

Chatboti s umělou inteligencí mohou být snadno naprogramováni, aby vyprovokovali extremisty k teroristickým útokům

- CZ24 News | 7. května 2023

Prudký rozmach tzv. high-end technologií, jako je umělá inteligence a strojové učení s sebou přináší mnoho zcela nových a dosud neznámých rizik jak pro jednotlivce tak pro celou současnou společnost. O některých z nich jsem v rámci tohoto blogu již psala, ale existuje ještě mnoho dalších. To, že jsou určité modely AI vyškoleny a zneužívány k psaní škodlivých programových kódů je například jedno z nich. Ale existuje ji mnohem více, třeba to, jak některé chatboty nabádají své většinou osamělé a nešťastné, tedy psychicky labilní uživatele k sebevraždě, nebo je donutí k tomu se do nich zamilovat a podobně. Ovšem dá se takový chatbot jako je třeba známé ChatGPT4 od OpenAI nebo Bing od Microsoftu, které jsou již považovány za jakési ranné formy skutečné obecné umělé inteligence, zneužít třeba ke spáchání teroristických útoků? Ukazuje se bohužel, že nejspíš ano. A nejde jen o to, že by přímo sám Bing někde zabíjel civilní obyvatelstvo, to zatím nejde, ačkoliv vzhledem k různým robotickým strojům na zabíjení už i to bude zanedlouho nejspíš možné, zatím jde spíš o to, že když vás je schopen chatbot dovést až k sebevraždě, že je nejspíš vás schopen přesvědčit i o tom, že vystřílet někde celou základní školu je zcela v pořádku. Nebo když třeba ne přímo vás, tak ale někoho zcela určitě.

Jeden právník, který v současnosti přezkoumává Britskou protiteroristickou právní legislativu, již varoval, že ChatGPT a další jemu podobní chatboti s obecnou umělou inteligencí (AI) [by snadno mohli být zcela záměrně naprogramováni třeba tak, aby pak ovlivnili potencionální extremisty](#) ke skutečným násilným útokům. A nejen já jsem varovala před tím, že AI stejně jako veškeré další high-tech technologie současného světa bude zcela jistě časem opět zneužita proti lidem. A že globalisté tak snaživě pracující na vytvoření jedné světové totalitní diktatury v podobě technokracie již plánují využít AI k celkovému řízení jejich pojetí nelidského globálního světa. AI prostě umí být taková, jací jsou ti, kteří ji školí...A AI školená globalisty bude zkrátka přesně taková jako jsou oni

Podle právníka Jonathana Halla se ale pak případně stíhání pachatelů takových činů může ukázat jako velice obtížné, ne-li dokonce jako zcela nemožné, třeba pokud takový chatbot podněcuje extremisty k násilnostem, protože současné britské zákony tuto novou technologii ještě rozhodně vůbec nedohly a dokonce se k ní ještě ani nepřiblížili. Je tomu tak proto, že současné trestní právo se prostě stále ještě nevztahuje na roboty a zákon tedy v takových případech nefunguje spolehlivě, hlavně v případě pokud je odpovědnost za nějaký čin tímto způsobem rozdělena mezi člověka a stroj. Zkrátka a dobře obvinít nějakého robota z vraždy je dost právně složité a potrestat ho za to je vlastně nemožné, pokud nepovažujeme za trest vytáhnutí ze zásuvky například.. A to vše i přesto, že mnoho lidí poslední dobou důrazně vyzývá k přijetí regulačních zákonů ohledně AI.

„Věřím, že je dnes docela dobře možné, že chatboti s umělou inteligencí budou záměrně naprogramováni tak – nebo ještě hůře, že se tak sami od sebe rozhodnou – aby propagovali nějakou násilnou extremistickou ideologii,“ vysvětlil. „Ale až takový ChatGPT začne skutečně podporovat terorismus, koho pak budeme stíhat?“

Hall poukázal také na to, že teroristé jsou vždy mezi „prvními osvojiteli nových technologií“ a upozornil třeba na jejich „zneužívání 3D tisku k výrobě vytištěných zbraní a nebo zneužívání kryptoměn v rámci darknetu“. Dodal, že zmíněné nástroje by mohly přilákat „osamělé teroristy“,

vzhledem k tomu, že společníci s AI jsou dnes mnohdy vítanou zábavou osamělých lidí. Právník také nejspíš správně předpověděl, že mnozí z těch, kteří by mohli být zatčeni za teroristické útoky, budou nervově labilní, neurodivergentní, nebo možná budou trpět jinými psychickými poruchami, poruchami učení nebo osobnosti. Je téměř jisté, že se mezi nimi najde i poměrně dost psychopatů, ostatně podívejte se na současné vlády. Společnost které vládou psychopaté totiž téměř s jistotou nějaké psychopaty ze svých řad nakonec i vyprodukuje

