural Language Inference Systems</i>	
Honoka Kobayashi, Hinari Daido and Daisuke Bekki	9
<i>Negation in Reasoning Traces: Interpretable Signals of Correctness and Provenance</i>	
Leon Lukas Hammerla and Alexander Mehler	19
<i>Discovering New Theorems via LLMs with In-Context Proof Learning in Lean</i>	
Kazumi Kasaura, Naoto Onda, Yuta Oriike, Masaya Taniguchi, Akiyoshi Sannai and Sho Sonoda	
40	
<i>Where Does Proof Search Spend Its Effort? A Visualization and Profiling Tool for Formal NLI</i>	
Koharu Saeki and Daisuke Bekki	50
<i>Visual Question Answering and the relation between language, perception and the world</i>	
Staffan Larsson, Bill Noble and Robin Cooper	60
<i>Semantic Distance for Presburger Formula Generation</i>	
Masaya Taniguchi and Hitomi Yanaka	71

Program

Monday, August 3, 2026

- 14:00 - 14:05 *Opening Remarks*
- 14:05 - 15:05 *Keynote: Reasoning in Language Models: Insights from Cognitive Science*
Raffaella Bernardi
- 15:05 - 15:30 *Discovering New Theorems via LLMs with In-Context Proof Learning in Lean*
Kazumi Kasaura, Naoto Onda, Yuta Oriike, Masaya Taniguchi, Akiyoshi Sannai
and Sho Sonoda

Tuesday, August 4, 2026

- 14:00 - 14:25 *Neural DTS: Integrating Hyperbolic Classifiers into Natural Language Inference Systems*
Honoka Kobayashi, Hinari Daido and Daisuke Bekki
- 14:25 - 14:50 *Visual Question Answering and the relation between language, perception and the world*
Staffan Larsson, Bill Noble and Robin Cooper
- 14:50 - 15:15 *Negation in Reasoning Traces: Interpretable Signals of Correctness and Provenance*
Leon Lukas Hammerla and Alexander Mehler

Wednesday, August 5, 2026

- 14:00 - 15:00 *Keynote: LLMs as a Synthesis between Symbolic and Distributed Approaches to Language*
Gemma Boleda
- 15:00 - 15:25 *Revisiting the Systematicity in Negation in the Era of In-Context Learning*
Hitomi Yanaka and Taisei Yamamoto

Thursday, August 6, 2026

- 14:00 - 14:20 (non-archival) *Lost in Transfer: Argument Mining with Large Language Models — From In-Context Learning to Policy Optimization*
Fatma-Zohra Rezkellah, Reto Gubelmann
- 14:20 - 14:40 (non-archival) *Neural Wani: Toward Accelerating the Automated Theorem Prover wani for Dependent Type Theory*
Nanako Miyagawa, Hinari Daido, Daisuke Bekki
- 14:40 - 15:00 *Where Does Proof Search Spend Its Effort? A Visualization and Profiling Tool for Formal NLI*
Koharu Saeki and Daisuke Bekki
- 15:00 - 15:25 (Special talk) *Mathematical NLI*
Lawrence S. Moss

Friday, August 7, 2026

- 14:00 - 15:00 *Keynote: Reasoning in Language Models: Insights from Cognitive Science*
Mateusz Lango
- 15:00 - 15:25 *Semantic Distance for Presburger Formula Generation*
Masaya Taniguchi and Hitomi Yanaka
- 15:25 - 15:30 *Closing Remarks*