

Experții în cybersecurity au analizat 800.000 de aptitudini ale AI, iar 25.000 au fost considerate suspecte (analiza)

Peste 25.000 de aptitudini ale Inteligenței Artificiale (AI) dintr-un total de 800.000 analizate au fost considerate suspecte de către experții în securitate cibernetică Eset, în urma efectuării unei ample analize.

Potrivit datelor prezentate într-un articol publicat, recent, pe blogul din România al producătorului de soluții antivirus, din martie 2026 au fost cercetate peste 800.000 de AI skills, fiind identificate mai mult de 25.000 suspecte, dintre care aproximativ 2.500 au fost clasificate drept malițioase.

'Unele dintre cele mai evidente exemple au fost concepute pentru a descărca pe dispozitivele utilizatorilor programe malware de tip infostealer, cu scopul de a colecta informații sensibile, precum chei API, parole sau date de plată pentru aplicații de criptomonede și servicii bancare. Altele conțineau troieni (programe malware ascunse) și backdoor-uri concepute pentru a oferi atacatorilor acces discret la sistemele victimelor', explică specialistul Eset, Phil Muncaster.

Un studiu separat, realizat în luna februarie a acestui an, a identificat peste 800 de aptitudini malițioase listate pe platforma ClawHub, ce conțineau programe infostealer, malware pentru furtul criptomonedelor și 'reverse shell'-uri care le permit hackerilor să preia controlul asupra dispozitivelor victimelor.

'Hackerii nu vizează doar AI skills, ci folosesc inteligența artificială și ca metoda de deghizare în contexte deja cunoscute, precum extensiile de browser. Microsoft a descoperit o campanie amplă în care extensii malițioase imitau asistenți AI legitimi. Au fost concepute pentru a colecta istoricul conversațiilor utilizatorilor, care poate conține informații sensibile, precum date de autentificare. Microsoft a identificat aproape un milion de descărcări doar în cadrul acestei campanii. Una dintre cele mai serioase amenințări pentru AI generativă și agenții AI este prompt injection. Funcționează în două moduri. Prompt injection direct presupune introducerea de către atacatori a unor instrucțiuni într-un model lingvistic mare (LLM), pentru a-l determina să ignore mecanismele de siguranță integrate. Însă, pentru agenții AI, mai relevant este prompt injection indirect. Instrucțiuni malițioase sunt ascunse în pagini web, postări, cod sau alte tipuri de conținut. Atunci când un agent interacționează cu acel conținut, acesta poate fi manipulat să efectueze activități malițioase', este de părere Muncaster.

În același timp, AI poate fi utilizată și pentru generarea în timp real a unor scripturi malițioase sau a unor componente de malware și ransomware.

În context, experții în securitate cibernetică vin cu o serie de recomandări pentru utilizatori, pentru a reduce riscurile de securitate asociate utilizării AI: precauție la orice instrument care pare să necesite mai multe permisiuni decât ar fi necesar; actualizări recente și remedierea problemelor de securitate; un semnal de alarmă apare atunci când pluginul trebuie să ruleze scripturi necunoscute sau să acceseze domenii/resurse suspecte în timpul instalării; monitorizarea agentului/pluginului și după instalare, pentru a identifica eventuale modificări de comportament realizate discret și malițios; menținerea întotdeauna a software-ului agentului AI actualizat, pentru a beneficia de cea mai sigură versiune disponibilă.

'În doar câțiva ani, inteligența artificială a devenit un instrument indispensabil pentru mulți utilizatori casnici. Acest fapt nu a trecut neobservat de infractorii cibernetici, care au dezvoltat diverse metode prin care pot transforma AI-ul utilizat zilnic într-un canal pentru distribuirea de malware, furtul de date și alte activități malițioase. În prezent, cele mai mari riscuri apar în domeniul emergent al agenților AI. Deși chatbotii sunt mai siguri decât în trecut, este important să nu renunțăm la măsurile de precauție. Din fericire, se conturează deja bune practici care pot contribui la utilizarea sigură a inteligenței artificiale. Combinarea lor cu funcții de securitate

oferite de furnizori de încredere poate contribui la menținerea unui nivel ridicat de protecție', se arata în articolul de specialitate.

În opinia specialiștilor, un chatbot AI legitim nu 'fura' datele utilizatorului, deși poate stoca și partaja cu alte părți informațiile introduse în conversații.

Eset a fost fondată în anul 1992 în Bratislava (Slovacia) și se situează în topul companiilor care oferă servicii de detecție și analiza a conținutul malware, fiind prezenta în peste 180 de țări.