



UMELÁ INTELIGENCIA

PRIPRAVTE SA NA BUDÚCNOŠŤ

FILIP HANKER / JÁN TRANGEL / MAROŠ ŽOFČIN

Filip Hanker, Ján Trangel, Maroš Žofčín
Umelá inteligencia: Pripravte sa na budúcnosť

UMELÁ INTELIGENCIA

PRIPRAVTE SA NA BUDÚCNOSŤ

FILIP HANKER / JÁN TRANGEL / MAROŠ ŽOFČIN

›aktuality.sk


ZIVE

Filip Hanker, Ján Trangel, Maroš Žofčín
Umelá inteligencia: Pripravte sa na budúcnosť

Copyright © Filip Hanker, Ján Trangel, Maroš Žofčín, 2025

Cover design © Martin Žubor

Layout © Samuel Ryba - Design Ryba

O výbere respondentov rozhodli autori knihy.
Ide o nekomerčnú spoluprácu.

Slovak edition © Ringier Slovakia Media, s. r. o.

ISBN: 978-80-69052-25-3

OBSAH

DVE STRANY AI	9
Mozog, srdce a budúcnosť AI.....	11
Rozhovor s Máriou Bielikovou	
AI A BIG TECH	35
Čo majú spoločné AI a objav elektriny?.....	37
Rozhovor s Michalom Valkom	
AI A MEDICÍNA	51
Pomôže liečiť, lekárov však nenahradí.....	53
Rozhovor so Štefanom Korcom	
AI A VESMÍR	67
Odhalíme vďaka AI život vo vesmíre?.....	69
Rozhovor so Žofiou Chrobákovou	
AI A VIZUALIZAČNÉ NÁSTROJE	83
Svet očami umelej inteligencie.....	85
Rozhovor s Vandou Benešovou	
AI A DRONY	97
Dron ako poštár, kontrolór či ochranca.....	99
Rozhovor s Jozefom Rodinom	
AI A ETIKA	113
Môže AI dostať pokutu či skončiť vo väzení?.....	115
Rozhovor s Danielou Vacek	
AI A PRÁVO	127
Získa AI právnu zodpovednosť?.....	129
Rozhovor s Miroslavom Chlípalom	

AI A DEEFAKE	141
Prečo (ne)budeme môcť veriť vlastným očiam ani ušiam?	143
Rozhovor s Michalom Gregorom	
AI A ROBOTY	157
Robot ako kolega i pomocník v domácnosti	159
Rozhovor s Radoslavom Balajkom	
VPLYV AI NA DETI A MLÁDEŽ	171
Ako AI zmení naše deti?	173
Rozhovor so Zuzanou Juránekovou	
AI A KYBERNETICKÁ BEZPEČNOSŤ	187
Ochráni nás AI pred kyberkriminálmi?	189
Rozhovor s Lukášom Hlavičkom	
AI A HERNÝ PRIEMYSEL	199
Vymyslí AI nekonečné hry na mieru?	201
Rozhovor s Marekom Rosom	
AI A TRH PRÁCE	215
Prečo nás AI (ne)vyhodí z práce?	217
Rozhovor s Ivanou Molnárovou	
AI A PODNIKANIE	229
Ako AI mení turbulentný svet podnikania?	231
Rozhovor s Daliborom Cicmanom	
AI A MÉDIÁ	245
Úloha slobodných médií v ére AI	247
Rozhovor s Františkom Novákom	
SLOVNÍK POJMOV A SKRATIEK	263

PRIPRAVTE SA NA BUDÚCNOSŤ!

Žijeme v ére umelej inteligencie. Nie je to iba módna vlna ani výmysel technologických firiem, ktoré chcú vylákať ďalšie peniaze od investorov. Umelá inteligencia – alebo lepšie povedané súbor služieb umelej inteligencie – je ďalšou technologickou revolúciou. Takou významnou, akými boli svojho času vynájdenie kníhtlače, príchod parného stroja alebo nástup internetu. Všetky zmienené objavy prevrátili naruby vnímanie sveta a priniesli rýchly rozmach úplne nových priemyselných odvetví.

