

WASSIM (WES) BOUAZIZ

✉ wassim.bouaziz@mistra.ai

🎓 Google Scholar

🔊 wesbz

🌐 wassim-bouaziz

RESEARCH EXPERIENCE

Mistral AI AI Scientist

Mistral AI, Safety team

Initiating research projects on AI security.

Nov. 2025 – Present

Paris, France

Meta & École polytechnique, X Research Assistant, PhD student

FAIR & CMAP department at École polytechnique

I explored a spectrum of research questions at the intersection of AI Security, Safety, and Interpretability.

- Published research work (either as first author or contributor) in top venues including ICLR, ICASSP and NeurIPS.
- Contributed to Llama 3 models by developing safety evaluations and safety data annotation campaigns.

May 2022 – Oct. 2025

Paris, France

Independant Machine Learning Consultant

Consulting and Teaching in Machine Learning

- Giving Deep Learning lectures and supervising projects at Le Wagon.
- Advising startups on project feasibility and state of the art in machine learning research on various topics.

Jan. 2021 – May. 2022

Paris, France

Facebook AI Research Research Intern

Noise-robust Automatic Speech Recognition

- Achieved 10% accuracy improvement using audio data augmentation to enhance robustness of speech recognition systems.

Jul. – Nov. 2020

Paris, France

ENS Research Intern

ENS CoML team

- Classification of child/adult directed speech for a world-wide study on child language acquisition (ACLEW).

Apr. – Aug. 2019

Paris, France & Montréal, Canada

CentraleSupélec Graduate Research Assistant

Graphs, Algorithms and Combinatorics team

- Contributed to improving the known time complexity of minimum cuts counting in graphs.

Oct. 2018 – Apr. 2019

Saclay, France

Snips (Acquired by Sonos) Embedded Software Intern - AI

Embedded systems - Real time inference

- Developed real time inference application using Snips' voice assistant with a < 30% CPU usage on a Cortex-M7 microcontroller.

Jun. – Aug. 2018

Paris, France

EDUCATION

Meta & École polytechnique, X PhD in Mathematics and Computer Science

Defended Dec. 2025. Towards Secure and Trustworthy Machine Learning:

From Data Poisoning to Ownership Verification

- Advisor: Eric Moulines & El Mahdi El Mhamdi (CMAP, École polytechnique),
Nicolas Usunier (Helsing), Max Nickel (FAIR, Meta)

May 2022 – Oct. 2025

Paris, France

Paris-Saclay University MSc in Mathematics, Vision and Learning ("MVA")

Among the World's best university in Mathematics in Shangai's ranking - with Highest Honors, GPA: 4.0/4.0

Paris, France

Université Paris-Saclay MSc degree in Computer Science and Machine Learning

Double degree with CentraleSupélec, GPA: 4.0/4.0

2018 – 2019

Paris, France

CentraleSupélec MEng degree in Mathematics and Computer Science.

Top Engineering school in Computer Science and Electrical Engineering. GPA: 3.9/4.0

2016 – 2019

Gif-sur-Yvette, France

Preparatory classes Lycée Blaise Pascal

Two-year undergraduate intensive program for the entrance exams to top French Engineering schools

2014 – 2016

Orsay, France

TECHNICAL SKILLS

AI Large scale training, Language models training (pre & post training),

Programming Python (PyTorch), Bash, Latex, C++⁽⁻⁾ SQL⁽⁻⁾, HTML/CSS/JS⁽⁻⁾,

Languages French (*native*), English (*C1*), Spanish (*B2*), Arabic (*A2*), Japanese (*A1*)