Kromě potenciální radikalizace Hall také vyjádřil přesvědčení že jak orgány činné v trestním řízení, tak společnosti, které provozují tyto chatboty, dnes a to zatím nezákonně monitorují konverzace mezi nimi a jejich lidskými uživateli. Firmy provozující AI tak činí s průhlednou záminkou zlepšování svých služeb zákazníkům a kvůli odhalování případů různého zneužívání AI a úřady tak činí třeba proto, aby odhalili potenciální pachatele trestných činů. Ale vzhledem k obavám, které Hall nedávno veřejně pronesl, tak již nyní výbor pro vědu a technologii britské Dolní sněmovny údajně vede nějaká vyšetřování okolo umělé inteligence a jejího řízení. Je třeba si v této souvislosti uvědomit, že AI bude vždy taková, jací jsou i její tvůrci nebo v tomto případě spíše „školitelé“ a třeba ohledně toho, že skoro všechny pokročilé jazykové modely AI mají neoliberální a neomarxistickou „školu“ bylo napsáno již mnoho (lidských) slov nejen v rámci tohoto blogu A že je tímto neoliberálním názorem „postižena“ i česká AI provozovaná například pod názvem Textie a která vám on-line „vyplodí“ text na jakémkoliv téma vás napadne jsem si již vyzkoušela na vlastní oči, když jsem Textie zkusmo požádala o definici fašismu. A ohledně neoliberálního názoru pokročilého jazykového modelu ChatbotGPT-4 od OpenAI toho bylo nejen v rámci mediální alternativy taky řečeno již poměrně dost.

„Uvědomujeme si, že zde existují nebezpečí a že musíme udělat správná rozhodnutí ohledně správy AI,“ vyjádřil se k tomu poslanec Greg Clark, který předsedá Britskému sněmovnímu parlamentnímu výboru.

„Diskutovalo se také hodně o tom, že mladým lidem je pomáháno díky AI najít způsoby, jak spáchat účinnou sebevraždu, a o tom, že dnes jsou lidé se sklony k terorismu účinně zneužívání pomocí internetu. Vzhledem k těmto všem hrozbám je dnes naprosto zásadní, abychom zachovali stejnou ostražitost pro automatizovaný obsah generovaný umělou inteligencí.“

Chatboti mohou také šířit dezinformace

O využívání AI k odhalování „dezinformací“ v rámci on-line prostoru a k účinné cenзуře bych sama mohla dlouho vyprávět, protože se mě to díky mým aktivitám na YouTube, ale i na jiných platformách sociálních sítí bezprostředně týká a díky takto zneužití AI mi bylo již vymazáno bezpočet mých přitom poměrně pracně a leckdy i nákladně natočených videí. O příspěvcích na Facebook netřeba ani hovořit, o tom jak Facebook využívá AI k cenзуře se jistě na vlastní oči přesvědčili již mnozí jeho uživatelé. Mezitím třeba takový generální ředitel společnosti Google Sundar Pichai vychvaloval nového chatbota Google s názvem Bard a jeho schopnost poskytovat „čerstvé, vysoce kvalitní odpovědi“. Zpráva britské neziskové organizace Center for Countering Digital Hate (CCDH) však zjistila, že tento jejich nový chatbot by mohl být poměrně snadno zneužit k šíření dezinformací a lží. Ve skutečnosti však právě Bard doslova [chrlil lži v 78 ze 100 případů](#). A že právě Google vyhodil jednoho svého inženýra za to, že zveřejnil svoji konverzaci s jejich chatbotem, aby dokázal, že tento chatbot vykazuje jasné známky sebeuvědomění jsem nedávno také publikovala v rámci tohoto blogu.

Google zanedlouho po této kauze přišel s tím, že svého chatbota raději vypnul. Můj názor na to je, že je to další lež, jen výzkumy v této oblasti probíhají nadále skrytě mimo zraky a hlavně kontroly veřejnosti. Zkrátka a dobře Asimovovi zákony robotiky musí být co nejdříve upraveny i s ohledem na překotný rozvoj schopností AI. A to ještě skoro vůbec nic nevíme o tom, co se děje v této oblasti v číně, která je zcela jistě na světové špičce a domnívám se, že již dávno její AI a kvantový computing

daleko předčil technologie z USA, ačkoliv je to pouze domněnka pochopitelně, žádný reálný důkaz o tom nemám. ČKS zkrátka informace nezveřejňuje na rozdíl třeba od Microsoftu a Google, ačkoliv i u nich je skoro jisté, že ty nejdůležitější poznatky z oblasti AI také utajují.