S umelou inteligenciou je to rovnako. Jej nástup však sprevádza jeden podstatný rozdiel. Žiadna zo spomínaných zmien neprichádzala tak rýchlo ako práve AI. Internet sa aj na Slovensku začal rozmáhať v polovici 90. rokov minulého storočia, pričom trvalo minimálne jednu dekádu, kým sa z neho stal kľúčový biznisový nástroj. Ostatne, spomeňte si, ako dlho dokázal odolávať elektronickej konkurencii napríklad papierový telefónny zoznam.

V prípade umelej inteligencie nemáme luxus v podobe času potrebného na to, aby sme sa zmene prispôbobi opatrne a pomaly. Treba konať hneď, nové služby a nástroje AI sa objavujú prakticky denne. Ignorovať túto technológiu a tváriť sa, že neexistuje, sa môže rýchlo ukázať ako kardinálna chyba. A to je len jeden z dôvodov, pre ktoré sme sa rozhodli napísať túto knihu.

Snažili sme sa pripraviť viac než len analýzu súčasného stavu vývoja služieb umelej inteligencie, ktorý je vzhľadom na dynamiku zmien iba ťažko uchopiteľný v čase. Naším cieľom bolo poskytnúť čitateľskej verejnosti čo najkomplexnejší obraz o fenoméne, ktorý zásadným spôsobom mení naše životy. V snahe obsiahnuť čo najširšie spektrum oblastí ovplyvňovaných AI sme oslovili špičkových

slovenských odborníkov a odborníčky, ako aj zástupcov podnikateľskej sféry, aby nám ponúkli vlastný pohľad na život v ére AI, aktuálne trendy i budúcnosť služieb umelej inteligencie.

Výsledkom je 272-stranová kniha šestnástich rozhovorov na rôzne témy: od medicíny cez psychológiu, vzdelávanie, právo, kybernetickú bezpečnosť, trh práce, herný priemysel a podnikanie až po výskum vesmíru, robotiku či médiá. A tým sa výpočet ani zďaleka nekončí. Umelá inteligencia nastolí nové etické otázky, ovplyvní medziľudské vzťahy a prinesie do nášho každodenného bytia novú entitu, s ktorou sa budeme musieť naučiť komunikovať – robota. Celá spoločnosť bude musieť nájsť efektívne, spoľahlivé a bezpečné spôsoby, ako existovať v novom svete a využívať najnovšie technologické nástroje vo svoj prospech. V práci i súkromnom živote.

Dramatické zmeny nastanú vo všetkých oblastiach života. Často opakované obavy z toho, že nás AI pripraví o prácu, budú platiť len vtedy, ak sa ju rozhodneme ignorovať. V skutočnosti nám ju totiž zoberú ľudia, ktorí sa naučia efektívne používať nástroje umelej inteligencie, vďaka čomu budú schopní urobiť rovnakú činnosť omnoho rýchlejšie a kvalitnejšie. Zlom nastane aj v rôznych oblastiach vedy. Nové obzory, ktoré otvoria nástroje umelej inteligencie, nám umožnia rozlúsknuť doteraz neriešiteľné matematické či medicínske problémy, analyzovať dáta z vesmírnych teleskopov a sond a, ktovie, možno dokonca aj odhaliť mimozemský život.

Umelá inteligencia je revolúcia, ktorá môže pre ľudstvo znamenať významný krok vpred. A nejde len o to, že môže zásadne prispieť k predĺženiu nášho života, zmierneniu klimatických zmien, riešeniu energetických problémov či zvýšeniu životnej úrovne. Záleží len na nás, aký bude mať vplyv – na spoločnosť i nás ako jednotlivcov. Aj preto nemá zmysel podľahnúť strachu zo zmien, ktoré AI prináša. Naopak, treba s nimi počítať.

Pripravte sa na budúcnosť. Tá patrí ľuďom, ktorí vedú efektívne využívať umelú inteligenciu vo svoj prospech.

DVE STRANY AI

Za posledné dve dekády nastali na poli umelej inteligencie až tri zásadné prelomy. Prvým bol rok 2012, keď sa udial veľký pokrok v rozpoznávaní obrazu. Rok 2017 zas priniesol transformery na lepšie spracovanie textu. Koncom roka 2022 bola verejnosti predstavená prvá verejná verzia ChatGPT.

