

Nenad Vertovšek

Zadar, Hrvatska

nenad.vertovsek@gmail.com**Izazovi umjetne inteligencije – ne-ljudska budućnost?!**

Sažetak: Upravo ljudi – od znanstvenih rasprava do medija – nameću svojevrsan strah od robotizacije, AI koja će nadmašiti ljudsku, svjesnosti suprotstavljene čovječanstvu. Sadašnji i budući izazovi umjetne inteligencije zapravo su izazovi ljudske (ne)inteligencije i prepreka koje ljudsko društvo i kultura moraju savladati kako bi otišli korak naprijed u društvenoj i psihološkoj evoluciji. Ponovo smo na raskrižjima koja čovječanstvo vode u slijepu ulicu vlastitih sebičnih, korporativnih i profitnih interesa nasuprot prevladavanju ograničenja koje čovjek nameće sebi i drugim ljudima. Hoće li umjetna inteligencija biti samo oružje globalnih lidera i bezobzirnih interesa ili jedno od oruđa za povezivanje ljudi u jednu novu razinu svijesti i ponašanja? Može li umjetna inteligencija dobiti i vlastitu svjesnost ne-ljudske etičnosti koja neće biti protiv čovjeka, već protiv čovjeka koji to više nije? Suživot čovjeka i AI ili nova vrsta koja će biti neko novo ne-ljudsko?

Ključne riječi: Umjetna inteligencija, ljudska svijest, društvena i tehnološka revolucija, etika umjetne inteligencije

Challenges of artificial intelligence - a non-human future?!

Abstract: It is people - from scientific debates to the media - who impose a kind of fear of robotization, AI that will surpass human, awareness opposed to humanity. The present and future challenges of artificial intelligence are in fact the challenges of human (non)intelligence and the obstacles that human society and culture must overcome in order to take a step forward in social and psychological evolution. We are once again at the crossroads that lead humanity into a dead

end of its own selfish, corporate and profit interests versus overcoming the limitations that man imposes on himself and on other people. Will artificial intelligence be just a weapon of global leaders and reckless interests or one of the tools to connect people to a new level of consciousness and behavior? Can artificial intelligence gain its own awareness of non-human ethics that will not be against man, but against man who is no longer human? Coexistence of humans and AI or a new species that will be some new non-human?

Keywords: Artificial intelligence, human consciousness, social and technological revolution, ethics of artificial intelligence

Budućnost ima mnogo imena – za slabe ona je nedostižna, za bojažljive nepoznata, za hrabre ona je šansa...

Victor Hugo

Prekrasna stvar u vezi tehnologije jest to da ljudi s njom naposljetku rade nešto različito od onoga za što je izvorno bila namijenjena...

Manuel Castells

Negdje između holivudskih dramatičnih zapleta 60-tih godina i razina umjetne inteligencije koje teško razumiju mase podložne medijskim iskrivljavanjima, danas se nalazi i svijest većine običnih ljudi kada razmišljaju o budućnosti umjetne inteligencije. Kažemo obični ljudi, jer većina, bez obzira koliko mislila da je naprednija od, primjerice, srednjovjekovnih pogleda na svijet, zapravo nije daleko od toga, kada se vrati na osnovne postavke „što i gdje sam ja? Ili „što će biti sa mnom i sa svijetom?“. Strahovi, frustracije, stereotipi, predrasude i manipulacije čine naš svijet naizgled svijetom blistavih reklama, no ujedno i sveprožimajuće paklene stvarnosti u kojoj većinom bježimo od posljedica naše nemoći, neznanja i prepuštanja vlastite budućnosti političarima i pohlepnim korporacijama.

Naravno, nije sve tako crno i sivo kako možda izgleda, no moramo poći od onog što je danas umjetna inteligencija uglavnom visokostručnih, profesionalnih ljudi, znanstvenika i stručnjaka koji se već desetljećima bave onim što nazivamo umjetnom inteligencijom. Dakako, nije to nikakav poziv da mislimo isto kao tzv. elita, jer ratove s umjetnom inteligencijom, onima koji se vode i koji će se tek voditi, ne kreiraju obični ljudi već visokointeligentni stručnjaci s diplomama najboljih sveučilišta, znanstveni *think tankovi*, trustovi mozgova znanstvenika i stručnjaka. Stoga bi već na početku trebalo naznačiti kako trenutno opasnost ne prijeti ljudima i čovječanstvu od umjetne inteligencije, već ljudske neinteligencije ili visoko razvijene ljudske inteligencije bez „suvišnih“ etičkih ograda.

Premda je budućnost ljudskog utjecaja na umjetnu inteligenciju relativno neizvjesna, ovisna o nekim neljudskim aspektima dijela ljudskih karaktera istraživača, znanstvenika i vojnih struktura, razvoj ne-ljudske inteligencije¹ upravo dobiva svoj zamah. Tu valja ubrojiti i samo naizgled apsurdan trend gdje dio tzv. ljudske elite čini strahove, nerazumijevanje i pritisak kako bi AI učinio odbojnim širim slojevima, dok gorljivo radi na zlouporabama umjetne inteligencije u vojne svrhe i interese vlastite kontrole društvenih mehanizama. Takva „šizofrenost“, ali prije svega društveno i znanstveno licemjerje, dostiže svoje kritičnu točku koja se razvija brže i nego što bi pomislili..

U takvom apsurdnom, praktički oksimoronskom okruženju zapravo je sve teže i složenije razabrati – naravno i razumjeti – što se uopće događa na polju istraživanja i razvoja umjetne inteligencije. Kako bi to mogli razumjeti u suvremenom kontekstu, trebamo dijelom posegnuti i za nekim religijskim, mitološkim obrascima, odnosno drevnom ljudskom mudrošću. Michio Kaku to metaforički, gotovo poetski i s mitološkim naznakama, naglašava kako su u drevnim mitologijama božanstva mogla svom božanskom moći oživljavati neživo.

„Mi smo danas poput Vulkana dok kujemo u našim laboratorijama strojeve i udahnujemo život u čelik i silicij. Hoćemo li, međutim, osloboditi ljudski rod ili pasti u sužanjstvo?“², pita se Kaku i dodaje pomalo ironično kako se prema naslovnica masmedija čini sigurnim da će kontrolu nad ljudskim rodом preuzeti njegove vlastite tvorevine, temeljeći to brojnim medijskim sadržajima i senzacionalističkim pristupima. Sudeći po medijskim napisima koji dobrim dijelom igraju na naše prizemne strahove i bojazni iz neznanja, čovječanstvu kao da broje zadnji trenuci, a roboti spremaju zatvore u kojima ćemo ovisiti o njihovoj milosti.

