



AI & LLM SECURITY ASSESSMENT

DELIVERED BY CERTIFIED CYBERSECURITY SPECIALISTS WITH EXPERIENCE ACROSS PUBLIC-SECTOR, FINANCIAL, AND OTHER HIGH-STAKES DIGITAL SYSTEMS

We have yet to see an AI agent that should be considered secure without dedicated testing — even in environments where teams believed the risks were already under control.

This document shows how we uncover real weaknesses, abuse paths, and data exposure risks before they become costly business or compliance failures. Read on.

An AI agent does not have to fail visibly to become dangerous. It can stay active, look normal, and still cause hidden business damage.



Why AI Security Testing Matters

AI agents are no longer just chatbots. They can access data, call APIs, trigger workflows, and act inside business systems — which also makes them a new attack surface. We help companies assess AI agents, LLM applications, RAG systems, and connected AI workflows before hidden risks become business, security, or compliance problems.

WHAT WE TEST

We assess more than model output. Our scope includes AI agents with access to tools, APIs, and internal systems; LLM applications, chatbots, and AI assistants; RAG systems connected to internal knowledge bases; and AI solutions integrated with CRM, HR, finance, support, or operational workflows.

RISKS WE ASSESS

We test for prompt injection and jailbreaks, unauthorized tool or API use, data leakage from prompts, memory, files, RAG, or connected systems, excessive agent permissions, workflow abuse, guardrail bypass, unsafe handling of files and user input, and exposure of confidential or regulated data.

HOW WE WORK

We perform practical security testing from the perspective of an attacker, malicious user, insider, or compromised external source. The goal is to determine whether an AI system can be pushed into unsafe actions, misuse connected tools, bypass restrictions, or create risk for adjacent internal systems and business processes.

WHAT CLIENTS GET

Clients receive a structured report with identified vulnerabilities, attack scenarios, severity level, business impact, technical evidence, reproduction steps, remediation guidance, and a prioritized action plan for security, engineering, and product teams. This also helps organizations prepare for growing regulatory expectations such as the EU AI Act, NIS2, GDPR, DORA, and the Cyber Resilience Act.



AI agents introduce a different level of risk: they do not just generate answers — they can trigger actions, expose data, misuse connected systems, and become a path into internal infrastructure. Testing them early reduces security, operational, and compliance risk before AI is scaled further.



This service is designed for companies that are already using AI in real workflows — or are preparing to scale it further. Once AI agents are connected to tools, internal data, customer-facing processes, or business systems, security testing becomes a business, operational, and compliance requirement rather than a technical nice-to-have.

WHO THIS IS FOR

! AI risk grows fast once AI moves from isolated experimentation into connected systems, internal workflows, and customer-facing operations.

AI IN WORKFLOWS

CONNECTED TO BUSINESS SYSTEMS

- Relevant for companies that:
- deploy AI agents or LLM-based solutions
 - connect AI to internal business systems
 - use AI across support, operations, finance, HR, legal, DevOps, or cybersecurity
 - rely on AI to assist decisions or trigger actions inside real workflows

SENSITIVE DATA

CONFIDENTIAL & REGULATED ENVIRONMENTS

- Especially important where AI systems interact with:
- personal data
 - confidential internal information
 - regulated business processes
 - financial or operational records
 - customer-facing systems that can expose business risk if misused

SCALING AI

RISK BEFORE EXPANSION

Most companies focus on functionality first and security later. This service helps reduce risk before AI is scaled more broadly across teams, systems, customers, or operational workflows — when failures become more costly and harder to contain.

COMPLIANCE RELEVANCE

WHY THIS NOW MATTERS

- AI security testing supports readiness for growing regulatory expectations, including:
- EU AI Act
 - NIS2 Directive
 - Cyber Resilience Act
 - GDPR
 - DORA for the financial sector
 - Product Liability rules affecting software and AI-enabled systems

This service is most valuable before AI is scaled further — when companies still have the chance to identify hidden weaknesses, reduce exposure, and improve trust in how AI behaves inside real business environments.



RELEVANT EXPERIENCE & SELECTED PROJECTS

Barabaka Production includes a dedicated Cybersecurity unit focused on practical AI security testing and security assessment. Our work is grounded in practical cybersecurity testing across complex digital environments — including web applications, APIs, mobile products, internal systems, and AI solutions. This experience gives us a strong foundation for assessing AI agents and connected AI systems in real operational contexts.

We have contributed to or supported large-scale projects where security, reliability, confidentiality, and system trust are critical. That experience is directly relevant to organizations deploying AI in business, public-sector, and regulated environments.

Trusted in high-stakes digital systems where failure, misuse, or data exposure carry real consequences. Identify your AI exposure early — schedule a scoping call.

Our security expertise is supported by internationally recognized professional certifications held by members of our team. Individual certification details can be provided upon request.

HTB Certified Penetration Testing Specialist (HTB CPTS)
Web Application Penetration Tester eXtreme (eWPTX)
Certified Ethical Hacker (CEH)
Certified Penetration Testing Engineer (CPTe)
ISO/IEC 27032 Lead Cybersecurity Manager — PECB



+43 677 64899747 AT www.barabaka.studio
+38 098 070 5695 UA info@barabaka.studio