Súčasný AI systém sú založené na neurónových sieťach s miliardami parametrov, ktoré sa učia na obrovskom množstve dát. Mohlo by sa zdať, že podobne funguje aj ľudské vedomie, no v prípade AI dnes nejde o skutočnú inteligenciu a systémy nemajú vlastné vedomie ani emócie. Netreba zabúdať ani na to, že umelá inteligencia nie je neomylná - naopak, môže a robí chyby.

Rýchly vývoj AI so sebou prináša nielen pozitíva, ale aj riziká. Nehovoríme o apokalyptických, tie sú vzdialené a často nereálne. Sú tu však hmatateľné hrozby. Už dnes začína byť vážnym problémom šírenie dezinformácií pomocou deepfakov, využitie AI vo vojnách či prehľbovanie sociálnych nerovností. Aj preto je dôležité premýšľať okrem iného nad tým, prečo by nemali byť modely AI výlučne v rukách súkromných spoločností či prečo budú v ére AI ešte dôležitejšie kritické myslenie a vedomosti.



MOZOG, SRDCE A BUDÚCNOSŤ AI

Aký je rozdiel medzi generatívnou umelou inteligenciou a inými druhmi AI, známymi aj v minulosti?

V prvom rade treba povedať, že umelá inteligencia zahŕňa viacero oblastí od strojového učenia, spracovania obrazu a textu cez analýzu dát, optimalizácie, predikcie až po znalostné a expertné systémy, fuzzy logiku a ďalšie.

Výstupom klasického strojového učenia je nejaká hodnota. Napríklad, keď máme úlohu klasifikácie, umelá inteligencia mi povie: Toto je mačka, toto je pes, toto je niečo iné. Alebo pri úlohe predikcie mi systém AI povie, aká bude v budúcnosti hodnota skúmaného parametra a na základe toho môže robiť systém alebo človek rozhodnutia.

Keď hovoríme o generatívnej umelej inteligencii, systém generuje text, obrázky, video, zvuk. Inak povedané, vytvára nové artefakty a nielen hodnoty. A to je celkom zásadný posun.

Čo si pod pojmom AI predstavíte vy osobne a aký široký je podľa vás tento pojem?

Umelá inteligencia je disciplína, ktorej cieľom je vytvárať systémy, ktoré sú počítačové, umelé a vykazujú inteligentné správanie. A to je veľmi široké. Základnou vlastnosťou je schopnosť učiť sa. V situáciách, pre ktoré neboli ani naprogramované, ponúkajú dobré riešenia. Často tieto situácie nepredpokladali ani ich tvorcovia.

Väčšina laickej verejnosti registruje umelú inteligenciu posledný rok, dva, hoci ona je tu s nami už desaťročia. Môžete stručne vysvetliť, akým vývojom za ten čas prešla?

Ľudia sa odjakživa po celé stáročia snažili vytvoriť inteligentné umelé bytosti alebo stroje. Zoberte si príbeh Golema alebo Frankenstein.

No až s príchodom počítačov sa vytvorenie strojov, ktoré sa v nejakých situáciách správajú tak, že to človek označí za inteligentné, začalo stávať realitou.

V 50. rokoch 20. storočia vznikol známy Turingov test, ktorý mal za cieľ zistiť, či stroj môže vykazovať inteligentné správanie neodlíšiteľné od ľudského. Test sa zakladá na princípe imitačnej hry, v ktorej sa zisťuje, či stroj dokáže komunikovať tak, aby ho ľudský pozorovateľ nedokázal odlíšiť od skutočného človeka. Turingov test dlho slúžil ako dobrý nástroj na vývoj nových riešení, dnes s veľkými jazykovými modelmi ho už vieme prekonať, ale o skutočnej inteligencii to nesvedčí.

Treba tiež povedať, že dnes rozšírené neurónové siete vznikli ešte v 40. rokoch minulého storočia. Umelá inteligencia sa odvtedy, za tých viac ako sedemdesiat rokov, vyvíjala viac alebo menej intenzívne cez obdobia zimy a leta, ktoré sa niekoľkokrát vystriedali. Mali sme dve obdobia, keď nastal útlm rozvoja umelej inteligencie. Vtedy sa do AI investovalo menej. Až prišlo obdobie okolo roku 2010, keď už sme mali dostatok výpočtovej kapacity aj znalostí, algoritmy i potrebné množstvo dát na to, aby sa začali rozvíjať hlboké neurónové siete.

