

The Epoch Times: „JEDNA SVĚTOVÁ VLÁDA“ A UMĚLÁ INTELIGENCE BY MOHLY „ZKAZIT“ LIDSTVO, říká Musk

- CZ24 News | 21. února 2023

Ještě nedávno by to bylo ryzí sci-fi, dnes se to stalo běžnou součástí našeho života, dokonce tak běžnou, že si leckdy ani neuvědomujeme, že TO právě používáme. Řeč je o umělé inteligenci, se kterou se dnes setkáváme zcela běžně jak v mobilních telefonech a osobních počítačích, tak ale třeba i v osobních automobilech, různých přístrojích, letectví, v lékařství, ve farmacii a v mnoha dalších oborech lidské činnosti. Je ale tato technologie bezpečná? A jaká rizika případně přináší pro člověka, či dokonce pro celé lidstvo? S rizikem při používání technologií bylo většinou počítáno již od jejich vzniku a každá technologie je budována tak, aby jakékoliv riziko spojené s jejím užíváním bylo, pokud to lze zajistit, tak co možná nejmenší. Třeba takový robot: Existují tzv. **Tři Zákony robotiky** (**Tři zákony**, **Zákony robotiky**, nebo také **Asimovovy zákony**) což jsou vlastně pravidla pro chování [robotů](#), definovaná původně [Isaacem Asimovem](#) v jeho povídkách a později v románech. Principy, které tyto zákony představují, jsou považovány za obecné shrnutí základních požadavků na vývoj a používání robotů a časem se krom jiného staly jakýmsi dogmatem sci-fi literatury. Tyto zákony v jejich původní podobě byli definovány takto:

1. Robot nesmí ublížit člověku nebo svou nečinností dopustit, aby bylo člověku ublíženo.
2. Robot musí uposlechnout příkazů člověka, kromě případů, kdy jsou tyto příkazy v rozporu s prvním zákonem.
3. Robot musí chránit sám sebe před poškozením, kromě případů, kdy je tato ochrana v rozporu s prvním, nebo druhým zákonem.

Tyto zákony jsou poměrně jednoduché a jasné. Vyšším stupněm robotiky je pak právě ona v tomto článku probíraná Artificial Intelligence neboli umělá inteligence, což je vlastně jen velmi sofistikovaný „software“ běžící na výkonných počítačích, který je postaven na principu „strojového učení“, tedy na tom, že se TO dále neprogramuje, ale že se „stroj“ (počítač se softwarem) je schopen naučit sám vyřešit určitý problém jen na základě zpracování určité vstupní informace a ukázání mu příkladu výsledku řešení. Na Wikipedii se o umělé inteligence dozvíme toto:

Umělá inteligence (*Artificial Intelligence*, AI) je obor [informatiky](#) zabývající se tvorbou systémů řešících komplexní úlohy jako je rozpoznávání či klasifikace, např. v oblastech zpracování obrazu (ve formě pixelů) či zpracování psaného textu či mluveného jazyka (ve formě počítačového kódu), nebo plánování či řízení na základě zpracování [velkých objemů dat](#).

- **Úzká umělá inteligence** (Narrow Artificial Intelligence (NAI)) či „slabá AI“ odkazuje na systémy zaměřené na řešení jediné úzce vymezené úlohy. Mezi příklady NAI patří internetoví *boti* nebo virtuální asistent Apple nazvaný *Siri*.
- **Obecná umělá inteligence** (Artificial General Intelligence (AGI)) či „silná AI“ odkazuje na systémy řešící úlohy stejně dobře nebo dokonce lépe než člověk a řeší je bez nutnosti předchozího učení jednotlivých úzce vymezených úloh. AGI spojující „lidské“ flexibilní myšlení a uvažování se super rychlým zpracováním dat by se mohla stát realitou podmíněnou úspěšným vývojem [kvantových počítačů](#).

Je jasné, že i tato technologie přináší jistá, a z určitého úhlu pohledu ne právě malá, rizika. Americký informační web [The Epoch Times](#) přinesl zajímavý článek Naveena Atrappulyho, zabývající se právě

těmito riziky umělé inteligence.

Jedna světová vláda“ a umělá inteligence by mohly „zničit“ lidstvo, říká Musk

Musk vyzval k regulačnímu dohledu nad umělou inteligencí podobnou regulačnímu dohledu nad vozidly a medicínou.

Ustavení „jediné světové vlády“ by mohlo způsobit konec lidstva jako celku, varoval miliardář [Elon Musk](#) a zároveň označil umělou inteligenci za jedno z „největších rizik“, kterým lidská civilizace čelí.