EXTERNAL ENGAGEMENT

Reviewer	ICLR (2023, 2025, 2026), NeurIPS (2024, 2025)
Invited Talks	• “Risks of AI agents in high-stakes environment” in the “House of Finance Days” seminar at Université Dauphine-PSL (2026). • “Towards Secure and Trustworthy Machine Learning: From Data Poisoning to Ownership Verification” in the “AI for social good” course of the “AI and society” PSL Master (2026). • “AI security & data poisoning” for <i>Hacienda</i>, a French tech community (2026). • “How can AI models be hacked?” in the French Underscore_tech show (2025) – >400k views on YouTube • “Harnessing AI’s Fragility: A Tool for Dataset Tracing” in PurrLab seminar (2025), INRIA Sierra seminar (2025), Callyope research seminar (2025) • “AI Security & Safety” in the “Reliable and Collaborative learning” course at École polytechnique (2024) • “Digital commons in AI research” in the “Introduction to digital commons and public policy making” course for Public Policy Master students at Sciences Po Paris (2023, 2024)

PUBLICATIONS

– First author publications

Byzantine Machine Learning: MultiKrum and an optimal notion of robustness <i>Gilles Bareilles, Wassim (Wes) Bouaziz, Julien Fageot, El Mahdi El Mhamdi</i>	October 2025 <i>arXiv</i>
Winter Soldier: Hypnotizing Language Models at Pre-Training with Indirect Data Poisoning <i>Wassim (Wes) Bouaziz, Mathurin Videau, Nicolas Usunier, El Mahdi El Mhamdi</i>	April 2025 <i>ICLR 2026</i>
Data Taggants: Dataset Ownership Verification via Harmless Targeted Data Poisoning <i>Wassim (Wes) Bouaziz, Nicolas Usunier, El Mahdi El Mhamdi</i>	October 2024 <i>ICLR 2025</i>
Targeted Data Poisoning for Black-Box Audio Datasets Ownership Verification <i>Wassim (Wes) Bouaziz, El Mahdi El Mhamdi, Nicolas Usunier</i>	October 2024 <i>ICASSP 2025</i>
Inverting Gradient Attacks Makes Powerful Data Poisoning <i>Wassim (Wes) Bouaziz, El Mahdi El Mhamdi, Nicolas Usunier</i>	October 2024 <i>TMLR Dec. 2025</i>
Clustering Head: A Visual Case Study of the Training Dynamics in Transformers <i>Ambroise Odonnat, Wassim (Wes) Bouaziz, Vivien Cabannes</i>	December 2024 <i>arXiv</i>
Easing Optimization Paths: a Circuit Perspective <i>Ambroise Odonnat, Wassim (Wes) Bouaziz, Vivien Cabannes</i>	December 2024 <i>Lecture ICASSP 2025</i>

– Collaboration publications

On Monotonicity in AI Alignment <i>Gilles Bareilles, Wassim (Wes) Bouaziz, Sadeh Farhadkhani, Julien Fageot, Lê-Nguyên Hoang, Sébastien Rouault, Peva Blanchard, El-Mahdi El-Mhamdi</i>	May 2025 <i>Under review</i>
On the strength of goodhart’s law <i>Adrien Majka, Wassim (Wes) Bouaziz, El Mahdi El Mhamdi</i>	May 2025 <i>2nd MoFA Workshop at ICML 2025</i>
Cultivating Pluralism In Algorithmic Monoculture: The Community Alignment Dataset <i>Lily H Zhang, Smitha Milli, Karen Long Jusko, Jonathan Smith, Brandon Amos, Wassim (Wes) Bouaziz, Jack Kussman, Manon Revel, Lisa Titus, Bhaktipriya Radharapu, Jane Yu, Vidya Sarma, Kristopher Rose, Maximilian Nickel</i>	May 2025 <i>Oral 2nd MoFA Workshop at ICML 2025</i>
The Llama 3 Herd of Models <i>Llama Team, AI @ Meta</i>	July 2024 <i>arXiv</i>
Iteration head: A mechanistic study of chain-of-thought <i>Vivien Cabannes, Charles Arnal, Wassim Bouaziz, Xingyu Yang, Francois Charton, Julia Kempe</i>	December 2024 <i>NeurIPS 2024</i>
On the Parameterized Complexity of Counting Small-Sized Minimum (S, T)-Cuts <i>Pierre Bergé, Wassim Bouaziz, Arpad Rimmel, Joanna Tomasiak</i>	December 2023 <i>SIAM Journal on Disc. Math.</i>