Ipak, stvarnost nije tako stupidno bajkovita, premda cijele priče o umjetnoj inteligenciji svoju popularnost zaslužuju nizu drugih ljudskih otkrića koja su premašila svoje prvotne ciljeve i „otuđili“ se od namjeravanih humanističkih funkcija i namjena. Atomska bomba, dronovi, cjepiva, te druga masovna oružja završila su u krivim rukama, kao i umovi nebrojenih ljudi uhvaćenih kroz razna razdoblja povijesti u zamke ratnih psihoza, pobjedničkih fanfara i uvjerenja o vlastitoj izabranosti i „božjoj naklonosti“. Znanost i stručnjaci nešto opreznije postavljaju odgovore na

¹ Koristimo relativno uvriježen naziv *umjetna inteligencija*, premda valja razlikovati pojam „robot“, odnosno „strojeva“, za koji je velikim dijelom zaslužna navedena holivudska iluzija metalnih strojeva koji samo čekaju pokoriti jedno čovječanstvo. S druge strane imamo nešto točnije definicije androida, kiborga i ostalih humanoidnih robota ne samo od metala, već i plastike, imitacije ljudske kože i organa, pokretanih čipovima ili digitalnim i minijaturiziranim sustavima. Na kraju toga, zapravo na početku budućnosti eksperimentiranja i istraživanja, slijede povezivanja neurona, neuroplastičnih stanica s digitaliziranim sustavima i kvantnim računalima. Budućnost koja, po nekima, vodi i do nove (novih) vrste, povezivanju organskog i anorganskog života s genetskim i neuronskim transformacijama.

² Kaku, Michio, *Fizika budućnosti*, (2011) Zagreb: MATE d.o.o., str. 65.

polu-medijska nagađanja o tome kada bi pametni „strojevi“ mogli prevazići i nadvladati čovjeka. Znanstvenici i prognosticirano znanstveno vrijeme nekog otkrića ili realizacije mnogo je složenije. Katkad euforično, ali na duže staze većinom opreznije u granicama i ograničenjima umjetne inteligencije, posebno u sferi postizanja svjesnosti.

Još pedesetih godina prošlog stoljeća (jedni od „očeva“ umjetne inteligencije) Marvin Lee Minsky i John McCarthy opisuju umjetnu inteligenciju kao svaku zadaću koju izvodi program ili stroj za koju bi, da to obavlja čovjek, mogli reći kako je čovjek morao/trebao primijeniti inteligenciju za izvršenje takvog zadatka. Naravno, postoje brojne definicije koje preciznije ili specijalizirano naglašavaju određene aspekte umjetne inteligencije. Nastojimo li, u filozofskom smislu, pokriti područje – ili područja – umjetne inteligencije, možemo dijelom u prvi plan dovesti promišljanja o umjetnoj inteligenciji³ ispred tehničkih elemenata i tehnologije izrade, povezivanja i programiranja.

Za promišljanja o tomu možemo se koristiti i podjelom koju je postavio, uz sve moguće dodatke ili kritike, računalni znanstvenik Nick Heath odvajajući umjetnu inteligenciju u užem i širem smislu. Za onu u užem smislu smatra sve što vidimo oko nas u računalima - inteligentne sustave koji su podučavani ili naučili kako provesti određene zadatke bez izričitog programiranja kako to učiniti. Ova vrsta strojne inteligencije vidljiva je u prepoznavanju govora i jezika virtualnog pomoćnika *Siri* na Apple iPhone uređaju, u sustavima prepoznavanja vida na vozilima za samo-vožnju, u sustavu preporuka koje predlažu proizvode koji bi vam se mogli svidjeti na temelju onoga što ste kupili u prošlosti. Za razliku od ljudi, ti sustavi mogu samo naučiti ili biti naučeni kako napraviti određene zadatke, zbog čega ih se zove AI u „užem smislu“.

Pod umjetnom inteligencijom u općem smislu imamo vrlo različita mjerila, odnosno specifičnosti. To bi bio, recimo tako, tip prilagodljivog intelekta koji se izvorno nalazi u ljudima, ali i u ljudskim tvorevinama koje zovemo UI/AI, bio bi fleksibilan oblik inteligencije sposoban za učenje kako obavljati nevjerojatno različite zadatke, bilo što od šišanja i brijanja do izrađivanja proračunskih tablica ili motivacije za široku lepezu tema na temelju akumuliranog iskustva. U sferama umjetne inteligencije u pravom smislu takvo što još ne postoji u cijelosti, uz stručne i znanstvene kritike hoće li takvo što, kada i kako postati stvarnost.

Sam Michio Kaku navodi dio stručnjaka koji o tomu razmišljaju – još 1993. godine. Vernon Vinge je smatrao da će unutar 30-tak godina postojati nadljudska inteligencija proizišla iz onog što

³ Naravno, ali do trenutka, odnosno razdoblja u kojem će promišljanja određivati i realizirati sama umjetna inteligencija, ne (samo) ljudi kao danas, ili vrhunski programi umjetne inteligencije kojima su osnovne postavke još uvijek odredili ljudski stručnjaci...

se zove umjetna inteligencija. Međutim, s druge strane, Douglas Hofstadter izjavio je kako bi bio iznenađen ako to bude moguće za kojih 200 do 300 godina. Međutim, razmišljanja znanstvenika – o čemu svjedočimo i danas kod otkrića i puštanja u rad određenih programa od Alpha Go do Chat GPT-a 4 i razvijenijih inačica na kojima se upravo radi – kreću se oko toga da će se svjesnost računala (kakva god ona bila) vrlo izvjesno dogoditi, no teško je zasad reći kada.

Uostalom, i znanstvenici i znanosti se bore s više stotina i gotovo tisuća članaka, rasprava i zaključaka, i konstatacija o tomu što znači svjesnost računala, svjesnost tzv. neorganskog života. Ljudska imaginacija je nešto produktivnija – primjerice, pojam *Matrixa* gdje ljudi ni ne znaju da im je umjetna inteligencija postala gospodar. Premda, slična realizacija jest i tzv. *deep state*, duboke države u kojima vladaju samozatajni pokretači fenomena i političara koje vidimo i koje biramo. Ili filmski pojam *Skyneta* koji postaje svjestan i odmah napada čovječanstvo, jer smatra da mu je samoobrana nužna te pokreće nuklearni rat protiv čovječanstva. Kako ksenofobija i strah od nepoznatog ide kroz povijest, mijenjaju se i uloge tzv. neljudskih inteligencija, od zlih vanzemaljaca do umjetne inteligencije koja počinje razumijevati dosadašnju glavnu ulogu čovjeka na planetu i nastoji ga eliminirati.

No i u takvim se simulacijama obično polazi od (pret)postavke kako umjetna inteligencija slijedi obrazac već viđenog – beskrupuloznih i pohlepnih ljudi, zavjerenika i korporacija da pokore svijet. Što učiniti? Dijelom su takve imaginacije i promišljanja ponavljanje obrasca viđenog kroz ljudsku povijest i dijelom ekstrapolacije prošlosti prema budućnosti. Zanimaju se dijelom slika i ciljevi razvoja umjetne inteligencije koji će biti možda vrlo drugačiji od onog što poznajemo, možda i etičniji od ljudske sitničavosti, kontroliranja i manipuliranja drugim ljudima.

Jedan od sigurno ne manjih elemenata zagonetke zvane budućnost umjetne inteligencije jest ljudsko očekivanje što i kako bi se moglo i kada bi se moglo dogoditi ili događati. Očekivanja znanstvenika, stručnjaka, ali i običnih ljudi, mogućih korisnika proizvoda „*made in AI*“, kao što je i to dosad bilo kod značajnijih otkrića, mogu već sada ubrzati ili usporiti cijeli proces pretvaranja budućnosti (najprije bliske) u sadašnjost. Nažalost, ono što je dosad ubrzavalo trendove najčešće je bila vojna primjena, mnogi ratovi ili njihovi dijelovi vodili su se radi iskušavanja novosti vojno-industrijskog kompleksa. Zadnji, ali ne posljednji, primjer je rat u Ukrajini, posebno kada je riječ o dronovima i korištenju prototipova vojnih robotskih sustava.

