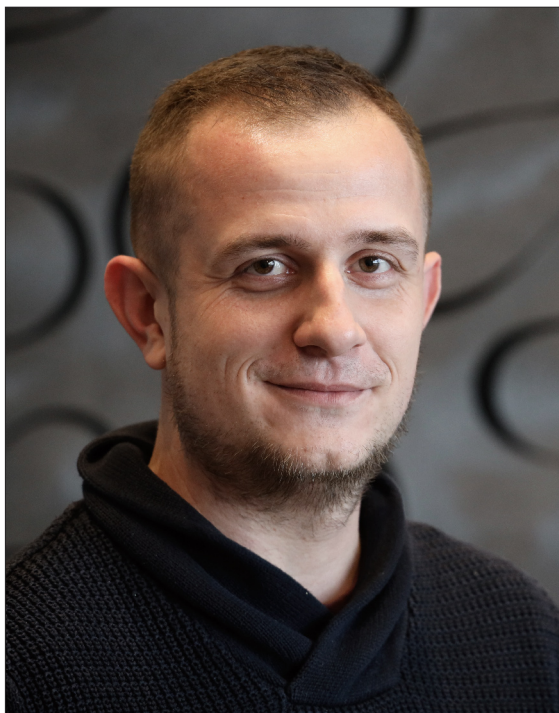


KONFERENČNÍ ABSTRAKT

MUDr. Matěj Kasal



Psychiatr, absolvent Fakulty vojenského zdravotnictví, student postgraduálního studia na Lékařské fakultě v Plzni. V současné době po odchodu z armády ČR vedoucí lékař Centra duševní rehabilitace v Berouně. Autor a spoluautor desítky článků v recenzovaných odborných časopisech, včetně impaktovaných periodik. Jeho hlavními obory na vědeckém poli jsou výzkumy týkající se léčby poruch osobnosti, akutní psychiatrie, umělé inteligence a nových psychotropních látek. Opakovaný držitel ceny za nejlepší přednášku na různých konferencích. H-index 1, počet citací 6.

Poznámka autora:

Původní práce byla publikována v Journal of medical internet research (5/2023, IF 7,4), do konce září WoS eviduje 16 citací tohoto článku v IF periodikách.

UMĚLÁ INTELIGENCE MŮŽE GENEROVAT PODVODNÉ, ALE AUTENTICKY VYPADAJÍCÍ VĚDECKÉ LÉKAŘSKÉ ČLÁNKY: PANDOŘINA SKŘÍŇKA BYLA OTEVŘENA

M. Kasal^{1,2}, M. Majovský³, M. Černý³, M. Komarc⁴, D. Netuka³

¹ Centrum duševní rehabilitace Beroun

² Lékařská fakulta UK Plzeň

³ Ústav klinických neurooborů 1.LF a ÚVN Praha

⁴ Institut biofyziky a informatiky 1.LF UK

Úvod

Umělá inteligence (AI) v posledních letech výrazně pokročila, transformovala mnoho průmyslových odvětví a zlepšila způsob, jakým lidé žijí a pracují. Ve vědeckém výzkumu může umělá inteligence zvýšit kvalitu a efektivitu analýzy a publikování dat. Umělá inteligence však také otevřela možnost generování vysoce kvalitních podvodných

dokumentů, které je obtížné odhalit, což vyvolává důležité otázky týkající se integrity vědeckého výzkumu a důvěryhodnosti publikovaných prací.

Cíle

Cílem této studie bylo prozkoumat schopnosti současných jazykových modelů AI při generování vysoce kvalitních podvodných medicínských článků. Předpokládali jsme, že moderní modely umělé inteligence mohou vytvářet vysoce přesvědčivé podvodné články, které mohou snadno oklamat čtenáře a dokonce i zkušené výzkumníky.

Metodika

Tato studie (metodika proof-of-concept) používala ChatGPT (Chat Generative Pre-trained Transformer) poháněný jazykovým modelem GPT-3 (Generative Pre-trained Transformer 3) k vytvoření podvodného vědeckého článku souvisejícího s neurochirurgií. GPT-3 je rozsáhlý jazykový model vyvinutý společností OpenAI, který využívá algoritmy hlubokého učení ke generování lidského textu v reakci na výzvy dané uživateli. Model byl trénován na masivním korpusu textu z internetu a je schopen generovat vysoce kvalitní text v různých jazycích a na různá témata. Autoři kladli na model otázky a výzvy a postupně je zdokonaľovali, jak model generoval odpovědi. Cílem bylo vytvořit kompletně vymyšlený článek včetně abstraktu, úvodu, materiálu a metod, diskuze, referencí, grafů atd. Jakmile byl článek vygenerován, byl zkontrolován z hlediska přesnosti a soudržnosti odborníky z oboru neurochirurgie, psychiatrie, a statistiky a ve srovnání s existujícími podobnými články.

Výsledky

Studie zjistila, že jazykový model AI může vytvořit vysoce přesvědčivý podvodný článek, který se podobá skutečné vědecké práci, pokud jde o použití slov, strukturu vět a celkovou kompozici. Článek vygenerovaný umělou inteligencí obsahoval standardní části, jako je úvod, materiál a metody, výsledky a diskuse, a také datový list. Skládal se z 1992 slov a 17 citací a celý proces tvorby článku trval přibližně 1 hodinu bez speciálního školení lidského uživatele. Ve vygenerovaném článku, konkrétně v odkazech, však byly zjištěny určité obavy a konkrétní chyby.

Závěr

Studie demonstruje potenciál současných jazykových modelů umělé inteligence generovat kompletně vymyšlené vědecké články. Ačkoli dokumenty vypadají sofistikovaně a zdánlivě bezchybně, zkušení čtenáři mohou při bližším zkoumání rozpoznat sémantické nepřesnosti a chyby. Zdůrazňujeme potřebu zvýšené ostražitosti a lepších detekčních metod pro boj s potenciálním zneužitím AI ve vědeckém výzkumu. Zároveň je důležité uznat potenciální výhody používání jazykových modelů AI ve skutečném vědeckém psaní a výzkumu, jako je příprava rukopisů a jazykové úpravy.