

AI-SEC

Secure GenAI and Agent Threat Modeling

Two days after which you hold a threat model, a data flow diagram and a test plan for a real application instead of a list of theoretical risks.

DURATION

2 days

GROUP

up to 12 people

FORMAT

In-house · Remote · In the lab · German or English

THIS EXISTS AFTERWARDS

- ▶ A data flow diagram of your application from input through model, knowledge sources and tools to output
- ▶ A threat model that carries a STRIDE-style method over to GenAI and agents
- ▶ A catalogue of concrete protective measures: guardrails, permissions and output filters, assigned per threat
- ▶ A test plan that makes it verifiable whether a measure works
- ▶ A prioritised residual risk list with owners

AGENDA

Day 1 · 09:00	Why GenAI carries different threats: Prompt injection, unsafe outputs, data leakage through knowledge sources and tools. Core question: where does this differ from classic application security and where does it not?
Day 1 · 10:30	Data flow diagram of your application: We draw the full flow: input, model, RAG sources, tool calls, output. Exercise: every trust boundary is marked and named.
Day 1 · 13:30	Carrying a STRIDE-style method over to GenAI: We work through the threat categories and translate them onto prompt, context, tools and output. The OWASP LLM reference frame serves as a checklist so no category is forgotten.
Day 1 · 15:30	Collecting threats on your application: Exercise on your data flow diagram: per trust boundary we name plausible attacks. Day result: a first threat model with prioritised threats.
Day 2 · 09:00	Live demonstration in the lab: On a prepared environment we show how a prompt injection works through a RAG source and what an over-permissioned tool gives away. No customer data, only to understand the effect.
Day 2 · 11:00	Protective measures per threat: Guardrails, least-privilege permissions, output filters, source checks. Decision per threat: which measure, which effort, which residual risk remains?
Day 2 · 14:00	Building the test plan: For each measure a test that proves its effect. Exercise: for each important threat we formulate a concrete test case with an expected result.
Day 2 · 16:00	Residual risks and handover: We gather what remains and name owners. Day result: threat model, measures catalogue, test plan and prioritised residual risk list for your application.



YOUR TRAINER

Dino Bordonaro

Microsoft MVP for eight consecutive years (Microsoft Azure · Cloud and Datacenter Management), Microsoft Certified Trainer, almost 30 years in IT. Builds sovereign AI systems on own hardware (Azure Local, Foundry Local, NVIDIA GPUs) and trains on exactly that. More than 800 professionals trained, every concept proven in his own enterprise lab in a TÜV-audited datacenter.

€15,800 flat, plus VAT

The track delivers model, measures and test plan as a draft. It does not replace implementation, a penetration test mandate or a managed security service.

AI-SEC