S druge strane smo imali možda hvalevrijedan, ali u osnovi uzaludan apel brojnih uglednih znanstvenika da se razvoj i unaprjeđivanje umjetne inteligencije zamrzne na barem šest mjeseci. Koliko god to izgledalo etički i humanistički kako bi se spriječile eventualne zlouporabe istraživanja, doradile karakteristike i uvjeti rezultata te etičke dvojbe kako i zašto istraživati, bila je

to uzaludna i besmislena akcija. Niti se tko od vojnih poglavara osvrnuo na to, niti su veliki znanstveni kompleksi za trenutak ili u nekom manjem udjelu usporili i zaustavili istraživanja. Danas znanost i vojska idu često jedni uz druge i ne pada im na pamet usporavati istraživanja, ratne i društvene eksperimente u kojima se dijeli moć i ekstra profit.

Uostalom, ni sama akcija nije ciljala prave fenomene, ni osobe, a kamoli korporacije. Jer premda se zna tko vodi glavne igre, bilo to korporacije ili pojedini investicijski fondovi kapital kojih se mjeri ne milijardama već i trilijunima dolara, *nomina odiosa sunt*, bolje je ne ciljati točno i izgubiti novce ili utjecaj na svoje redovne djelatnosti.

Nešto slično bila je svojedobno i prava medijska senzacija 1997. godine kada se svjetski šahovski prvak Gari Kasparov suprotstavio tada najjačem računalu *Deep Blue* i izgubio u zadnjoj partiji nakon neriješenih 2,5 : 2,5 bodova. Dotad je dio uglednih znanstvenika, stručnjaka i šahista bio uvjeren da čovjek neće nikad izgubiti od računala. Ili će se to možda dogoditi u daljoj budućnosti? Mediji su tada poludjeli, a tržišna kapitalizacija IBM-a porasla je za vrtoglavih 18 milijardi dolara!

Međutim, valja poslušati glas znanstvenog razuma. Stuart Russell ističe „kako iz perspektive znanosti o umjetnoj inteligenciji ta partija šaha nije predstavljala baš nikakav pomak. Ona je samo nastavila trend koji je već desetljećima bio primjetan. Temeljni dizajn algoritama za igranje šaha još je 1950. godine oblikovao Claude Shannon i uvelike je poboljšana u ranim 60-tim. Nakon toga su se šahovske vještine najboljih programa postupno poboljšavale, uglavnom kao posljedica sve bržih računala koja su dopuštala programima da rade predviđanja u daljoj budućnosti.“⁴

Nakon što je izrađen dijagram vrijednosti i poboljšavanja računalnih šahovskih programa, ekstrapolacija linearnih poboljšanja pokazala je još 1994. godine kako će bodovna vrijednost računalnog programa nadmašiti Kasparovih 2805 bodova 1997. godine. Upravo kada se to i dogodilo! Sljedeći udarac za loše prognozere nenadmašnog ljudskog mozga prema računalima bio je 2015. godine kada je računalno pobijedilo svjetskog prvaka u igri Go, složenijoj od šaha. To je nakon toga učinilo i s nekoliko prvih na ljudskoj rang listi. Ostale su praktički samo strateške video igrice s više ljudi sudionika i protivnika, no i tu noviji programi sve lakše pobjeđuju, još ne u stopostotnom omjeru, ali najnoviji službeni program Alpha pobijedio je već stopostotno prethodni stariji računalni program...

⁴ Russell, Stuart, *Kao čovjek*, (2022) Zagreb: Planetopija, str. 70

Ljudi se više ne igraju s računalima jer ne mogu pobjeđivati, no stvari su već izvan sfere računalnih igara. Ono što se medijski prati uglavnom je još s okusom senzacija i čudnovate priče o „robotima“ koji opet već vuku na već navedene podsvjesne strahove zajedno s tim kako će nas vjerojatno pogoditi meteorit kojeg nećemo izbjeći. Poput vijesti o robotima koji su postali građani Saudijske Arabije ili o robotima koji će postati zvijezde porno industrije.

No Russell, jedan od značajnih aktera u razvoju AI, ističe kako se u istraživačkim laboratorijima događa nešto drugačije – „istraživanje se uglavnom svodi na mnogo razmišljanja, razgovore i ispisivanja matematičkih formula na pločama. Neprestano se stvaraju, napuštaju i ponovno otkrivaju ideje. Čak i zaista dobra ideja, katkad i revolucionarna, često se neće zamijetiti i možda će se tek kasnije shvatiti kako je pružila temelje za značajan napredak na području umjetne inteligencije, kad je netko možda preoblikuje i predstavi u pogodnijem trenutku.“⁵

Dakako, takve su aktivnosti nekako nevidljive ili manje vidljivije, što se tiče masmedija relativno i nezanimljive. Zanima ih više neka priča koja će u skladu s *post truth* vremenom snažno stimulirati emocije gledatelja, slušatelja i čitatelja; više nije pitanje je li neistina ili istina, jer – istina ionako više nije važna. Unatoč takvim „promicateljima“ znanstvenih činjenica ili dokaza, istraživanja idu dalje, ponovno smo u vremenu kada otkrića na polju umjetne inteligencije sustižu jedno drugo, više je pitanje koliko obični ljudi ili publika mogu to i kako podnijeti. U znanstvenim krugovima, bez obzira na njihova ograničenja oko financiranja ili usmjeravanja rezultata prema određenim industrijama, ipak sjede i promišljaju ljudi koji, metaforički rečeno, počinju lagano otvarati vrata budućnosti.

Kako dobro upućuje Mustafa Suleyman u vezi tehnološkog napretka i otkrića – „Tehnologija ima jasnu neizbježnu krivulju: masovno širenje s kotrljajućim valovima. To vrijedi za najranije kremeno i koštano oruđe jednako kao i za najnovije modele umjetne inteligencije...mogućnosti se povećavaju. Eksperimentiraj, ponovi, upotrijebi. Rasti, napreduj, prilagodi se. Ovo je neizbježna evolucijska priroda tehnologije.“⁶ Za njega, od obrade kamena i vatre, preko industrijske revolucije, razvoj se kreće u (nadolazećim) valovima. Njih treba znanošću predosjetiti, istraživati, otkrivati trendove i predviđati širenja, skupljanja, usložnjavanja i nestajanja.

Naravno, nije riječ o ujednačenim procesima poput plime ili oseke, mogu se i nekontrolirano i iznenadno miješati, slabiti i pojačavati. Posljednji dozreli nadolazeći val valja promisliti i razmotriti kako bismo jasnije dobili nagovještaj onog što slijedi, smatra Suleyman.

⁵ Ibid., str. 71.

⁶ Suleyman, Mustafa, *Nadolazeći val*, (2024) Zagreb: Grafički zavod Hrvatske, str. 28.