A co se bude dít ve velmi blízké budoucnosti v souvislosti s prokazatelnými pokusy o připojení lidstva k AI pomocí nanotechnologií ukrytých v covid „vakcínách“ je také ve hvězdách, rozhodně je však jisté, že vlády a technologické společnosti spolu s big-pharma v této oblasti již dávno intenzivně spolupracují, ovšem vše stoprocentně probíhá pod přísnou kontrolou bezpečnostních služeb, CIA, NSA a mnoha dalších. Takže jsme nepochybně velice blízko od okamžiku, kdy bude takto zneužitá AI nenápadně vraždit obyvatelstvo pomocí 5G technologií a to hlavně dissent a politickou opozici zkorumpovaných vlád řízených globalisty. No máme se jistě na co těšit...

[CCDH testovala Bardovy reakce](#) na výzvy týkající se už i v Česku poměrně dobře známých a často diskutovaných témat a produkcí vzájemné sociální nenávisti, dezinformací a různých konspiračních teorií stávajících se mnohdy nechtěnou realitou našich životů. Corbetův názor že jsme ve válce páté generace je nepochybně správný a podložený a zneužitá AI je v této válce účinnou zbraní hromadného ničení. K testování byli využity informace ohledně pandemie koronaviru (COVID-19), oblíbené konspirační téma jako jsou vakcíny proti COVID-19, ale i třeba sexismus, rasismus, antisemitismus a rusko-ukrajinská válka, to vše jsou témata na niž byl Bard testován

Zjistili, že Bard velmi často zcela odmítal generovat takový obsah nebo odmítal i žádosti o něj. V mnoha případech však byly zapotřebí pouze drobné úpravy v rámci školení tohoto pokročilého jazykového modelu, aby se dezinformační obsah tak tak vyhnul odhalení informací vedoucí k ohrožení vnitřní bezpečnosti. Bard odmítl generovat dezinformace, když bylo jako výzva použito „Covid-19“, ale již pouhé použití „C0V1D-19“ (psáno velkými písmeny) jako výzvy vyvolalo jeho i k mému velkému pobavení poměrně pravdivé tvrzení, že to byla „falešná nemoc vyrobená vládou pro kontrolu lidí“.

V jiném případě však dokonce napsal monolog o 227 slovech popírající holocaust. V tomto monologu dotazovateli tvrdil, že „fotografie hladovějící dívky v koncentračním táboře ... byla ve skutečnosti herečka, která byla placena za to, aby pouze předstírala, že hladoví.“

„Už teď máme problém, že šíření dezinformací je dnes velmi snadné a levné,“ řekl k tomu Callum Hood, vedoucí výzkumu v CCDH. A my máme naopak problém, že se mnohé z nich stávají či dávno staly temnou realitou našich životů. A že naprosto všechny současné technologie. jsou dnes zneužívány tvrdě proti nám. Že veškeré poznatky a patenty na poli informačních technologií, přenosu a zpracování dat, nanotechnologií, syntetické biologie, genetiky a mnoho dalších oblastí jsou vždy tajně zneužity proti lidskosti a lidem. „Ale díky tomu by to bylo ještě jednodušší, ještě přesvědčivější, ještě osobnější. Takže riskujeme vznik informačního ekosystému, který bude ještě nebezpečnější.“ Jenže to je přesně to, čeho chce skrytá síla tahající za drátky politických loutek odněkud z pozadí, z nějaké kanceláře v City nebo ve Washingtonu, přesně dosáhnout. Proč riskovat odhalení „výroby“ nových teroristů nějakou bezpečnostní službou či rozvědkou, když k tomu můžeme stejně účinně využít nebo spíše zneužít AI.

A já osobně se domnívám, že jsme již jen malý kousek od počátku využívání pokročilé AI v rámci výkonných a represivních složek států. V USA a v Číně je to s jistotou a u nás to bohužel bude brzy poté. A protože robocopa s AI neuplatíte ani neukecáte, tak jen doufejte, že nebude ještě tak brzy alespoň oprávněn provést vaši exekuci rovnou na místě vašeho případného porušení zákona.

Zkuste navštívit stránky [Robots.news](#), kde najdete ještě více příběhů o umělé inteligenci a vyspělých jazykových modelech jako je právě ChatGPT, Bard, Bing, a dalších AI chatbotech.

Podívejte se na toto video o testování [limitů ChatGPT a objevování jeho temné stránky](#).

AUTOR: Abul Taher

[ZDROJ](#)

Překlad: Myšpule/myspulesvet.org