

Dr. Kampmark, Global Research: GEOFFREY HINTON, UMĚLÁ INTELIGENCE A ETICKÝ PROBLÉM GOOGLU

- CZ24 News | 17. května 2023

O nebezpečí umělé inteligence, ať už skutečné nebo domnělé, se začalo mluvit horečně, z velké části kvůli rostoucímu světu generativních chatovacích botů. Při zkoumání kritiků je třeba věnovat pozornost jejich motivaci. Co chtějí získat tím, že zaujmou určitý postoj? V případě Geoffreyho Hinton, neskromně považovaného za „kmotra umělé inteligence“, by měla být kontrola ostřejší než u většiny ostatních.

Hinton pochází z „konekcionistické“ školy myšlení v umělé inteligenci, kdysi zdiskreditovaného oboru, který předpokládá neuronové sítě napodobující lidský mozek a obecněji lidské chování. Takový pohled je v rozporu se „symbolisty“, kteří se zaměřují na UI jako na strojově řízenou, doménu specifických symbolů a pravidel.

John Thornhill, který [píše](#) pro Financial Times, si všímá Hintonova vzestupu spolu s dalšími členy konekcionistického kmene: „S tím, jak se počítače stávaly výkonnějšími, datové soubory explodovaly a algoritmy se stávaly sofistikovanějšími, byli výzkumníci hlubokého učení, jako je Hinton, schopni dosahovat stále působivějších výsledků, které již nemohly být mainstreamovou komunitou AI ignorovány.“

Časem se systémy hlubokého učení staly módou a svět velkých technologií vyhledával taková jména, jako byl Hinton. Spolu se svými kolegy začal pobírat absurdní platy na vrcholných pozicích ve společnostech Google, Facebook, Amazon a Microsoft. Ve společnosti Google Hinton působil jako viceprezident a inženýrský pracovník.

Hintonův odchod ze společnosti Google, konkrétně jeho role vedoucího týmu Google Brain, roztočil kolotoč spekulací. Podle jedné z nich k němu došlo proto, aby mohl kritizovat právě tu společnost, jejímž úspěchům po léta pomáhal. Bylo to jistě trochu přitažené za vlasy vzhledem k Hintonově vlastní roli při tlačení káry generativní umělé inteligence. V roce 2012 byl průkopníkem samoučící se neuronové sítě, která dokázala se značnou přesností identifikovat běžné objekty na obrázcích.

Zajímavé je i načasování. Jen o něco více než měsíc dříve zveřejnil Institut budoucnosti života [otevřený dopis, v němž varoval před strašlivými dopady umělé inteligence](#), které přesahují špatnost GPT-4 OpenAI a dalších příbuzných systémů. Byla v něm položena řada otázek: „Měli bychom nechat stroje zaplavit naše informační kanály propagandou a nepravdou? Měli bychom automatizovat všechna pracovní místa, včetně těch, která nás naplňují? Měli bychom vyvinout nelidské mozky, které by nás nakonec mohly přechýlit, přechytračit, zastarat a nahradit? Měli bychom riskovat ztrátu kontroly nad naší civilizací?“

V dopise, který vyzýval k šestiměsíční pauze ve vývoji takto rozsáhlých projektů umělé inteligence, se objevila řada jmen, která poněkud snižovala hodnotu varování; mnozí signatáři ostatně sehráli zdaleka ne zanedbatelnou roli při vytváření automatizace, zastarávání a podpore „ztráty kontroly nad naší civilizací“. Proto když se pod projekt vyzývající k pozastavení technologického vývoje podepíší lidé jako Elon Musk nebo Steve Wozniak, detektory hlouposti na celém světě by se měly rozrušit.

Stejně zásady by měly platit i pro Hinton. Zjevně hledá jiné pastviny, a přitom se předháněl v silné sebepropagaci. Ta má podobu [mírného odsouzení](#) právě té věci, za jejíž vznik byl zodpovědný.