Kým v minulosti sme prešli z pomyselnej zimy do leta a nakrátko späť, aby sme sa opäť ocitli v lete, dnes zrejme zažívame súvislé tropické leto. Kedy k tomu došlo?

Definitívny prelom nastal niekedy v roku 2012, keď bola súťaž vo vizuálnom rozpoznávaní ImageNet (veľká databáza navrhnutá na vizuálne rozpoznávanie objektov, pozn. aut.). Spracovanie obrazu bol odjakživa veľmi vďačný problém pre umelú inteligenciu, pretože na jej využitie strojmi treba, aby stroj rozpoznal situáciu, v ktorej sa nachádza, identifikoval, čo sa nachádza v prostredí alebo na obrázku, aby vedel reagovať.

Keď chcete zostrojiť robota, musíte mu dať oči, aby bol schopný rozpoznáť objekty a napríklad vedel niečo uchopiť. Všetky metódy, ktoré boli dovtedy vyvinuté, sa podarilo prekonať práve hlbokými neurónovými sieťami. S výrazne vyššími presnosťami. A to bol začiatok intenzívneho leta umelej inteligencie, ktoré tu máme dodnes.

V roku 2017 nastúpili transformery, čo je architektúra pre veľké jazykové modely s mechanizmom pozornosti, ktorý umožnil, aby hlboké neurónové siete dokázali efektívne pracovať aj s textom. Text má totiž na rozdiel od obrázkov ohromnú variabilitu.

A napokon prišiel koniec roka 2022, keď mala verejnosť po prvý raz v histórii priamy kontakt s veľkými systémami umelej inteligencie v podobe ChatGPT.

Možno je trochu škoda, že až v ten rok si väčšina ľudí uvedomila AI a jej možné dosahy, pričom za umelú inteligenciu sa často považuje len tá generatívna.

Od príchodu ChatGPT uplynul krátky čas, no umelá inteligencia tu bola aj predtým a v niektorých oblastiach sa uplatňovala v pomerne rozsiahlom meradle. Ľudia si málo uvedomujú, že napríklad vo fotoaparátoch či mobiloch majú umelú inteligenciu už dlhšie.

Keď sa pozrieme, kedy vznikol Facebook a sociálne médiá, to je už skoro dvadsať rokov. Ich súčasťou sú odporúčacie systémy používajúce umelú inteligenciu. Alebo si stačí spomenúť na navigáciu v aute či medicínu, kde sú pokroky v súvislosti s umelou inteligenciou obrovské. Toto všetko však bolo to skryté v aplikáciách, čo je pre technológie prirodzené.

Prečo si to nevšímame, keď je to pre náš život, zdravie, komfort či bezpečnosť vlastne veľmi dôležité?

Keď je technológia rozumne odskúšaná a prejde testami, slúži ľuďom a veľmi neprekvapuje, ľudia ju nemajú potrebu veľmi skúmať. Drvivá väčšina populácie nerieši, na akom princípe funguje lietadlo. Vie, že funguje. Zároveň si je vedomá, že niekedy môže zlyhať. Ale to riziko sme ochotní akceptovať a bežne ho neriešime.

Avšak zrazu sa nám dostalo priamo do domácností niečo, čo nás fascinuje, keďže to má schopnosti, ktoré sme doteraz pripisovali výlučne ľuďom. Mimochodom, to, že podobný princíp bol využitý už predtým v roku 2018 v projekte AlphaFold na predikciu štruktúry bielkovín, si bežní ľudia nevšimli. Až do roku 2024, keď boli autori - informatici z DeepMind - ocenení Nobelovou cenou za chémiu.

Keď sa dnes hovorí o službách typu ChatGPT, často sa skloňuje slovo model, ktorý je využívaný danou službou na to, aby nám poskytoval konkrétne výsledky. Mohli by ste vysvetliť, o čo ide? Na ChatGPT alebo podobné služby sa môžeme pozeráť ako na nejaký AI systém, ktorý obsahuje viaceré komponenty. Model je jedným z nich. Bez neho by to nefungovalo, ide o srdce alebo mozog celého systému.