Razvoj računala su tako u prvim desetljećima poticali novi matematički koncepti, ali i sukobi velikih sila, no kasnije je razvoj, odnosno val, postao neizbježan uz stalno širenje i ubrzano tehnološko napredovanje od tranzistora do vakuumskih cijevi i čipova. Ali i nezaustavljivo – „Alan Turing i Gordon Moore nikada ne bi predvidjeli, a kamoli promijenili uspon društvenih mreža, memova, Wikipedije ili kibernetičkih napada....Neizbježan izazov tehnologije jest u tomu da njezini stvaratelji brzo gube kontrolu nad smjerom i intenzitetom njihova izuma kad se jednom otisne u svijet.“⁷

Umjetna inteligencija, njezina istraživanja i primjena već sada se mijenjaju, razvijaju, šire i nadopunjavaju. Jedan od bitnih elemenata jest i da se u tehnološkom razvoju uz čovjeka i matematičke formule sada u punoj snazi pojavljuju i računala koja pomažu čovjeku da i dalje širi, uslođnjava i ubrzava valove istraživanja i otkrića, smanjuje vrijeme potrebno od ideje do primjene. Na neki način umjetna inteligencija i sama ubrzava svoj razvoj što dosad nije baš bio slučaj s tehnologijama; glavnu ulogu u tomu imao je čovjek. Sada je umjetna inteligencija sposobna i upravljati, kontrolira i popravlja svoje *software* i *hardware*. Može sudjelovati i u stvaranju i opisivanju ideja, projektiranju, premda će se još uvijek ljudi natjecati u tomu da joj pripisuju ili ospore svijest ili razum. No da već može, voljom programera ili slaganjem (vlastitih?) algoritмова određivati brzinu i snagu (re)volucije, svojevrsna je činjenica, htjeli to neki priznati ili ne.

Kada govorimo o umjetnoj inteligenciji, onda za publiku postoji jedna „ugrađena“ greška – pretpostavka kako je AI nekakav računalni program, sustav čipova ili algoritama, izumljen, napravljen i tako ga treba prihvatiti. Međutim, umjetna inteligencija i njen utjecaj danas, posebno u nadolazećim godinama, sve je složeniji „ekosustav“, okruženje u različitim sferama bez kojih opći i dublji razvoj umjetne inteligencije ne bi bio moguć, ali i sam razvoj AI utječe sve više na sve veći broj ukupnog tehnološkog, ekonomskog i društvenog svijeta.

„U početku je okoliš u kojem je većina računala djelovala bio u osnovi prazan i bezobličan: njihovi su jedini unosi dolazili na bušenim karticama, a njihova je jedina metoda proizvodnje učinka bila tiskati znakove na linijskom pisaču.“⁸ Russell konstatira kako je tada (možda i od tada) velik broj znanstvenika pa i javnosti inteligentna računala doživljavao kao puko sredstvo dobivanja nekih odgovora. Prihvatanje ideje ili koncepta da su ti „strojevi“ i mogući subjekti koji percipiraju okruženje i djeluju na temelju zaključaka nije se proširilo sve negdje do 80-tih godina 20.stoljeća.

Međutim, sve više činjenica i realizacije ideja otvorili su nova vrata – od početaka 90-tih godina to je razvoj World Wide Web-a, *soft botova*, sekvence znakova poput URL-a, tzv. dot.com

⁷ Ibid., str. 39.

⁸ Russell, S., *Kao čovjek*, str. 72.

mjehur(1997 – 2000.) koji je omogućio nagli rast i zarade sve brojnijim kompanijama, sustava što su pogodovali potrošačkom društvu i orijentaciji, sve do mobilnih telefona, usluga Amazona ili Google-a, virtualnih asistenata..⁹ Ono što je možda još fascinantnije i uzbudljivije jest kada znanstvenici i stručnjaci zamišljaju i opisuju moguće i po vjerovanju zasad nemoguće aspekte i područja koje će negdje do potpunosti preobraziti umjetna inteligencija.

Ako „preskočimo“ dio dosadašnjeg modela razvoja i pojedinosti u razvoju inteligentnih ne-ljudskih sustava, te se usmjerimo na znanstveni pogled budućnosti čovjeka, tehnologije i društva, moramo najprije konstatirati značajno nesuglasje oko pravaca razvoja, širine i dubine utjecaja i posljedica, posebno kada je riječ o svjesnosti umjetne inteligencije, onoga što, recimo, sada tako nazivamo svijest i opažanje te onog kako definiramo sebe, ljudska bića.

Michio Kaku, primjerice, ukazuje kako ne postoji opće prihvaćena definicija pojma svjesnosti, tako da to onemogućuje i preciznu kvantifikaciju. Prema njemu, početak rješavanja jest analiziranje „tri temeljne komponente: 1. osjećanje i prepoznavanje okoline; 2. svijest o sebi; 3. planiranje budućnosti postavljanjem ciljeva i planova, to jest, simuliranje budućnosti i sastavljanje strategije.“¹⁰ Kod prve komponente tzv. inteligentni roboti imali bi neke primitivnije obrasce prepoznavanja okoline jer imaju osjetilne mogućnosti pojmiti okolinu i bolje od ljudi, ali još uvijek ima poteškoća u prepoznavanju ili razumijevanju onog što vide.

Svijest o sebi problem je i kod životinjskih vrsta; bolje prolaze majmuni, slonovi ili dupini, dok će kod umjetne inteligencije značenje dobiti rješavanje toga kako i na koji način utiskivati obrasce svijesti o sebi ili tumačenja samosvjesnih androida ili kiborga. Spoj neuronskih mreža i kvantnih računala jedan je od puteva prema savladavanju druge komponente. Treća razina nalazi se u blažim oblicima kod nekih životinja, grabežljivaca svjesnijih od plijena, primata koji mogu percipirati kratko u budućnost dijelom i na temelju genetskih pretpostavki, prošlost je temeljena na razumijevanju iskustava.

Planiranje budućnosti, smatra Kaku, zasad je u punom smislu osobina jedino ljudske inteligencije, kao planovi o budućem, razmišljanja „što bi bilo kad bi bilo, te planiranje strategija života i umovanja. Znanost bi trebala u razvoju umjetne inteligencije nastojati podrazumijevati i usavršavati sve tri komponente. „ Planiranje budućnosti zahtijeva zdrav razum, intuitivno razumijevanje onoga što je moguće i konkretne strategije za postizanje specifičnih ciljeva. Da bi

⁹ Više o razvoju i napredovanju AI u društvu i prema ljudskim potrebama vidi u cijelom izuzetno analitičkom, temeljitom i lucidnom poglavlju Russellova knjige: *Kako bi umjetna inteligencija mogla napredovati u budućnosti?*, str. 70. - 110.

¹⁰ Kaku, M., *Fizika budućnosti*, str.97.

robot mogao simulirati stvarnost i predviđati budućnost, potrebno je da savlada milijune zdravorazumskih pravila o svijetu oko njega. Zdrav razum nije dostatan. Zdrav razum pruža samo pravila igre, ali ne i pravila strategije i planiranja.“¹¹

Sve ovo poziva na široku raspravu i promišljanje kakva bi to bila funkcija planiranja i strategije umjetne inteligencije. Najprije u odnosu na ljudsku inteligenciju, a u budućoj naprednijoj formi i sadržaju (o čijem uspostavljanju većina znanstvenika ne spori) i mogućem suživotu AI i ljudske inteligencije. U konačnici i spoju neuro organske tvari i anorganskih kvantnih kompjutora. Poimanje multidimenzionalne stvarnosti – više perspektiva, detalja i dubine – možda bi bio partnerski posao i prožimanje ljudske i nove neljudske inteligencije. Zatim, još preciznije simulacije vremena prošlog, sadašnjeg i budućeg i otkrivanje novih obrazaca - etičkih, socijalnih, ekonomskih, tehnoloških.