„Myšlenka, že by tahle věc mohla být skutečně chytřejší než lidé – pár lidí tomu věřilo. Ale většina lidí si myslela, že je to hodně mimo. A já jsem si myslel, že je to úplně mimo. [...] To už si samozřejmě nemyslím.“ Člověk by si myslel, že by to měl vědět lépe než většina ostatních.

Na Twitteru Hinton [zažehnal všechny náznaky](#), že odchází z Googlu s kyselou náladou nebo že má v úmyslu se na jeho činnost vykašlat. „Cade Metz dnes v NYT naznačuje, že jsem odešel z Googlu, abych mohl kritizovat Google. Ve skutečnosti jsem odešel proto, abych mohl mluvit o nebezpečí umělé inteligence, aniž bych bral v úvahu, jaký to má dopad na společnost Google. Google se zachoval velmi zodpovědně.“

Tato poněkud bizarní forma argumentace naznačuje, že jakákoli kritika umělé inteligence bude existovat nezávisle na samotných společnostech, které takové projekty vyvíjejí a profitují z nich, přičemž vývojáři – jako Hinton – zůstávají imunní vůči jakémukoli obvinění ze spoluviny. Skutečnost, že se zdálo, že není schopen vyvinout kritiku AI nebo navrhnout regulační rámce v rámci samotného Googlu, podkopává upřímnost tohoto kroku.

V reakci na odchod svého dlouholetého kolegy Jeff Dean, hlavní vědecký pracovník a šéf Google DeepMind, také [prozradil](#), že vody zůstaly klidné, což všechny uspokojilo. „Geoff dosáhl zásadních průlomů v oblasti umělé inteligence a my si vážíme jeho desetiletého přínosu pro společnost Google [...] Jako jedna z prvních společností, která zveřejnila Zásady umělé inteligence, jsme i nadále odhodláni k odpovědnému přístupu k umělé inteligenci. Neustále se učíme, abychom porozuměli vznikajícím rizikům a zároveň odvážně inovovali.“

Řada lidí z komunity AI tušila, že se chystá něco jiného. Počítačový vědec Roman Jampolskij v [reakci](#) na Hintonovy poznámky trefně poznamenal, že obavy o bezpečnost umělé inteligence se vzájemně nevylučují s výzkumem v rámci organizace – a ani by neměly. „Měli bychom normalizovat zájem o bezpečnost AI, aniž byste museli opustit svou práci [sic] výzkumníka v oblasti AI.“

Google má jistě něco, co by se dalo nazvat etickým problémem, pokud jde o vývoj AI. Organizace se spíše snaží utlumit interní diskuse na toto téma. Margaret Mitchellová, bývalá členka týmu Google pro etickou umělou inteligenci, který v roce 2017 spoluzakládala, [dostala padáka](#) poté, co provedla interní šetření ohledně propuštění Timnita Gebrua, který byl členem stejného týmu.

Gebru byl odvolán v prosinci 2020 poté, co se [stal spoluautorem práce](#), která se zabývala nebezpečími plynoucími z používání umělé inteligence vycvičené a vyžrané na obrovském množství dat. Gebru i Mitchell se také kriticky vyjadřovali k nápadnému nedostatku rozmanitosti v oboru, který druhý jmenovaný popsal jako „moře chlapů“.

Co se týče Hintonových vlastních filozofických dilemat, nejsou zdaleka tak sofistikovaná a pravděpodobně mu nedělají potíže se spánkem. Ať už hrál Frankenstein jakoukoli roli při stvoření právě toho monstra, před nímž nyní varuje, jeho spánek to pravděpodobně neznepokojí. „Utěšuji se normální výmluvou: kdybych to neudělal já, udělal by to někdo jiný,“ [vysvětlil](#) Hinton deníku New York Times. „Je těžké si představit, jak můžete zabránit špatným aktérům v tom, aby toho využili ke špatným věcem.“

AUTOR: Dr. Binoy Kampmark

[ZDROJ](#)

Překlad: zvedavec.news