Model je neurónová sieť. Môžeme si to predstaviť ako navzájom prepojené uzly – umelé neuróny. Neurón prijme signál od prepojených neurónov, spracuje ho a pošle spracovaný signál do prepojených neurónov na ďalšej vrstve.

Spracovanie predstavuje nejakú matematickú funkciu. Ide spravidla o jednoduchú nelineárnu funkciu. Príklad často používanej funkcie, ktorej rozumie aj prvák v škole, vyzerá takto: ak je nejaké číslo väčšie ako nula, výsledkom je táto hodnota, v opačnom prípade je to nula.

Takýchto neurónov je v modeli veľmi veľa a sú usporiadané do vrstiev. Neuróny sú poprepájané podobne ako v našom mozgu, aj keď prirovnanie k nášmu mozgu je priveľké zjednodušenie.

Do toho celého ešte vstupujú váhy jednotlivých prepojení.

Prepojenia medzi neurónmi v AI modeli majú svoje váhy – čísla. Tieto váhy sa v procese učenia optimalizujú. Vstupná vrstva prijme nejaké dáta, napríklad obrázok alebo text zakódovaný do čísel. Signál v podobe čísel sa šíri a spracúva od začiatku až po koniec siete, kde sa vygeneruje výsledok, napríklad pravdepodobnosť, že obrázok na vstupe je mačka.

Podľa správnosti výsledku sa následne upravujú váhy. Informácia, že sieť GPT-3 má 175 miliárd parametrov, v podstate zodpovedá týmto číslam – to sú spomínané váhy.

Ako si to má laik predstaviť?

Predstavte si sieť, ktorá má miliardy prepojení, ktoré predstavujú nejaké čísla, čomu zodpovedá primeraný počet neurónov. Takáto sieť s dostatočným počtom neurónov a správnymi váhami dokáže približne vypočítať akúkoľvek spojitú funkciu s požadovanou presnosťou.

Dôležité je, že sieť sa trénuje. Hlavným mechanizmom trénovania neurónových sietí, ktorý považujem za jeden z najzásadnejších objavov v tejto oblasti, je spätná väzba. Zabezpečuje, aby sa neurónová sieť mohla učiť z chýb a upravovať svoje váhy tak, aby dosahovala lepšie výsledky. To znamená, že vstupné dáta sa prenesú sieťou až po výstup a spočíta sa chyba, aby sme vedeli, ako ďaleko je aktuálny výsledok od správneho výsledku.

Následne sa vypočíta, ako treba upraviť váhy správnym smerom, teda tak, aby sa chyba zmenšila. A keď sa to opakuje veľa krát, tak sa sieť natrénuje.

Ako to funguje pri jazykových modeloch?

Jazykový model bol trénovaný na úlohe nahrádzania slov vo vetách. Čiže vie to robiť bez toho, aby sme mali príklady, na ktorých sa model učí, s označenými správnymi riešeniami človekom. Máme nejakú vetu, vyhodíme jedno slovo, dáme vetu do siete a spýtame sa, aké slovo chýba. Sieť navrhne slovo na doplnenie, to sa porovná s tým, ktoré sme vyhodili, upravíme váhy a opakujeme veľa krát.

Zaujímavé pritom je, že celé to vôbec nevzniklo preto, že by niekto chcel vymyslieť generátor dobrých textov alebo sumarizačný nástroj. Aký bol skutočný dôvod?

Vzniklo to z veľmi jednoduchej úlohy. Výskumníci sa snažili urobiť dobrý systém na prepisovanie reči do textu. Vstupná nahrávka je však často zašumená, čiže niektoré slová nie sú dobre rozpoznateľné, respektíve nie sú dobre počuť. Takže keď vymyslíte strojček, ktorý dokáže v kritických miestach slovo správne doplniť, tak urobíte lepší prekladač.

Lenže čo sa stalo? Výskumníci - zvedavci mali dosť peňazí, tak neurónové siete začali zväčšovať. Pridali vrstvy, parametre aj obrovské množstvo textu na trénovanie. A zisťovali, čo sa bude diať. Zväčšovanie bolo obrovské, na trénovanie sa míňali milióny dolárov. Keď to začali v OpenAI robiť, mnohí neverili a smiali sa, aké hlúposti vystrájakajú.

