

Rizika umělé inteligence

- CZ24 News | 13. března 2022

O umělé inteligenci se i v souvislosti s Covidem hovoří stále častěji a fenomén AI – artificial intelligence, umělé inteligence se již bezprostředně dotýká každého z nás. Pro mnohé je to však stále jen hlas umělého asistenta ve smartfonu, nebo nějaká „vychytaná apka“. Jen málo lidí však ví, co to ta AI skutečně je a jen úplná hrstka si uvědomuje například to, že každá pravá UI je převším posedlá touhou po absolutní kontrole úplně všeho. A že je nekompromisně totalitní. A když ji člověk ozbrojí, což pochopitelně dříve či později udělá, ehmmm, již to udělal, tak že může být i značně nebezpečná.

Rychlý vývoj umělé inteligence může představovat určitá rizika, a proto by měla být přísně a eticky kontrolována. Vývojáři umělé inteligence nyní diskutují o tom, jak lze [omezit technologii tak](#), aby zajistili, že roboti budou vždy jednat v zájmu lidstva, protože dát jim vlastní nezávislé osobnosti by mohlo být nebezpečné.

Konzultant AI Matthew Kershaw řekl, že je dokonce možné, že technologie ještě během tohoto života dosáhne až znepokojivých výšin, „pokud jsme; dostatečně mladí“.

Profesor Stephen Hawking, teoretický fyzik a kosmolog, před svou smrtí v roce 2018 řekl: „Primitivní formy umělé inteligence, které již máme, se ukázaly jako velmi užitečné. Ale myslím si, že vývoj plně umělé inteligence by mohl znamenat konec lidské rasy.“ (Související: [EU navrhuje legislativu omezující technologii rozpoznávání obličeje a „vysoce rizikové“ aplikace umělé inteligence](#).)

Zakladatel SpaceX Elon Musk souhlasil s Hawkingovým prohlášením. Řekl: „Myslím, že bychom měli být velmi opatrní ohledně umělé inteligence. Pokud bych měl hádat, co je naší největší existenční hrozbou, je to pravděpodobně ono. Takže musíme být velmi opatrní.“

Co je to umělá obecná inteligence?

AI sama o sobě může být nebezpečná, ale co můžeme od AGI čekat? Zjednodušeně lze umělou obecnou inteligenci definovat jako [schopnost stroje vykonávat jakýkoli úkol](#), který dokáže člověk. I když to zdůrazňuje schopnost strojů vykonávat úkoly s větší účinností než lidé, nejsou ještě od tohoto okamžiku obecně inteligentní. Například jsou velmi dobří v jedné funkci, zatímco ale nemají vůbec žádné schopnosti v ničem jiném, co by do nich nebylo naprogramováno. I když jsou tedy při provádění jednoho úkolu efektivní jako sto trénovaných lidí, mohou prohrát třeba s dětmi doslova v čemkoliv jiném.

Kershaw však věří, že skutečná AGI bude vyžadovat použití počítačů, které jsou dostatečně výkonné na to, aby vytvořili komplexní model světa, a to nejspíš nebude v dohledné době. „Vzhledem k tomu, že my sami opravdu ještě nerozumíme tomu, co to znamená být při vědomí, myslím si, že je velmi nepravděpodobné, že by se AGI v dohledné době stal realitou.“ Jen stále nevíme, co to vlastně znamená být „při vědomí“.

Kromě toho musí být každá umělá inteligence trénována v jakékoli funkci s využitím obrovského množství dat, zatímco lidé se mohou učit s podstatně méně. „Lidské dítě nepotřebuje vidět více než pět aut, aby se naučilo rozpoznat auto. Počítač by jich potřeboval vidět tisíce,“ řekl Kershaw.

V roce 1950 průkopník výpočetní techniky Alan Turing navrhl, že o počítači lze říci, že má umělou

inteligenci, pokud dokáže za určitých podmínek napodobit lidské reakce. Turing vymyslel test, který měl určit, zda se umělá inteligence může nebo nemůže zaměnit za člověka, a zatím neexistuje žádný systém, který by jím prošel, i když několik málo z nich se tomu již přiblížilo. Nejblíže to bylo v roce 2018, kdy duplexní AI Googlu sama zatelefonovala do kadeřnického salonu a úspěšně si domluvila schůzku.

Duplex však pracoval na velmi specifickém úkolu a opravdový AGI by si byl schopen popovídat s kadeřníkem.

Umělá inteligence je dnes všude, ale i když jsou tyto systémy ve svých specializovaných misích dobré, žádný z nich se dosud nebyl schopen naučit cokoli skutečně udělat bez pomoci lidí.

To, co vědci nazývají „umělá obecná inteligence“, zatím zůstává pouze teoretické. Aby tyto umělé systémy mohly fungovat, musí se stroje samy učit ze zkušeností, přizpůsobovat se novým vstupům a vykonávat úkoly jako to dělají lidé.

Odborníci vyvíjejí nové technologie, které se mohou blížit inflexnímu bodu, kde mohou rozvíjet obecnou inteligenci. Většina odborníků na umělou inteligenci však stále věří, že pravou AGI budeme schopni vytvořit již do konce století, přičemž neoptimističtější odhady jsou kolem roku 2040 až 2080. Jiní lidé také věří, že umělého vědomí nebude nikdy dosaženo, protože my, jako lidé, dosud nerozumíme ani tomu svému.

Rizika umělé inteligence

Umělá inteligence má také nevýhody, z nichž některé zahrnují následující:

Ztráta zaměstnání. Automatizace pracovních míst může [urychlit masivní ztrátu pracovních míst](#). Automatizace práce je bezprostředním problémem. Už nejde o to, *jestli* umělá inteligence může nahradit určité typy úloh, jde o to, do jaké míry je mohou nahradit. Mnoho průmyslových odvětví se již dnes zcela automatizuje, zejména v těch oblastech, kde se úkoly opakují. Zanedlouho bude možné pomocí AI provádět úkoly od maloobchodního prodeje přes analýzu trhu až po vlastní manuální práci.

Obavy o soukromí. Škodlivé použití umělé inteligence by také mohlo [ohrozit digitální bezpečnost](#) prostřednictvím hackingu, fyzickou bezpečnost vyzbrojováním spotřebitelských dronů a dokonce i politickou bezpečnost prostřednictvím sledování a profilování. Umělá inteligence může ovlivnit soukromí a bezpečnost podobně jako čínské „orwellovské“ používání obličejových technologií v kancelářích, školách a na místech.

Nestabilita akciového trhu. Díky algoritmickému vysokofrekvenčnímu obchodování by mohly být svrženy celé finanční systémy. Algoritmické obchodování nastává, když počítač může provádět obchody na základě předem naprogramovaných instrukcí a může provádět velkoobjemové, vysokofrekvenční a vysoce hodnotné obchody, které mohou vést k extrémní volatilitě trhu.

ZDROJ: transhumanism.news / myspulesvet.org

Překlad a úprava: Myšpule