

Paper ID Paper Title

Regular Papers	
627	SecurPoAgent: A Hybrid LLM Agent for Automated Python Static Vulnerability Repair
646	LLM2v: Contextual Prompt Whitelist Principles for Agentic LLMs
649	Detection of LLM Hallucinations Using Late Internal Representations
650	Breaking and Securing Vision Language Models: An Adversarial Robustness Study
675	Improving Robustness of Large Language Models Used in Healthcare Through Prompt Engineering
687	Jailbreaking Large Vision Language Models in Intelligent Transportation Systems

Authors

Jugal Gajjar (The George Washington University)*, Kamalakaran Subramaniakuppasamy (The George Washington University), Rishi Puthal (The George Washington University), Kaustik Ranaware (The George Washington University)
Tom Pawelek (Mississippi State University)*, Raj Patel (The University of Alabama), Charlotte Crowell (The University of Alabama), Noorbakht Amir Gollari (The University of Alabama), Sudip Mitral (The University of Alabama), Shahram Rahimi (The University of Alabama)
Sabbawar Hossain (University of North Carolina Greensboro)*, Ivo Dene (University of North Carolina Greensboro)
Md Zarf Hossain (Florida Atlantic University), Abdul K. Shahid (Southern Illinois University Carbondale), Ahmed Imbaj (Florida Atlantic University)*
Jonathan O'Herry (University of North Florida)*, Indika Kahanda (University of North Florida), Upulie Kanewala (University of North Florida)
Bashan Chandra Das (Florida International University)*, Md Tasmin Javed (Florida International University), Md Jusal Mia (Florida International University), M. Hadi Amin (Florida International University), Yingshao Wu (Florida International University)