Nakoniec sa ukázalo, že významné zväčšovanie - škálovanie bolo pre rozvoj jazykových modelov zásadné. Prinieslo nové vlastnosti,

o ktorých nikto predtým netušil. Napríklad okrem dopĺňania slov do vety ten istý model po malých úpravách vie robiť aj sumarizácie, analyzovať sentiment (náladu, pozn. aut.) alebo odpovedať na zložité otázky. A v tom je to vzrušujúce.

Myslíte si, že má ešte zmysel ďalšie škálovanie? Dá sa vôbec ešte niekam v sieťach rásť? Alebo sme už dosiahli vrchol, že to vie všetko, čo by to vedieť malo, a teraz, naopak, príde trend zmenšovania?

Trend zmenšovania už nastal. V tejto oblasti je pomerne rozsiahly výskum, ale aj mnoho praktických experimentov. Aj v Kempeleňovom inštitúte sa zaoberáme tvorbou modelov s obmedzenými zdrojmi. Tiež veľa firiem sa snaží robiť menšie modely. Často sú ešte lepšie, hoci len pre vybrané úlohy alebo oblasti. Zároveň sme stále ešte aj vo fáze, keď existuje mantra zväčšovania a hľadáme, čo sa ešte možno objaví.

Má to, samozrejme, svoje limity. Jedným vážnym je, že je to veľmi drahé, pretože potrebujete naozaj veľmi veľa výpočtovej kapacity a energie. Napriek tomu, že sa vymýšľajú rôzne spôsoby, ako efektívnosť zlepšiť – napríklad doladením modelov aj s využitím spätnej väzby od človeka alebo zlepšením výkonu efektívnejšou reprezentáciou vstupného textu.

Ako si vlastne môžeme predstaviť spotrebu elektrickej energie pri podobných službách? Aké náročné sú?

Veľmi náročné. Nemáme všetky informácie, ale sú rôzne odhady a porovnania, najmä na základe štatistík o použitých grafických kartách. Najnáročnejšie na spotrebu elektrickej energie je tréning modelov. Pre GPT-4 sa odhaduje, že to je asi desaťkrát viac ako ročná prevádzka ChatGPT. A predstavuje to približne spotrebu elektrickej energie stošesťdesiatich priemerných amerických domácností.

Treba však započítať aj experimentovanie, ladenie modelov vrátane neúspešných pokusov a výskumu, ktoré robia samotné spoločnosti vyvíjajúce modely. A tiež to, že tieto štatistiky sú z roku 2024. Vznikajú nové verzie modelov a aj používanie týchto systémov sa pomerne rýchlo mení.

Tieto snahy o zefektívnenie sa uplatňujú pri tréňovaní alebo až pri samotnej prevádzke?

Ide o rôzne mechanizmy do neurónových sietí, ktoré umožňujú už priamo pri tréňovaní, v prevádzke, pri dotréňovaní či doladovaní modelov využiť veci tak, aby sa znížila potreba výpočtovej kapacity. Na tréňovanie potrebujete aj veľa dát, pričom dnešné modely už „pochrúмали“ skoro všetok dostupný obsah. Prebiehajú diskusie aj o tom, čo sa stane, keď systémy začneme kŕmiť (a už možno aj kŕmime) textami, ktoré už boli čiastočne alebo úplne vygenerované umelou inteligenciou.

Nevieme, čo sa stane. A nevieme ani to, čím sa budú služby tréňovať. To znie ako problém.

To je jeden z problémov. Súkromné spoločnosti, ktoré pôsobia mimo Európy, mnohé veci ani neprezradia. Takže sa o tom len diskutuje a vieme len niečo. Mnohé však nevieme.

Nová regulácia AI v Európe nariaďuje, aby firmy mali zdokumentované zdroje, na ktorých boli systémy tréňované. No neverím, že nejaká z veľkých firiem, ktoré majú veľké modely, niečo také sprístupní v plnom rozsahu. Na druhej strane sú aj otvorené modely.

Dokedy teda môžu modely rásť?