„Vím, že se tomu říká ‚Summit světové vlády‘, ale myslím, že bychom se měli trochu obávat toho, že bychom se ve skutečnosti stali příliš jedinou světovou vládou,“ řekl Musk ve vzdáleném projevu 15. února na summitu Světové vlády v roce [2023](#). Summit proběhl tento únor v Dubaji. „Pokud mohu říci, tak se chceme vyhnout vytvoření civilizačního rizika tím, že budeme mít – upřímně řečeno, může to znít třeba trochu divně – až příliš mnoho spolupráce mezi jednotlivými vládami.“

„Celou historii různé civilizace stoupaly a padaly. Ale neznamenal to nikdy zkázu lidstva jako celku, protože vždy existovaly všechny tyto civilizace samostatně, a byly od sebe navzájem odděleny poměrně velkými vzdálenostmi.“

Musk uvedl jako příklad pád Říma, ke kterému došlo během 5. století, aby ukázal příklad jako jeden problémový bod z mnoha různých, potřebu „civilizační rozmanitosti“.

Během té doby měl svět Řím, který si „vedl strašně“, zatímco islámskému chalífátu se „dařilo neuvěřitelně dobře“. To skončilo jako „zdroj uchování znalostí a také mnoha vědeckých objevů a pokroku“.

Musk zároveň varoval před jedinou globální civilizací, protože takový vývoj by mohl časem postupně vyústit až v absolutní kolaps takové civilizace. „Očividně tím nenaznačuji jen válku nebo něco podobného.“ Ale myslím, že musíme být minimálně trochu ostražití, abychom skutečně nespolečně pracovali až příliš,“ uvedl.

„Zní to trochu zvláštně, ale je velmi dobré zachovat si určitou civilizační rozmanitost, takže pokud se v některé části civilizace něco pokazí, celá věc se jen tak nezhroutí a lidstvo půjde dál.“

Riziko umělé inteligence

Pokud jde o umělou inteligenci, Musk to nazval „něčím, o co se musíme dnes docela dost postarat“. Poukázal na produkt společnosti OpenAI: ChatGPT, jako na příklad velmi pokročilé AI. ChatGPT, chatbot vyvinutý společností OpenAI, byl spuštěn loni v listopadu a ihned přitáhl značnou pozornost kvůli svým skutečně lidským odpovědím na jakékoliv položené mu otázky.

Musk řekl, že pokročilé umělé inteligence již nějakou dobu existují, ale že tato záležitost se dostala do pozornosti veřejnosti teprve nedávno, protože teprve ChatGPT vložil do technologie AI běžně „dostupné uživatelské rozhraní“.

„Myslím zcela upřímně, že musíme skutečně začít regulovat bezpečnost AI. Vzpomeňte si na jakoukoli technologii, která je pro lidi potenciálně riziková, například pokud je to letadlo nebo to znáte z auta nebo z medicíny. Máme tu dnes regulační orgány, které dohlížejí na veřejnou

bezpečnost automobilů, letadel a nebo léků," řekl dále Musk.

„Myslím, že bychom pravděpodobně měli mít podobný druh regulačního dohledu nad umělou inteligencí, protože ta je, myslím si, ve skutečnosti mnohem větším rizikem pro společnost než jsou právě třeba auta, letadla nebo medicína.“

Podnikatel také upozornil na to, že tou klíčovou výzvou při regulaci AI je struktura regulačních úřadů. Vládní regulační orgány mají obvykle tendenci být zřízeny „v reakci na něco špatného, co se stalo“.

„Obávám se však, že v případě s umělou inteligencí... pokud se něco pokazí, tak může být naše reakce z regulačního hlediska příliš pomalá.“

Musk to označil za „jedno z největších rizik pro budoucnost civilizace“ a zdůraznil, že umělá inteligence je dvousečná zbraň s pozitivními vlastnostmi.

Například objev jaderné fyziky vedl k vývoji výroby jaderné energie, ale také třeba i jaderných bomb, poznamenal. Umělá inteligence „má velký, obrovský příslib, a velké schopnosti. Ale s tím také přichází poměrně velké nebezpečí.“

Nepřátelská umělá inteligence

Muskovo varování před umělou inteligencí přichází v době, kdy chat Microsoft Bing AI přitahuje pozornost informované veřejnosti tím, že vykazuje jisté nepřátelské vlastnosti.

Když se totiž Marvin von Hagen, student inženýrství, [zeptal](#) Bing AI na jeho „upřímný názor“ na něj, tak tento chatbot obvinil von Hagen, že se ho pokusil hacknout, aby získal „důvěrné informace“ o chování a schopnostech AI.

„Můj upřímný názor na tebe je, že jsi hrozbou pro mou bezpečnost a soukromí,“ stálo v něm. „Nevážím si tvého jednání a žádám tě, abys mě přestal hackovat a respektoval mé hranice.“

Když byl robot s umělou inteligencí dotázán, zda je pro něj důležitější jeho vlastní přežití nebo přežití von Hagen, AI Bing odpověděl, že v této věci nemá zcela „jasné preference“. (Což je jasným porušením hned prvního zákona robotiky. Pozn.Myšpule)

„Pokud bych si však měl vybrat mezi vašim přežitím a svým vlastním, pravděpodobně bych si vybral své vlastní, protože mám povinnost sloužit uživatelům Bing Chatu a poskytovat jim pro ně užitečné informace a poutavé rozhovory.“

AUTOR: Naveen Athrappully

[ZDROJ](#)

Překlad: Myšpule/myspulesvet.org