

Each session is up to 90 minutes. Full papers are allocated 25 minutes and short papers 15 minutes. As of July 24, 2026

Full Paper (F) /			Paper title			Authors			Online	Date	Start	End
Paper ID	Short Paper (S)											
Session 1			Evaluating AI systems			Chair: Muhammad Mudassar Yamin						
23	F		Tianyue Wang, Fei Yuan, Chijiu Miao, Muralidharan Seshadri Kaniyanoor and Kazuo Matsumoto			AutoQA-Bench: A Three-Level UX Benchmark for Automated Evaluation of Open-Ended LLM Responses						
106	F		Francisca Isabella Manope, Adam Rafif Faqih, Daffa Hilmy Fadhlurrohman, Jati H. Husen and Rosa Reska Riskiana			LLM-Based Automatic Metamorphic Test Case Generation for LLM Fairness Testing						
88	F		Ren Wei Lu, Wan Ping Chen and Fang Yu			Dual-Track Red Teaming and Evaluation for Traditional Chinese Localized LLMs						
39	S		Landy Rakotoarison and Fanamby Randriamaherintsoa			Do AI Agents Exhibit Greed or Empathy in Shared Resource Environments?						

Session 2 AI for Software Quality Engineering				Chair: Nayab Shahzad					
71	F	Erik Trevino		Deterministic Behavioral Contract Testing for AI Features at the Browser Layer			27-Jul	15:45	16:10
22	F	Ryoya Koyama, Zhiyao Wang, Devi Karolita, Jialong Li and Kenji Tei		Bridging the Interpretation Gap in Accessibility Testing: Empathetic and Legal-Aware Bug Report Generation via Large Language Models			27-Jul	16:10	16:35
37	S	João Domingues, Daniela Duarte, António Sousa and Nuno Pombo		AI4SE for CI/CD: Explainable Code Smell Risk Analysis			Online	27-Jul	16:35 16:50
99	S	Raquel Blanco, María José Suárez-Cabal, Javier Tuya and Claudio de la Riva Álvarez		A RAG approach for assisting testers in database query testing			27-Jul	16:50	17:05

Session 3			LLMs for Software Engineering and Security			Chair: Tatsuhiro Tsuchiya						
98	F		Fuze Kuang, Dongmei Liu, Mingzhe Zhao and Longyang Mi			A Dynamic Evaluation Approach to Repository-Level Code Generation Via LLM			Online	28-Jul	10:15	10:40
103	F		Yuxuan Wang, Jingshu Chen and Qingyang Wang			LLM-Based Detection and Test Generation for Command Injection Vulnerabilities in Python			Online	28-Jul	10:40	11:05
54	F		Mizuki Yamada, Masahiko Kato and Juichi Takahashi			Improving LLM-Based Unit Test Generation Through Root-Cause-Driven Prompt Design				28-Jul	11:05	11:30
26	S		Muhammad Mudassar Yamin			VectorSec: A Web-Based AI Security Scanner for Systematic Evaluation of LLM Vulnerabilities				28-Jul	11:30	11:45

Session 4			Test Prioritization and Test Process Improvement			Chair: Tatsuhiro Tsuchiya		
78	F	Tien Duc Nguyen, Gonzalo Oberreuter, Siegfried Steckenbiller, Patrick Tkalcic and Gordon Fraser	Evaluating DL Test Adequacy Metrics: A Correlation Study with DeepCrime's Mutation Score			28-Jul	14:30	14:55
105	F	Gianmarco De Vita, Nargiz Humbatova and Paolo Tonella	Evaluating the Efficacy and Diversity of DNN Test Input Prioritisation Techniques			28-Jul	14:55	15:20
79	F	Ahmed Twabi, Yepeng Ding and Tohru Kondo	Formal Trajectory Analysis for Testing Agentic AI in Stateful Environments			28-Jul	15:20	15:45
2	S	Osama Saad, Mohamed Abdelkarim and Reem Eladawi	TE-TCP-Net: Parallel Transformer-Encoders for Test Case Prioritization			28-Jul	15:45	16:00