Novi zadaci ne-ljudskog inženjeringa ili podizanja svijesti umjetne inteligencije jesu i anticipiranje i prepoznavanje obrazaca koje ljudi zanemaruju, te prepoznavanje i razumijevanje istine i laži, uz sve moguće kategorije poluistina, „poželjnih“ laži ili „nepoželjnih“ istina. To znači i preciznije definiranje samih ciljeva istraživanja i eksperimentiranja, te otkrivanje koliko jesu ili mogu biti snažni otisci ljudskog programiranja. Temeljni ciljevi za čovječanstvo jesu i bit će koliko i kako umjetna inteligencija može biti pomoć ljudima ili bi se moglo pokazati kako su ljudi prepreka (društvenom i tehnološkom) razvoju, pa čak i sebi samima.

U ovom dijelu rasprave o umjetnoj inteligenciji valja uvesti i aktualnu relativno nezapaženu raspravu o etici umjetne inteligencije. To dijelom proizlazi iz već navedenih predrasuda i površnih konstatacija o tome kako je AI vrsta robota s malo više znanja i matematike, a nimalo svijesti. No u teorijskim raspravama o definiciji, dosezima i mogućnostima, filozof John Seal već je 1980. godine formulirao razliku između „slabe“ i „jake“ umjetne inteligencije. Kod tzv. slabe, riječ je o umjetnoj inteligenciji „realiziranoj u računalu ili robotu, koja je u stanju na temelju svojeg programa simulirati ljudsku inteligenciju i učinkovito obavljati neki poseban zadatak. Tipični primjeri su programi za igranje šaha, roboti za kontrolu kvalitete proizvoda na trakama, sustavi za predviđanje želja potrošača u *online* trgovini i filteri za odvajanje neželjene (*spam*) elektroničke pošte. Novi termin „slaba“ umjetna inteligencija može ostavljati pogrešan dojam. Spomenuti primjeri nerijetko su i snažniji – od ljudske inteligencije...“¹²

¹¹ Ibid., str. 99.

¹² Bracanović, Tomislav, *Etika umjetne inteligencije*, (2022) Zagreb: Institut za filozofiju, str. 4.-5.

Slabija umjetna inteligencija podrazumijeva zapravo obavljanje manjeg broja ciljeva i zadataka na nekom ograničenom prostoru i području djelovanja, pa se ponekad upotrebljava i naziv „uska“ ili „uža“ (*narrow*) umjetna inteligencija. Ono što bi se onda nazivalo „jaka“ umjetna inteligencija, zapravo je predmet najveće dvojbe. Ne samo prema Searleu ili preteći McCarthyju pa i dijelu sadašnjih znanstvenika, gdje bi se koristila u nazivu za one strojeve i umjetnu inteligenciju koja bi uglavnom simulirala ljudski um i mogućnosti, već bi za takve „strojeve“ mogli reći da – misle (razmišljaju i promišljaju!). Rasprave već otvaraju pitanja „osobito povezanih sa svojstvima ljudske inteligencije, intencionalnih stanja, fenomenoloških doživljaja, odnosa uma i tijela i drugog. Zapravo je upitno je li takva inteligencija barem načelno moguća i bismo li ikada mogli jednoznačno utvrditi da ona postoji u računalu ili stroju.“¹³

Bracanović vrlo temeljito otvara raspravu napomenama o etici uopće, zatim normativnoj i primijenjenoj praktičnoj etici kako bi precizno utvrdili gdje su i kolike mogućnosti za sučeljavanje već postojećih ljudskih normi i prakse s onima umjetne inteligencije. Vremena je mnogo manje nego u proteklim stoljećima kada su se etički sustavi, kriteriji i norme sporije razvijali, utvrđivali i nalazili primjenu u promjenama društva i ljudskog promišljanja kroz filozofiju i srodne znanstvene discipline.

Međutim, promjene koje već dolaze s ulaskom umjetne inteligencije – i inačica od androida, kiborga, neuromreža i neuroračunala – takve su naravi da će uvelike mijenjati naš način života i razumijevanje, čak i s pojavom suprotstavljenih logika, ideologija i novih „luditskih“ pokreta koji su uzalud nastojali zaustaviti industrijsku revoluciju i parne strojeve. Ovog puta promjene će, najblaže rečeno, biti izazovnije te će otvarati mnoge nove puteve i promišljanja, osobito u etičkoj sferi i načinu poimanja, predočavanja i utemeljenja moralnih stavova i pravila. Posebno jer su sve dosadašnje revolucije i otkrića, tehnička ili društvena, bila pod ljudskom kontrolom, dok se ovdje radi o tehnologijama i načelima koji se mogu odvijati i bez izravnog ljudskog uplitanja, odnosno učinci su brži i bolji i nisu podložni ljudskim greškama i ograničenjima.

Naravno, nisu svi uređaji što ne trebaju ljudsko stalno kontroliranje umjetna inteligencija (primjerice bankomati, termostati, samodijavni uređaji...), ali postavlja se pitanje što ćemo s onima koji neće biti unaprijed programirani da samostalno obavljaju i rješavaju zadatke ili popravljaju slične strojeve i još ujedno uče iz vlastitog iskustva i donose samostalne odluke na temelju ispravnijeg i/ili svrsishodnijeg poteza i djelovanja. Novi etički izazovi jesu i bit će „tko je odgovoran za eventualne štetne posljedice takvih sustava? Bi li takve sustave trebalo *etički*

¹³ Ibid.- str. 7.

programirati i ograničiti njihovo djelovanje neakvim etičkim kodeksom? Smiju li takvi sustavi donositi odluke o ljudskim sudbinama ili čak i o životu ili smrti?“¹⁴

Još jedna bitna karakteristika rasprave o razvoju umjetne inteligencije jest i pitanje koliko smo dosad obratili pozornost usavršavanju jezika kojim se služi umjetna inteligencija na području od virtualnih asistenata do korištenja pisane riječi u odgovorima poput onih Chat GPT-a. Roger Penrose, u svojim razmišljanjima o sadašnjosti i mogućim budućnostima računalne snage i ciljeva širi i produbljuje sferu na nužno razmatranje i povezivanje načela razuma i razumnog djelovanja s principima teorijske i kvantne fizike. Također, tu je i Jacques Hadamard koji je istražujući kreativno mišljenje i njegovu primjenu u znanstvenim disciplinama zaključio kako verbalizacija nije neophodna za inteligentno mišljenje. Penrose konstatira kako bez obzira na vezu i činjenicu da razmišljamo u riječima i kroz riječi smatra „riječi gotovo nekorisnim za matematičko razmišljanje. Druge vrste razmišljanja, poput filozofiranja, čine mi se daleko bolje prilagođene verbalnom izražavanju. Možda zbog toga brojni filozofi smatraju kako je jezik bitan za inteligentnu ili svjesnu misao!“¹⁵

Penrose detaljno i temeljito razmatra i razlučuje načela i funkcioniranje ljudskog uma i umovanja koje bi izviralo iz usavršenih algoritama i proračuna. Premda želi ukazati na neodrživost stajališta i pretpostavki uspoređivanja našeg razmišljanja s djelovanjem vrlo složenog računala. To i jest jedno od pitanja budućnosti, da li izvođenje vrlo složenih algoritama na neki način stimulira i izaziva buđenje svjesnosti u sustavu.

