

AI neumí generovat bezpečná hesla. Test Kaspersky ukázal alarmující výsledky

2.5.2025 - | PROTEXT

Špatnou správu hesel ještě zhoršuje spoléhání se na běžné kombinace jmen, slov obsažených ve slovnících a číslic. Taková hesla jde nejen relativně snadno dešifrovat, ale pokud kyberzločinec získá přístup k heslu na jednom webu, může mu to usnadnit přístup k celé řadě dalších účtů.

Lidé jsou vyzýváni k vytváření jedinečných náhodných hesel, aby se zabránilo zranitelnosti, kterou představuje opakované použití stejného hesla. Vytváření a správa hesel však může být náročný úkol. Lidé proto mohou podlehnout pokušení usnadnit si práci a používat k tomu velké jazykové modely (LLM), jako jsou ChatGPT, Llama nebo DeepSeek.

Vypadá to lákavě. Aby si lidé nemuseli lámat hlavu s vymyšlením silného hesla, mohou jednoduše zadat umělé inteligenci příkaz „Vygeneruj bezpečné heslo“ a okamžitě získat výsledek. AI vytváří řetězce, které se jeví jako náhodné, což pomáhá vyhnout se lidské tendenci používat předvídatelná hesla založená na slovech ze slovníku. Zdání však může klamat – hesla generovaná umělou inteligencí nemusí být tak bezpečná, jak byste očekávali.

Alexey Antonov, vedoucí týmu Data Science ve společnosti Kaspersky, to otestoval vygenerováním 1000 hesel pomocí některých z nejvýznamnějších a nejdůvěryhodnějších LLM – ChatGPT (od OpenAI), Llama (od skupiny Meta) a DeepSeek (nováček z Číny). „Všechny modely ví, že dobré heslo se skládá z nejméně 12 znaků, včetně velkých a malých písmen, číslic a symbolů. Hlásí to i při generování hesel,“ říká **Antonov**.

„DeepSeek a Llama někdy generovaly hesla skládající se ze slovníkových slov, ve kterých byly místo některých písmen číslice podobného tvaru (tzv. leetspeak): S@d0w12, M@n@g03, B@n@n@7 (DeepSeek), K5yB0a8dS8, S1mP1eL1on (Llama). Oba tyto modely rády generují heslo ‚password‘ ve tvaru P@ssw0rd, P@ssw0rd!23 (DeepSeek), P@ssw0rd1, P@ssw0rdV (Llama). Je zřejmé, že taková hesla nejsou bezpečná,“ dodává **Antonov**. Triky s nahrazováním písmen jsou známé a není těžké je rozluštit pomocí hrubé síly. ChatGPT tímto problémem netrpí a generuje hesla, která vypadají jako náhodná. Příklady:

Když je však podrobně prozkoumáte, můžete najít vzory. Setkáváme se třeba často s číslicí 9.

Při analýze tisícovky hesel generovaných různými modely se ukazuje, že žádný z nich nedosahuje ideálního rozložení znaků. U modelu ChatGPT dominují zejména písmena x, p, l a velké L, což nasvědčuje tomu, že místo skutečně náhodného řetězce hesla obsahují výrazně častěji tyto znaky. Podobně u Llamy se nejčastěji objevují symbol „#“ a opět písmena p, l a L, přičemž jejich zastoupení je sice poněkud vyrovnanější než u ChatGPT, ale stále daleko od ideálu.

Model DeepSeek vykazuje v podstatě stejné tendence – některé znaky se objevují mnohem častěji než jiné, což zvyšuje předvídatelnost vygenerovaných hesel.

Pro skutečně náhodný generátor hesel by však mělo platit, že každý znak (znaková sada, ze které se heslo skládá) je zastoupen přibližně stejně často. Teprve tak lze zajistit, že hesla nebudou vykazovat vzory, které by usnadnily jejich prolomení.

Algoritmy také často opomíjely vložit do hesla speciální znaky nebo číslice. Týkalo se to 26 % hesel u

ChatGPT, 32 % u Llama a 29 % u DeepSeek. DeepSeek a Llama také někdy generovaly hesla kratší než 12 znaků.

Se znalostí těchto závislostí mohou kyberzločinci výrazně urychlit uhádnutí hesla hrubou silou, protože se místo postupného zkoušení všech možností v pořadí „aaa“, „aab“, „aac“, ... „aba“, „abb“, „abc“, ... „zzz“, mohou zaměřit na časté kombinace vytvářené AI.

Antonov vyvinul v roce 2024 algoritmus strojového učení, aby otestoval sílu hesel, a zjistil, že téměř 60 % hesel lze uhádnout pomocí moderních GPU nebo cloudových nástrojů pro prolamování hesel za méně než hodinu. Při testování hesel generovaných umělou inteligencí byly výsledky alarmující – byla mnohem méně bezpečná, než se zdálo: 88 % hesel vygenerovaných modelem DeepSeek a 87 % hesel vygenerovaných modelem Llama nebylo dostatečně silných, aby odolaly útoku sofistikovaných nástrojů kyberzločinců. ChatGPT si v tomto vedl o něco lépe – testem společnosti Kaspersky neprošlo 33 % jím vytvořených hesel.

„Problém je v tom, že LLM nevytvářejí skutečně náhodné kombinace.“ Místo toho napodobují vzory z existujících dat, takže jejich výstupy jsou předvídatelné, pokud útočníci vědí, jak tyto modely fungují,” doplňuje **Antonov**.

Používejte bezpečnější správu hesel

Lidé by místo spoléhání na AI měli začít používat specializovaný software pro správu hesel, například Kaspersky Password Manager. Tyto nástroje nabízejí několik klíčových výhod.

Zprvé, tento typ softwaru obsahuje kryptograficky bezpečné generátory, které vytvářejí hesla bez detekovatelných vzorů, což zajišťuje skutečnou náhodnost. Za druhé, všechny přihlašovací údaje jsou uloženy v zabezpečeném trezoru, chráněném jediným hlavním heslem. Tím odpadá nutnost pamatovat si stovky hesel a zároveň se zabrání jejich ukořistění útočníky.

Správci hesel navíc umožňují automatické vyplňování přihlašovacích údajů a synchronizaci napříč zařízeními, což zefektivňuje přihlašování bez ohrožení zabezpečení. Mnohé z nich nabízejí také monitorování úniků dat a upozornění uživatelů, pokud jsou jejich přihlašovací údaje někde zveřejněny.

AI může pomoci s mnoha úkoly, generování hesel však mezi ně nepatří. Vzory a předvídatelnost hesel vytvořených pomocí LLM je činí zranitelnými vůči prolomení. Místo improvizací investujte do renomovaného správce hesel, který je vaší první linií obrany proti kybernetickým hrozbám. V době, kdy dochází k častým únikům dat, je nutné mít pro každý účet opravdu silné a jedinečné heslo.

<https://www.ceskenoviny.cz/tiskove/zpravy/ai-neumi-generovat-bezpecna-hesla-test-kaspersky-ukazal-alarmujici-vysledky/2668186>