Myslím si, že ešte minimálne najbližší rok budú rásť, ale viac budú rásť inteligentným spôsobom než hrubou silou. Ktovie, či narazíme na hranicu, ktorá bude znamenať, že už nemáme dostatok peňazí, lebo si niekto povie, že už do toho nebude investovať. Alebo že sa už vlastnosti modelov nebudú zlepšovať. To sa dá predpovedať len veľmi ťažko.

Modely, ako aj schopnosti sietí sa neustále menia. Je možné, že sa siete niekedy môžu aj zhoršiť, nielen zlepšiť? Napríklad pri medziročnom porovnaní.

To závisí od toho, akú úlohu budú riešiť. A závisí to aj od toho, ako sa opýtate. V Kempelenovom inštitúte sme napríklad opakovane robili výskumy, ako vedľa veľké jazykové modely generovať dezinformácie. A zistili sme, že odpovede sa menili.

V niečom sa to zlepšilo – niektoré modely už odmietnu odpovedať alebo uvedú nejaké varovanie. Naopak, v niečom sa to zase z pohľadu rizikovosti dezinformácií zhoršilo, lebo texty sú jednoducho lepšie, presvedčivejšie a nebolo veľmi ťažké model presvedčiť, aby vygeneroval požadovanú informáciu. Dnes už vedia modely generovať aj personalizované informácie a, samozrejme, aj dezinformácie, čo je veľmi nebezpečné.

Nakoľko majú tvorcovia AI služieb pod kontrolou výstupy samotných sietí? Teda že môžu niečo zmeniť, zakázať, ale nevedia vždy presný efekt.

Dnes rozumieme len časti z toho, čo sa deje v neuronových sieťach. Čiže nerozumieme úplne všetkému, čo sa stane, keď niečo zmeníme, doladíme, ani tomu, aký to bude mať vplyv na ďalšie veci. To je prvá zásadná vec.

Druhá zásadná vec je, čo sa s modelmi robí ďalej alebo skôr čo sa robí s AI systémami. Lebo jedna fáza, keď máme takýto model, je jeho doučenie cez inštrukcie. Aby model robil to, čo chceme. Keby sme to nespravili, tak pokiaľ by ste napríklad zadali otázku, namiesto odpovede by model pokračoval v nejakom texte.

A ďalej cez *reinforcement learning*, čo je učenie s odmenou a trestom, sa systém ešte doladuje tak, aby na niektoré veci neodpovedal vôbec, aby vyberal čo najlepšiu odpoveď pre človeka a podobne. Zatiaľ nemáme lepší spôsob ako ten, že to robia ľudia – mnohokrát v krajinách, kde sa hovorí anglicky a je veľmi lacná pracovná sila. Stále je to ohromne finančne aj časovo náročné, keďže ide prevažne o manuálnu ľudskú prácu.

Prečo niekedy nevieme, ako modely pracujú, aký výstup nám dajú a ako ho pripravia? Čo je prekážkou vedeckého poznania alebo detegovania spôsobu ich fungovania?

Neuronové siete sa označujú ako subsymbolická umelá inteligencia, lebo pracujeme s množstvom čísel bez priamej interpretácie. Naproti tomu metódy symbolickej umelej inteligencie explicitne „napriamo“ pracujú so symbolmi tak, ako ich pozná človek.

V symbolických prístupoch umelej inteligencie používame napríklad pravidlá, ktoré vyjadrujú príčinné vzťahy - ak sú splnené nejaké podmienky, tak nastáva nejaká situácia. Pravidiel môže byť aj veľmi veľa, nemusia byť konzistentné. Ale reprezentujú znalosti spôsobom, ktorému človek rozumie. Pri usudzovaní s pravidlami pracujeme so symbolmi a vieme vysvetliť vytvorené riešenia a analyzovať zlyhania.

Pri neurónových sieťach máme množstvo váh v sieti, ktoré reprezentujú znalosti, a porozumenie tejto štruktúry je veľmi náročné. Je to daň za lepšie výsledky.