Međutim, na temelju svojih razmatranja o prirodi i prirodnim zakonima, te zakona koji prožimaju svemir i specifičnih „proračuna“ dijelovi kojih su i naš mozak i um, Penrose konstatira i kako su fizički i makrofizički svijet, odnosno „svi mi samo mali dijelovi svijeta kojim upravljaju detaljni (premda s tek malo vjerojatnosti) matematički zakoni. Samim našim mozgom, koji, izgleda, upravlja svim našim akcijama, također upravljaju isti točni zakoni. Međutim, zaista postoji nešto bitno što nedostaje u bilo kojoj čisto proračunskoj slici. Zahvaljujući prirodnim znanostima i matematici jednom ćemo nužno dospjeti do velikih napredaka u razumijevanju uma. Postoji mnogo zagonetnosti i ljepote koliko se samo može poželjeti u točnom platonovskom svijetu matematike, nalazi se u pojmovima koji leže izvan relativno ograničenog izračunljivog dijela u kojem stanuju algoritmi i izračuni.“¹⁶

¹⁴ Ibid., str. 15.

¹⁵ Penrose, Roger, *Carev novi um*, (2004) Zagreb: Izvori, str. 456.

¹⁶ Ibid., str. 478.

Dok se etika umjetne inteligencije kao posebna disciplina tek treba uspostaviti šire i temeljitije, obzirom na budućnost koja je zapravo vrlo blizu, ono što se događa upravo sada, ovom dijelu godine, vrlo je izazovno i obećavajuće. Posebno jer povezuje, možda i više nego očekujemo ili smo očekivali, teoretski i praktični dio, odnosno sfere razvoja umjetne inteligencije. Upravo je OpenAI, istraživački laboratorij za umjetnu inteligenciju predstavila novi sustav klasifikacije koji prati napredak istraživanja i dostignuća prema stvaranju AI softvera sposobnog nadmašiti ljudske sposobnosti.

Novi sustav klasifikacije je prvo interno predstavljen zaposlenicima kompanije, a iz njega je dobro vidljivo na koji način se zapravo postavljaju smjernice razvoja i težišta u istraživanju, razvoju i korištenju umjetne inteligencije.

Razine AI prema OpenAI¹⁷

Razina 1	<i>Chatbotovi</i> , AI s konverzacijskim jezikom
Razina 2	<i>Mislioni</i> , rješavanje problema na ljudskoj razini
Razina 3	<i>Agenti</i> , sustavi koji mogu poduzimati akcije
Razina 4	<i>Inovatori</i> , AI koja može pomoći u izumima
Razina 5	<i>Organizacije</i> , AI koja može obavljati posao cijele organizacije

Razine napretka kreću se od trenutno dostupne AI tehnologije koja može komunicirati s ljudima na konverzacijskoj razini (Razina 1) do AI sustava sposobnih obavljati posao čitave organizacije (Razina 5). Prema klasifikaciji OpenAI-ja, treća razina na putu prema općoj umjetnoj inteligenciji naziva se "Agents" (Agenti) i odnosi se na AI sustave koji mogu djelovati u ime korisnika. Četvrta razina opisuje AI koji može stvarati inovacije, dok se najnaprednija razina naziva "Organizations" (Organizacije).

Razine nisu postavljene sa čvrstim granicama, upravo kako se očekuje da bi tekao i budući napredak u istraživanju i primjenama. Bitno je znati kako u kompaniji vjeruju kako se aktualno nalaze na prvoj razini, ali i pred dostizanjem druge razine koju nazivaju "Reasoners" (Mislioci). Ova razina odnosi se na sustave koji mogu rješavati osnovne probleme jednako dobro kao i ljudi s doktoratom, ali bez pristupa dodatnim alatima.

¹⁷ Vidi Podnar, Ivan, *OpenAI klasificirao pet razina razvoja umjetne inteligencije: od konverzacije do organizacije*, <https://www.bug.hr/umjetna-inteligencija/openai-klasificirao-pet-razina-razvoja-umjetne-inteligencije-od-konverzacije-do-42264>, Pristup 24.11.2024.

Ovdje treba spomenuti, u skladu s dosad razmatranim, kako znanstvenici i stručnjaci koji se bave umjetnom inteligencijom, pored matematičkih, algoritamskih i tek dodirnutih etičkih aspekata, već duže vremena razmatraju i odabir elemenata kriterija određivanja kada bi umjetna inteligencija, uvjetno rečeno, mogla dostići pa i premašiti ljudske sposobnosti.

Tako je u mjesecu studenom 2023. godine dio istraživača iz Google DeepMind-a predložio i okvir od pet uzlaznih razina AI-ja, uključujući kategorije poput "stručnjak" i "nadjudski". Ove razine podsjećaju na sustav koji se često koristi u automobilske industrije za procjenu stupnja automatizacije autonomnih vozila, ali se odnosi i na proširene zadatke, odnosno ciljeve razvoja viših razina onog što se naziva umjetna inteligencija.

Također, tada je vodstvo tvrtke predstavilo i istraživački projekt koji uključuje njihov GPT-4 AI model. Demonstracija je, prema izvoru bliskom Bloombergovim novinarima, pokazala neke nove vještine koje se približavaju ljudskom rasuđivanju. Dakako, ono što se ne tako često navodi kod ovih tema jesu i ekonomski i investicijski, odnosno profitni plan i ciljevi – Google je već krenuo u podjelu ovih informacija s investitorima i drugim zainteresiranim stranama izvan kompanije.

Prema agenciji Reuters, OpenAI, kreator virtualnog pomoćnika ChatGPT, radi na novom pristupu svojoj tehnologiji umjetne inteligencije. Kao dio projekta kodnog naziva *Strawberry*, firma koju podržava Microsoft pokušava se drastično poboljšati sposobnosti rasuđivanja AI modela. Naravno, način na koji *Strawberry* radi je "strogo čuvana tajna" čak i unutar samog OpenAI-a, a projekt uključuje "specijalizirani način" obrade AI modela nakon što je prethodno obučen za opsežne skupove podataka. Cilj je omogućiti umjetnoj inteligenciji ne samo da generira odgovore na upite, već i dobroj mjeri planira unaprijed kad sprovede tzv. "duboko istraživanje", autonomno i pouzdano navigacijom Internetom.

Glasnogovornik OpenAI-ja rekao je za Reuters: „Želimo da naši modeli umjetne inteligencije vide i razumiju svijet više kao mi ljudi. Kontinuirano istraživanje novih mogućnosti umjetne inteligencije uobičajena je praksa u industriji, uz zajedničko uvjerenje kako će se ovi sistemi vremenom poboljšati u zaključivanju.”¹⁸ Istraživači-znanstvenici koji su razgovarali za Reuters, rekli su kako je rasuđivanje, koje je do sada izmicalo modelima umjetne inteligencije, ključno za postizanje ljudske ili tzv. nadjudske razine umjetne inteligencije.

