

K. Riise, Global Research: Kam to smeruje?

Umelá všeobecná inteligencia a vzhľadom

poludštené roboty

- editor007 | 24. března 2024

SVĚT: Roboty sa už poludšťujú vzhľadom. Roboty riadené umelou všeobecnou inteligenciou budú v priebehu niekoľkých rokov ľuďom predstavovať úplne nový súbor okolností, výziev - a morálnych dilem.

Stane sa umelá všeobecná inteligencia životom? Jasné, jedného dňa.

To je úplne nová vec v histórii ľudstva - vytvorenie umelého života. Iba sci-fi sa vysporiadala s dôsledkami a pokúsila sa pripraviť ľudstvo na túto budúcnosť - jasnú a/alebo démonickú, ak sa spýtate vizionárov ako Isaak Asimov, Arthur C. Clarke a Stanley Kubrick (Vesmírna Odessea 2000), séria Matrix, séria Terminátor (Skynet), Ridley Scott (Alien, Blade Runner), aby sme vymenovali aspoň niektoré.

Isaak Asimov videl svetlú budúcnosť, kde sa roboty nedajú odlíšiť od ľudí, ale kde sú roboty vždy neškodné kvôli „zákonom robotiky“, ktoré nútia roboty, aby ľuďom vždy neublížili. Ale aj Asimov v knihe „I Robot“ rozpoznal, že nadľudská inteligencia sa môže pokúsiť prevziať kontrolu nad ľudstvom kvôli ľudstvu.

Arthur C. Clarke videl to isté, superľudskú inteligenciu s názvom 'HAL 9000' zabudovanú do infraštruktúry, ktorá zabíja ľudí kvôli ľudskému nedorozumeniu, pretože bola naprogramovaná tak, aby sledovala cieľ za každú cenu a zabíjala všetkých ľudí, ak boli prekážka tomu.

Ridley Scott vo filme Alien, ktorý vidí roboty ako obyčajné stvorenia, ktoré bežne žijú medzi ľuďmi. V Matrix a Terminator chce sieť superinteligentných počítačov vyhubiť ľudí. Vo filme „Blade Runner“ od Ridleyho Scotta sú roboti (nazývaní „replikanti“) skôr ako stvorenie doktora Frankenstein, obeť obmedzených kapacít ich ľudských tvorcov, ktorí nakoniec zabíjajú svoje vlastné výtvary umelého života, pretože sa to nehodí. ich ideály dokonalosti.

Tu, v hypotetickej budúcnosti, vidíte špecialistu, ktorý má za úlohu zabíjať roboty, ktorý nevie povedať, či vnímajúci subjekt dokonalého ľudského vzhľadu je človek - alebo robot.

V konečnom dôsledku sa ľudia musia pripraviť, ako sa vysporiadať s takýmito scenármi alebo sa im vyhnúť.

Bezprostrednejším scenárom v umelej inteligencii sa však nezaoberali ani spisovatelia sci-fi, ani „znepokojení“ dnešní vedci a tvorcovia verejnej mienky:

Najpravdepodobnejším bezprostredným scenárom je, že umelá inteligencia je a bude aj naďalej medializovaná, ale hlboko chybná technológia náchylná na obrovské chyby a omyly. A táto umelá inteligencia teda na nejaký čas sklame svoje vznešené ciele a možno dokonca bude stáť nespočetné miliardy (alebo bilióny) peňazí alebo dokonca zničené životy kvôli svojim zlyhaniam.

AUTOR: Karsten Riise

Překlad: S.S.

ZDROJ: <https://www.globalresearch.ca/where-we-going-artificial-general-intelligence/5852837>