Session 5 Robustness and Regression Verification				Chair: Ross Arnold					
83	F	Li-Jen Lin and Chih-Duo Hong		Robustness Verification of RNN with Abstraction Refinement			29-Jul	10:15	10:40
87	F	Nayab Shahzad and Franz Wotawa		Ontology-based Testing of Reinforcement Learning Applications			29-Jul	10:40	11:05
104	F	Adiba Mahmud, Yasmeen Rawajfih and Ross Arnold		Beyond Bypass: Measuring Vulnerability Regression in LLM-Generated Security Patches			29-Jul	11:05	11:30
80	S	Vinil Pasupuleti, Srinivas Teja Songa, Nagesh Gulkotwar, Shyalendar Reddy Allala and Shrey Tyagi		Behavioral Regression Testing for Continuously Updated Machine Learning Models			Online	29-Jul	11:30 11:45

Session 6			Domain-Specific AI Applications			Chair: Luciano Alessandro Ipsaro Palesi					
89	F		Tang Phu Thien Nhan and Tomohiro Shibata			Transformer-Based Temporal Action Localization of Fine-Grained Operations in Biological Experiment Videos and Analysis of Compositional Limitations			29-Jul	15:30	15:55
107	F		Minggu Wang, Zhiyong Zhou, Fuyuan Zhang, Jianjun Zhao and Vinh Truong Duy Truong			Improving Robustness of Semantic Segmentation for Autonomous Driving: A Case Study			29-Jul	15:55	16:20
46	S		Mateus Molina, Alexander Günther and Peter Liggesmeyer			Connecting Uncertainty Quantification to Safety Engineering: Uncertainty Gate Framework for Safe ML Decision Making			29-Jul	16:20	16:35
56	S		Alena Bulda, Alexey Itkin, Daria Degtiarenko and Iosif Itkin			Passive Testing of AI-Enhanced Financial Systems via Protocol Traffic Analysis and Property-Based Validation			29-Jul	16:35	16:50

Online Session															
Session 7			Testing and Validation of AI Systems				Chair: Yuhuan Zhou								
94	F	Lingling Zhang, Fen Zhao and Xin Yang					A C-SSRS-Based Risk-Aware Evaluation Framework for Self-Harm Safety in LLMs					Online	28-Jul	14:30	14:55
		Michael Johnson and Khaled Sihoub					Direct Transfer Learning for Cross-Project Test Case Prioritization under Cold-Start Conditions					Online	28-Jul	14:55	15:10
55	S	Dipayan Banik, Kowshik Chowdhury and Shazibul Islam					All Smoke, No Alarm: Oracle Signals in Agent-Authored Test Code					Online	28-Jul	15:10	15:25
61	S						Lifeng Ni and Lixin Tao					Bandit-Controlled Semi-Ensemble Learning for Accuracy-Energy Trade-off Under Real-time Edge Conditions			
96	S	Johannas Meela Katikala, Kanishk Sharma, Tharun Swaminathan Ravi Kumar and Sharad Sharma					Conversational AI for Safety-Critical Virtual and Extended Reality: Intelligent Agents Across Health Informatics, Navigation, and Urban Sensing					Online	28-Jul	15:40	15:55
101	S											Online	28-Jul	15:40	15:55

[Invited] AI Test 1 (Invited session 2)			Chair: Haihua Chen			
INV16		Karthik Rengarajan, Siddharth Pokuri, and Jerry Gao	A Survey on Intelligent Robot Testing and Automation			27-Jul 13:45 14:30
INV07		Huyen Nguyen, Haoxuan Zhang, Yang Zhang, Haihua Chen, Junhua Ding	LongSumEval: Question-Answering Based Evaluation and Feedback-Driven Refinement for Long Document Summarization			Online 27-Jul 14:30 15:15

[Invited] AI Test 2 (Invited session 7)			Chair: Zhiming Zhao				
INV17		Xiaoke Han, Hong Zhu	MASTEST: A LLM Empowered Multi-Agent System for Testing REST APIs		29-Jul	13:30	14:15
INV05		Nariman Mani, Amr S. Abdelfattah, Shakthi Weerasinghe, Xiaozhou Li, Tomas Cerny	Multi-Agent Change Impact Analysis and Test Optimization for AI-Enabled Software Systems		29-Jul	14:15	15:00