¹⁸ Vidi Mišo T., *OpenAI, kreator virtualnog pomoćnika ChatGPT, radi na novom pristupu svojoj tehnologiji umjetne inteligencije*, <https://www.logicno.com/zivotni-stil/proizvodac-chatgpt-a-tajno-razvija-novu-vrstu-ai.html> Pristup 24.11.2024.

Ali u vezi s ovim, Yoshua Bengio, jedan od vodećih svjetskih stručnjaka za umjetnu inteligenciju i pionir *deep learninga* u AI, ponovo je upozorio na “mnoge rizike”, uključujući moguće velike opasnosti za čovječanstvo, koje predstavljaju korporacije utrkujući se postići ljudsku razinu umjetne inteligencije. Istakavši pitanje: “*ako postoje entiteti koji su pametniji nego ljudi i imaju svoje ciljeve, jesmo li sigurni da će djelovati na našu dobrobit?*” Bengio, profesor Sveučilišta u Montrealu i znanstveni direktor Montrealskog Instituta za algoritme učenja (MILA), pozvao je znanstvenu zajednicu i društvo da ulože “masovne kolektivne napore” za držanje pod kontrolom buduću, napredniju umjetnu inteligenciju.

Bez obzira na sve dosadašnje fascinantne mogućnosti i primjene poput Chat GPT-a, iako se na njih tek navikavamo, kao i naznake budućih sposobnosti AI sustava, sve su to prema znanstvenicima i stručnjacima zapravo tek počeci velike promjene koja slijedi. Nazivaju je tihom revolucijom, a to su tzv. Generativne kontradiktorne mreže ili GANs (engleski *Generative Adversarial Networks*).

Riječ je o doslovno genijalnom novom pristupu u domeni AI-ja koji dovodi dvije neuronske mreže (modeli računalnog učenja koji su kreirani po uzoru na mozak živih bića) na način da su jedna drugoj suprotstavljene. Kroz taj sraz te dvije mreže mogu kreirati nevjerovatno stvarne ishode, bilo da je konačni proizvod slika, zvuk, video zapis, ali i puno više od toga. GAN-ovi predstavljaju temeljnu promjenu u samom načinu kako percipiramo, ali i kako koristimo moć umjetne inteligencije.¹⁹ Recimo ovako: zamislite da imate dva prijatelja. Jedan je umjetnik, a drugi je detektiv. Njih dvoje se nađu i posao umjetnika je, jasno, da slika slike. Detektiv pak ima zadatak otkriti jesu li te slike koje je naš umjetnik naslikao stvarne ili “lažne”, odnosno predstavljaju li ono što mu je zadano kao tema. Pritom je jedan detalj jako važan - umjetnik želi navesti detektiva da misli da slike jesu stvarne!

U svojoj osnovi tako GAN-ovi funkcioniraju. Jedan dio mreže kojeg smo nazvali umjetnikom, a stručno se zapravo zove “generator”, stvara nešto - možda slike, glazbu, svašta... U isto vrijeme detektiv, odnosno “diskriminator” (stručni naziv), neumorno radi na tome da definira je li uradak stvaran ili lažan. Zapravo, uče jedan od drugog i kreću se prema konačnom proizvodu koji izgleda stvarno.

Prva mreža, generator, uzima nasumičnu “buku” podataka kao input i nastoji je pretvoriti u sintetički proizvod, recimo sliku, zvuk, video zapis ili pak trodimenzionalne modele. Konačni cilj

¹⁹ Vidi Marko K. *Tiba revolucija puno veća od umjetne inteligencije s kojom smo već upoznati*: <https://www.advance.hr/tekst/tiha-revolucija-puno-veca-od-umjetne-inteligencije-s-kojom-smo-vec-upoznati-generativne-kontradiktorne-mreze-kljuc-su-jos-nevidene-kreativnosti-i-moci/>

je doći do toliko realističnog proizvoda koji se više ne može razlikovati od stvarnih podataka i primjera. Riječ je o ogromnom broju međusobnih komunikacija u kojima diskriminator stalno procjenjuje konačan rezultat po pitanju je li isti stvaran ili "lažan" - i pritom svaki put daje svoju ocjenu vjerojatnosti stvarnosti. Generator na svaku povratnu informaciju uzvraća boljim primjerom stvarnosti i taj broj komunikacija se stalno povećava. Na kraju će se izjednačiti, umjetno i stvarno će se "preklopiti" u jedno. Znači, "tiha revolucija" donosi nam budućnost u kojoj nećemo znati gledamo li fotografiju stvarne ili "izmišljene" mačke ili psa?

Vrijedi spomenuti i napredak izvan kompanija OpenAI, odnosno Googlea, riječ je o kineskom napretku - DeepSeek, tvrtke za istraživanje umjetne inteligencije, koja je omogućila pristup svom modelu umjetne inteligencije DeepSeek-R1. Za razliku od većine modela te vrste, modeli rasuđivanja učinkovito sami provjeravaju činjenice trošeći više vremena na razmatranje pitanja ili upita. To im pomaže u izbjegavanju nekih od zamki koje inače stvaraju probleme. Tako i DeepSeek-R1 razmišlja kroz zadatke, planira unaprijed i izvodi niz radnji koje pomažu u dolaženju do odgovora. To može potrajati. DeepSeek-R1 može 'razmišljati' nekoliko desetaka sekundi prije nego što odgovori, ovisno o složenosti pitanja. DeepSeek tvrdi kako njihov model radi jednako kao i OpenAI-jev model na dva alata za mjerenje: AIME i MATH. AIME koristi druge modele za procjenu izvedbe modela, dok je MATH zbirka problema s riječima.²⁰

Dakle, bez obzira u kojoj se fazi nalaze danas najaktualniji proboji u buduću svijet umjetne inteligencije (ako pretpostavimo da je glavni dio ledenog brijega AI istraživanja ipak prikriven), slijede otkrića i objave novosti ne samo u jezičnim, odnosno verbalnim napredovanjima „strojeva“, također i oko moći rasuđivanja i zaključivanja AI, a čini se već i oko razumijevanja dijela emocija i proizvođenja sličnih reakcija. Primjerice, možda i nije teško predviđati i novi pristup AI gdje dvije neuronske mreže (modeli računalnog učenja po uzoru na živa bića) jesu suprotstavljene jedna drugoj, ali i komuniciraju i „osjećaju“ jedna drugu?!

Nadalje, pojavljuju se mogućnosti kreiranja nevjerojatnih ishoda – riječi, zvuka, slike, pokreta. Dolazi i doći će do temeljnih promjena u percipiranju AI-ja (i njegovog percipiranja nas!), što vodi i do korištenja različitih mogućnosti, dakle i moći. Hoće li to biti predmet najjačih nadmetanja ljudi i viših razina umjetne inteligencije? Ako u navedenom slučaju dio mreže kao

²⁰ Wranka Miroslav, *I Kinezi imaju svoj OpenAI, pogledajte što su ponudili*, <https://www.tportal.hr/tehnolo/clanak/i-kinezi-ima-ju-svoj-openai-pogledajte-sto-su-ponudili-20241121> Pristup 23.11.2024.

umjetnik slika mačku, drugi dio mreže kao detektiv otkriva jesu li one „laž” ili stvarne, a umjetnik navodi detektiva da misli kako su stvarne, što je onda stvarnost u mnogo ozbiljnijim stvarima nego jest slikanje mačke?

