

PERSONAL INFORMATION

Name	Marco Simoni Researcher, Ph.D. in Artificial Intelligence — Reinforcement learning for post-training and alignment of Large Language Models · Applied AI for cybersecurity
Address	Pisa, Italy
Telephone	+39 366 992 4397
E-mail	marco.simoni0711@gmail.com
Web profiles	LinkedIn — marco-simoni GitHub — winstonsmith1897 winstonsmith1897.github.io Google Scholar ORCID 0009-0000-4170-503X
Nationality	Italian
Date and place of birth	16 August 1998, Grosseto (GR), Italy
Sex	Male
Professional profile	Ph.D. in Artificial Intelligence (Cybersecurity curriculum) at Sapienza University of Rome, specialised in reinforcement learning for post-training and alignment of Large Language Models. Research focused on LLM-based AI agents supporting experts in cyber threat intelligence, and on distributed IoT for smart environments. Author of peer-reviewed publications at ACL, ARES, SECRIPT, SAC and in international journals.

WORK EXPERIENCE

June 2026 - ongoing	Researcher Scuola Superiore Sant’Anna , Pisa — Department of Excellence in Robotics and AI Research on reinforcement learning for post-training and alignment of Large Language and Multimodal Foundation Models: policy optimisation, reward design and training stability.
Oct. 2025 - Sep. 2026	Horus Project — Predictive Cybersecurity Foundational Model AI Researcher Project Collaborator — NetGroup - Engineered from scratch a hierarchical Transformer (HAT), <i>~10M parameters</i>, that predicts the next attack from an attack sequence and autoregressively generates the continuation of the campaign. - Trained on 80,000 samples over four features (country, threat group, industry, time), reaching <i>F1 = 0.60</i> on next-attack prediction.
May - Dec. 2025 Mar. - Apr. 2026	Automated Translation of Natural Language into XACML Policies AI Researcher Project Collaborator — CNR-IIT - Developed a two-step LLM-powered framework converting unstructured natural language commands into fully compliant U-XACML access and usage control policies, via an intermediate structured JSON representation. 93% accuracy in policy generation, 91% on ambiguous or noisy inputs and 98% agreement with expert-defined policies; taxonomy-driven dataset generalising across smart home, smart office and healthcare settings. - Hybrid LLM + dedicated-library design validated for reliable on-device inference under constrained hardware.

EDUCATION

1 Nov. 2022 - 29 Jan. 2026	Ph.D. in Artificial Intelligence — Curriculum: Cybersecurity Sapienza University of Rome , Italy — Final grade: <i>Excellence</i> Thesis: <i>“Toward Reliable and Adaptive Large Language Models in the Cybersecurity Domain”</i> Full text
Nov. 2020 - Sep. 2022	M.Sc. in Artificial Intelligence and Data Engineering University of Pisa , Italy — Final grade: <i>110/110 cum laude</i>
Sep. 2017 - Nov. 2020	B.Sc. in Electronic Engineering University of Pisa , Italy

PUBLICATIONS

Journal articles	- Nayab, S. et al. (2026). “Concise Thoughts: Impact of Output Length on LLM Reasoning and Cost”. <i>Information Sciences</i> 757, p. 123939. ISSN 0020-0255. - Simoni, M. et al. (2026). “GTPO: Trajectory-Based Policy Optimization in Large Language Models”. <i>Transactions of the Association for Computational Linguistics</i> (accepted). - Alajramy, L. et al. (2025). “On-Device Derivation of IoT Usage Control Policies: Automating U-XACML Policy Generation from Natural Language with LLMs in Smart Home Environments”. <i>Future Generation Computer Systems</i>, p. 108067.
------------------	--

