



ANCHOR CYBER SECURITY LLC

STRIDE + AI Threats

Microsoft's six threat categories, plus what changes when the component being threat-modeled is an LLM instead of a conventional service.

\$ stride --categories

The Six Categories

LETTER	THREAT	VIOLATES	ASK
S	Spoofing	Authentication	Can someone convince the system they're a different, legitimate identity?
T	Tampering	Integrity	Can data or code be modified without authorization, in transit or at rest?
R	Repudiation	Non-repudiation	Could someone perform an action and later credibly deny doing it?
I	Information Disclosure	Confidentiality	Can data reach someone who isn't supposed to see it?
D	Denial of Service	Availability	Can the system be made unavailable to legitimate users?
E	Elevation of Privilege	Authorization	Can someone do something the system shouldn't let them do?

```
$ stride --component llm
```

Where STRIDE Bends for an LLM Component

Tampering, sideways

Data & Model Poisoning

Tampering usually means "modify data in transit." Here it can mean corrupting training or fine-tuning data long before the attack surfaces — the tampered artifact is the model's weights, not a record in a database.

Not really Spoofing

Prompt Injection

The attacker doesn't spoof an identity — they smuggle an instruction inside content the model treats as trustworthy. Closer to Tampering-with-intent than classic Spoofing. OWASP's #1 LLM risk for two years running.

Information Disclosure, amplified

Sensitive Info Disclosure & System Prompt Leakage

A model can memorize and reproduce training data verbatim, or be talked into revealing its own system prompt — a disclosure vector with no equivalent in a normal CRUD app.

Elevation of Privilege, agentic

Excessive Agency

An agentic AI wired to tools/APIs it doesn't need for its actual task is a standing privilege-elevation risk even with zero code changes — the excess permission is the vulnerability.

```
$ stride --apply
```

Applying This

Run STRIDE first against the product's conventional architecture the normal way — every service, every trust boundary, every data flow. Then run the AI-specific pass separately against any component that inspects, summarizes, or makes a decision using an LLM — that's a second, different threat surface layered on top of the first, not a replacement for it.

src: Microsoft STRIDE · OWASP Top 10 for LLM Applications 2025 (top-ranked items)

exit 0