Novi ili izmijenjeni fenomeni umjetne inteligencije, ponajviše u kontrakciji i/ili suživotu ljudskih bića s androidima, kiborzima i humanoidnim bićima sličnim čovjeku, tražit će nove pristupe i definicije. „Na praktičnoj razini nužna je nova i drugačija paradigma mišljenja, usmjerena na način, granice i primjene znanja novih tehnoloških područja, gdje su etičke vrijednosti i norme u prvom planu. Znanstvenik ne može biti moralno neutralan jer je upravo moderna fizika pokazala kako su obrasci koje znanstvenici zapažaju u prirodi povezani s obrascima i vrijednostima njihova duha. Isto tako, treba utvrditi ciljeve i definirati norme koji će prisiliti sve one koji su uključeni u kreiranje novih postupaka i materijala i njihovo implementiranje u ljudski organizam, te sve one koji iz različitih interesnih i medicinskih razloga donose odluke o poboljšanju čovjeka natjerati i da prihvate neophodnu količinu etičke odgovornosti.”²¹

Korak dalje zasigurno je i metaforično-alegorijski prikaz suvremenosti i budućnosti medija kroz pojam izvanzemaljskog ili vanzemaljskog, kao i vanvremenskog, računajući od sadašnjosti, što u prvom značenju i ne obuhvaća samo nešto što nije porijeklom sa Zemlje. Taj je pojam zapravo dio ključa za razrješenje zagonetke otuđenih svjetova i otuđenog društva i pojedinaca koje proizvodi.

Sam svijet i (živi i neživi) mediji su zbroj i spoj takvog izvanzemaljskog (nezemaljskog?) i ljudskog bića, sadašnjosti i medijskih mogućnosti u budućnosti gdje bi se Čovjek trebao osjećati i ponašati sve više kao vanzemaljac (a i danas takav osjećaj nije više rijetkost) u odnosu na klasične i dosada prihvaćene norme i pravila (mas)medijskog svijeta. Promjene koje dolaze stvarat će strance i/ili vanzemaljce (*aliene* sa Zemlje!) od već postojećih jedinki što će opet dovesti do potpuno novih i “još novijih” vrsta bića i njihovih medija, posrednika između svijeta iznutra i svijeta izvana i oko nas/njih. Negdje s ruba i na rubu znanstvene fantastike dolazi tehnologija koja će nas vjerojatno izmijeniti i više nego što smo mislili. Hoće li tehnologija sama i dalje biti ovisna o ljudskim nedostacima i hoće li se analizirati dosada uobičajenom pojmovima? Ili će trebati – nije li ovo zapravo retoričko pitanje – i jednu drugačiju, izmijenjenu filozofiju, filozofiju budućih vrsta i medija, drugačiju filozofiju čovjeka i medija?²²

²¹ Vertovšek, Nenad; Greguric, Ivana, *Filozofija budućeg: ogledi o neljudskom*, (2021) Zagreb: Jesenski i Turk. Str. 151. Vidi i cijelo poglavlje *Ontološka pitanja kiborizacije čovjeka i nužnost prolegomene o kiborgoetici*.

²² Vertovšek, Nenad, *Drveno željezo medija*, (2020) Nikšić: Medijska kultura, str. 127. Vidi i cijelo poglavlje *Ako umjetna inteligencija postane stvarna filozofija*.

Ovdje nije riječ (samo) o nagovještajima robotskih inteligencija ili robotskih filozofija jer bi pogled koji seže samo do *strojeva-koji-misle-kao-ljudi* imao jednu osnovnu „grešku“ ili nesavršenost sustava – još bi uvijek bio human(oidan). Što s mogućim novim apsurdom ili, još gore, oksimoronom – je li pred nama budućnost posredovanja i s tim vezane filozofije koji nisu ljudski? Ako budu kiborške ili vanzemaljske, kako će i koliko sadržavati ljude? Ali recimo još jednom, to neće biti nužno dobri ili loši ljudi kao mediji (pa ni filozofija) ili dobre ili loše promjene koje će oni donijeti, posebno prema sadašnjim “zastarjelim” medijskim mjerilima promjena. Bit će sigurno – vrlo drugačije promjene. Gotovo ih neće moći prepoznati ni tada već izmijenjeni ljudi...

Literatura i Internet izvori :

- Bracanović, Tomislav, *Etika umjetne inteligencije*, (2022) Zagreb: Institut za filozofiju
- Davenport, H. Thomas, *Prednost umjetne inteligencije: kako iskoristiti revoluciju umjetne inteligencije?*, (2021) Zagreb: MATE d.o.o.
- Denning, J. Peter; Tedre, Matti, *Računalno razmišljanje*, (2021) reb: MATE d.o.o.
- Kaku, Michio, *Fizika budućnosti*, (2011) Zagreb: MATE d.o.o.
- Marinković, Renato, *Inteligentni sustavi za poučavanje*, (2004) Zagreb: Hrvatska zajednica tehničke kulture
- Penrose, Roger, *Carev novi um - razmišljanja o računalima, razumu i zakonima fizike* (2004) Zagreb: Izvori
- Russell, Stuart, *Kao čovjek*, (2022) Zagreb: Planetopija
- Stipaničev, Darko; Šerić, Ljiljana; Braović Maja, *Uvod u umjetnu inteligenciju*, (2021) Split: Fakultet elektrotehnike, strojarstva i brodogradnje
- Suleyman, Mustafa, *Nadolazeći val*, (2024) Zagreb: Grafički zavod Hrvatske
- Vertovšek, Nenad, *Drveno željezo medija*, (2020) Nikšić: Medijska kultura
- Vertovšek, Nenad; Greguric, Ivana, *Filozofija budućeg: ogledi o neljudskom*, (2021) Zagreb: Jesenski i Turk
- Vertovšek, Nenad; Greguric, Ivana, *Philosophies of the Future and the Non-human: from Cyberspace to Human Enhancement*, (2023) Cambridge: Cambridge Scholars Publishing
- Marko K., Tiha revolucija puno veća od umjetne inteligencije s kojom smo već upoznati:
<https://www.advance.hr/tekst/tiha-revolucija-puno-veca-od-umjetne-inteligencije-s-kojom-smo-vec-upoznati-generativne-kontradiktorne-mreze-kljuc-su-jos-nevidene-kreativnosti-i-moci/>
Pristup 21.11.2024.
- Mišo T., OpenAI, kreator virtuelnog pomoćnika ChatGPT, radi na novom pristupu svojoj tehnologiji umjetne inteligencije,
- Podnar, Ivan, OpenAI klasificirao pet razina razvoja umjetne inteligencije: od konverzacije do organizacije, <https://www.bug.hr/umjetna-inteligencija/openai-klasificirao-pet-razina-razvoja-umjetne-inteligencije-od-konverzacije-do-42264> Pristup 21.11.2024.
- Wranka Miroslav, I Kinezi imaju svoj OpenAI, pogledajte što su ponudili,
<https://www.tportal.hr/tehnolo/clanak/i-kinezi-imaju-svoj-openai-pogledajte-sto-su-ponudili-20241121> Pristup 23.11.2024.