Conference proceedings

- Fontana, A., Simoni, M. et al. (July 2026). "On the Hidden Objective Biases of Group-Based Reinforcement Learning". *Proc. of the 64th Annual Meeting of the ACL (Vol. 2: Short Papers)*, pp. 109–121. ISBN 979-8-89176-391-3.
- Fontana, A. and Simoni, M. (2025). "Unmasking Model Behavior: How LLMs Reason on Vulnerability Detection". *Int. Conf. on Availability, Reliability and Security (ARES)*, Springer, pp. 316–333.
- Nayab, S. et al. (2025). "Leveraging Knowledge Graphs and LLMs for Structured Generation of Misinformation". *ARES*, Springer, pp. 334–350.
- Simoni, M. and Saracino, A. (2025). "MATRIX: A Comprehensive Graph-Based Framework for Malware Analysis and Threat Research". *Proc. of the 22nd Int. Conf. on Security and Cryptography (SECRYPT 2025)*, pp. 495–502.
- Simoni, M., Saracino, A. et al. (2025). "MoRSE: Bridging the Gap in Cybersecurity Expertise with Retrieval Augmented Generation". *Proc. of the 40th ACM/SIGAPP Symposium on Applied Computing*, pp. 1213–1222.
- Simoni, M. and Saracino, A. (2024). "Cybersecurity with LLMs and RAGs: Challenges and Innovations". *Security and Privacy in Communication Networks (SecureComm)*.
- Simoni, M. et al. (2023). "Graph-Based Android Malware Detection and Categorization through BERT Transformer". *Proc. of the 18th Int. Conf. on Availability, Reliability and Security*, pp. 1–7.

TEACHING

Ph.D. course (May 2026)

Lecturer — **Trustworthy AI, Ph.D. programme**

10 hours of lectures on **Transformers from scratch**, covering both the theoretical foundations and the full implementation of the architecture.

EUROPEAN PROJECTS

Sep. 2025 - ongoing

TURING — Trustworthy Unified Robust Intelligent Generative Systems → [Website](#)

TASK LEADER *Adversarially Robust Model Training*

Led the development of adversarially robust training pipelines for **Multimodal Foundation Models** applied to complex physical-system simulations.

Nov. 2022 - Nov. 2024

SIFIS-HOME — Secure Interoperable Full-Stack IoT for Smart Homes → [Website](#)

PROJECT CONTRIBUTOR *H2020 programme*

Implemented core modules of the Data Analytics Toolbox: *System Protection Manager, Application Manager and Intrusion Detection System*.

PERSONAL SKILLS

Languages

Italian — mother tongue · **English** — C1 (Writing · Listening · Speaking)

LLM training & retrieval

Policy optimisation — PPO, GRPO, GTPO¹ · Unsloth¹ · verl¹ · TRL¹ · Sparse MoE, RoPE, GQA³ · LangChain² · HuggingFace² · vLLM¹

AI for cybersecurity

CTI⁴ · Knowledge Graphs — NetworkX, Neo4j⁴ · MITRE ATT&CK^{2,4} · CAPEC · Metasploit²

Core ML & engineering

Python^{1,2,3,4} · PyTorch^{1,4} · JAX/Flax³ · TensorFlow · SQL · Docker · Linux · Git

Research & organisational

Task leadership in EU-funded consortia; scientific writing and peer-reviewed publishing; end-to-end research from algorithm design to large-scale training and evaluation.

Portfolio projects

applications of the skills marked 1–4

[1] GTPO — Trajectory-Based Policy Optimization for LLM Post-Training → [GitHub](#)

Stack: Python, PyTorch, Unsloth, verl

Designed a KL-free policy optimisation algorithm mitigating LLM policy collapse; up to **+15%** reasoning performance on OOD benchmarks (AIME 2024, AIME 2025, AMC) vs. GRPO.

[2] MoRSE — Mixture-of-RAG-Security-Experts → [GitHub](#) · [YouTube](#)

Stack: Python, LangChain, HuggingFace

Dual-cascaded RAG framework with **7 parallel retrievers** for cybersecurity Q&A; outperformed *GPT-4* by **15%** in response accuracy on general and multi-hop questions, **84%** accuracy in CVE identification.

[3] DantinoX — A Unified Framework for Multi-Paradigm Language Modeling → [GitHub](#) · [Docs](#) · [YouTube](#)

Stack: Python, JAX, Flax

Built an LLM architecture featuring **Sparse MoE, RoPE** and **GQA**; maximised hardware throughput via Sliding Window Attention, Static KV Caching and Gradient Checkpointing.

[4] TITAN — Context-Aware Reasoning for Cyber Threat Intelligence (CTI) → [GitHub](#)

Stack: Python, PyTorch, NetworkX

Architected a Knowledge Graph reasoning framework for CTI, automating complex threat analysis by modelling relationships across IoCs, TTPs and CVEs.

** 1,2,3,4 denote the application of each skill in the corresponding portfolio project.*

Pisa, August 11, 2026



Marco Simoni