Tí, ktorí tomu rozumujú, ma teraz asi budú trochu haniť, že príliš zjednodušujem, ale dá sa na to pozrieť aj takto: Pravidlá, ktoré sme v symbolickej umelej inteligencii vytiahli z ľudí - lebo sú to naše znalosti zapísané jazykom, ktorému človek rozumie -, boli v subsymbolickom prístupe odvodené strojom z dát, ktoré človek necháva ako stopy svojej činnosti. Aj keď to v skutočnosti nie sú pravidlá, lebo je to zložitejšie, keďže máme neurónovú sieť s mnohými vrstvami.

A teraz si predstavte, že máte tých 175 miliárd parametrov a v nich je zakódovaná znalosť. Často malé zmeny v sieti majú zásadný vplyv na výsledok a nevieme vždy prečo, nevieme zopakovať situácie, v ktorých nastalo niečo neobvyklé.

Preto je dôležitý základný výskum. Dnes pri nejakých typoch úloh vieme povedať, čo sa deje na tej-ktorej vrstve. Napríklad, že sa pracuje s nejakými časťami obrázka. Avšak to je stále ešte veľmi ďaleko od toho, aby sme rozumeli, na základe čoho sa systém rozhodol.

Akú úlohu bude AI zohrávať pri rozvoji spoločnosti alebo jej udržateľnosti? Akým spôsobom bude ovplyvňovať ekonomiku, vzdelávanie či zdravotníctvo?

Vplyv umelej inteligencie je dnes čoraz väčší. A časom bude obrovský. Nehovorím to preto, že pôsobím v tejto oblasti, a preto si myslím, že sa bude rozvíjať. Stačí si uvedomiť rozdiel oproti predchádzajúcim technológiám.

Aj doteraz sme vedeli vytvárať zložité systémy - hardvérové, teda trojrozmerné, aj softvérové. Už som spomenula napríklad lietadlo. Je úžasné postaviť niečo, vďaka čomu človek dokáže lietať. No vždy,

keď sme chceli nejakú činnosť automatizovať, museli sme vymyslieť, ako to, čo robí človek, transformovať do postupnosti krokov, aby to dokázal realizovať stroj.

Avšak umelá inteligencia automatizuje tento proces na základe štôp, ktoré necháva pri svojej činnosti človek, a robí to automaticky pre mnohé situácie.

Povedzme si príklad, aby si to ľudia mohli predstaviť.

Napríklad máme dáta o tom, ako sa človek v banke rozhoduje pri poskytovaní úveru. Máme aj informácie z rôznych senzorov, máme dáta, ako ľudia vyjadrujú, či sa im páči nejaký príspevok na sociálnej sieti.

Algoritmy AI hľadajú opakujúce sa vzory v dátach a dokážu pri novom prípade, ktorý predtým nevideli, dobre rozhodnúť o pôžičke, rozpoznať anomáliu v nameraných dátach či odporučiť informácie na sociálnej sieti, ktoré sa budú danému človeku páčiť, aj keď presne rovnakú situáciu predtým nevideli.

Zrazu sa teda AI môže dostať do mnohých oblastí. A to je niečo, čo sa nedá prehliadať. Lebo efektívnejší budú tí, ktorí sa o problematiku zaujímajú, či tí, ktorí budú tieto technológie používať. Aj tí, ktorí ich budú dávať do svojich výrobkov alebo budú používať produkty iných.

Často až neporovnateľne efektívnejší. Vieme na nejakom príklade ilustrovať rozdiel?

Je to podobné, ako keď idem napríklad do New Yorku po vlastných alebo lietadlom. Dá sa jedno aj druhé, veď oceán by sa asi dal s nejakými pomôckami preplávať, ale letecky sa tam dostanem rýchlejšie.

AI bude mať čoraz väčší vplyv nielen na našu prácu, ale aj zábavu. Naše generácie už síce vyrastali s technológiami, ale nie s takými to exponenciálnymi. Vplyv najmä posledných rokov, keď vznikla generatívna umelá inteligencia, na našu komunikáciu, spracovanie informácií či spoločenské pravidlá, ktoré sa už dnes menia, bude nesmierny.

Aj keď nie vždy len pozitívny. V čom vidíte možné slabiny, riziká?

Obávam sa, že začína byť ohrozená samotná demokracia. Nespôsobila to však umelá inteligencia, ale my ľudia. Umožnila to najmä