

Intelligenza artificiale e stoltezza umana*

Stefania Sinigaglia

A partire da un articolo a proposito dell'ultima versione del robot parlante del novembre 2022, ChatGPT4, si sviluppa una riflessione sui risvolti di questo uso della tecnologia di I.A. a fini commerciali e quindi di profitto e se ne sottolineano le implicazioni, i limiti e i pericoli nell'assenza di un severo quadro normativo. Ciò in contrasto con l'impiego indispensabile e proficuo dell'IA in campo scientifico e per fini socialmente utili. Infine si commenta una lettera aperta di vari esponenti del mondo tecnologico della IA che propugna una breve sospensione dello sviluppo della stessa, e si illustra l'ideologia che sta dietro a questa richiesta, ideologia che si ritiene pericolosa e nefasta per vari motivi.

“In 1950, the English computer scientist Alan Turing devised a test he called the imitation game: could a computer program ever convince a human interlocutor that he was talking to another human, rather than to a machine?”¹

Tale gioco fu poi battezzato: test di Turing.

Così inizia un articolo di Billy Perrigo del dicembre 2022 sulla rivista *Time*, che di seguito menziona il fatto che Google aveva qualche mese prima licenziato uno dei suoi dipendenti proprio perché questi si stava convincendo che uno dei suoi *chatbots* (robot capaci di intrattenere una conversazione plausibile) fosse diventato “senziente”, cosciente quindi. Il *chatbot* aveva superato il test di Turing, pur non essendo certamente cosciente di ciò che diceva. Il giornalista riproduce in seguito la sua conversazione con un *chatbot* diventato famoso negli ultimi mesi, ChatGPT, presumibilmente ChatGPT4, sfornato nel novembre 2022 da Open AI, una start-up nata negli Stati Uniti nel 2015 nel cui Consiglio di Amministrazione erano presenti sia Elon Musk che Sam Altman, quest'ultimo diventato ora Amministratore Delegato. A partire dal chiasso che il lancio sul web di questo robot parlante ha fatto, sembra che la questione serissima del ruolo preponderante che l'intelligenza artificiale sta assumendo in ogni campo *ora* (e le direzioni in cui procede) si stia trasformando artatamente in un dibattito stampa-



* L'articolo è tratto dal blog *La Croce e l'Orsa*, <http://croceorsa.blogspot.com/>

1. Perrigo B, *AI chatbots are getting better. But an interview with ChatGPT reveals their limits*, *Time*, 5-12-2022 <https://time.com/6238781/>

chatbot-chatgpt-ai-interview/

lato sulla minaccia esistenziale che queste abili macchine chiacchierine (e spesso pasticciona) o disegnatrici, auto-trasformatesi in loro avatar ben più potenti *in futuro*, potrebbero rappresentare. Mutatesi in mostri super-intelligenti potrebbero diventare un pericolo per la sopravvivenza della specie umana, in tot anni, variabili a seconda degli autori delle previsioni apocalittiche, molti dei quali hanno contribuito alla creazione dei possibili futuri mostri.

Nel presente, il *chatbot* con il quale si intrattiene Billy Perrigo dice chiaramente di essere semplicemente un modello linguistico statistico², che quindi formula i suoi enunciati sulla base di associazioni probabilistiche fondate su miliardi di dati (input) ingurgitati dal web, costruito con miliardi di parametri con una tecnologia di *deep learning* (apprendimento profondo) usando reti neurali artificiali che imitano l'apprendimento umano³. Quindi non sente, non pensa, non è cosciente di quel che dice, è un abile pappagallo stocastico⁴. L'acronimo GPT significa "generative pre-trained transformer", cioè indica che il robot è in grado di generare autonomamente frasi appropriate a seconda della domanda, selezionando dalle montagne del suo database le risposta più probabile sulla base di una tecnologia di *machine learning*, algoritmi che imparano a partire da altri algoritmi di apprendimento.

C'è anche un livello superiore ed è quello che suscita gli incubi futuribili di mostri super-intelligenti, più intelligenti degli umani: il *meta-learning* di questi ordigni, cioè il fatto che imparando ad imparare si possano perfezionare *ad libitum*. Infatti si prospetta che l'intelligenza artificiale (AI in acronimo inglese) potrebbe trasformarsi in AGI, cioè in intelligenza artificiale generale, non limitata a compiti specifici prede-

finiti. Da pappagalli stocastici, modelli statistici controllabili, i robot si trasformerebbero in super-macchine che ragionano autonomamente con potenza cognitiva superiore a quella di menti umane, prevaricandole, con conseguenze potenzialmente catastrofiche. Vedremo chi sono i preveggenti àuguri di sventure virtuali. Se ChatGPT4 crea già oggi testi, saggi, articoli su richiesta che rivaleggiano con testi scritti da professionisti, un altro tipo di robot (come *Midjourney*) crea e manipola immagini e video su indicazioni (*prompts*) via via più precise, fino al prodotto finale desiderato. È chiaro che esiste un grosso problema autoriale nella produzione e pubblicazione di questi contenuti, linguistici e grafici, non solo di copyright. E di responsabilità. Si è ingaggiata una nobile gara di corsa tra i colossi informatici privati Google, Microsoft e Amazon, ognuno con i propri ordigni di intelligenza artificiale⁵, mentre varie start up in poco tempo vengono quotate in borsa e aumentano il valore delle loro azioni, come appunto Open AI, che dopo il successo di ChatGPT4 è stata comprata da Microsoft. L'addestramento delle macchine intelligenti consuma un'immensa quantità di energia e costa molti miliardi di dollari. Siamo sicuri che ne valga la pena? A che scopo chiacchierare con un robot? La traduzione istantanea è utile, ma perché far creare loro testi o video? Per risparmiare soldi, forza lavoro, dipendenti, cos'altro? A chi giova? Analisi costi/benefici? Analisi dei rischi? Quali problemi si vogliono risolvere con questa costosa linea di ricerca⁶? E i comuni mortali come possono dire la loro e intervenire su decisioni che comunque li toccano già da vicino: privacy, cyber-sorveglianza, modelli di informazione, pubblicità, spettacoli, cultura, ecc.? Il mercato del lavoro, il giornalismo, la medicina, la politica, l'Università funzionano

2. Più precisamente, un prodotto dei *Large Language Models*, reti neurali artificiali con miliardi di parametri, vedi Wikipedia per approfondire. https://it.wikipedia.org/wiki/Large_language_model

3. Si noti che qui per apprendimento si intende l'assorbire acriticamente tutto lo scibile propinato da montagne di dati presenti sul web, il che farebbe rabbrivire non solo pedagogisti come i compianti Freinet e Freire, ma anche chi sia dotato di un minimo di cognizioni

in materia o anche solo di buon senso. E il nostro cervello non funziona solo con neuroni.

4. La definizione è di Timnit Gebru, ricercatrice informatica ex dipendente di Google.

5. Microsoft ha comprato Open AI e si è aggiudicato ChatGPT, Google ha sviluppato *Bard*, Amazon ha il suo *Bedrock*, Apple ci sta pensando, come META, ecc.

6. Si presume che ogni ricerca debba partire da un problema da risolvere.

sempre più con algoritmi di intelligenza artificiale. L'apprendimento scolastico ne è stato modificato e forse non in direzioni tali da coltivare attitudini critiche.

Dopo avere letto sull'argomento cercando di districarmi tra i tecnicismi, mi sembra di essere arrivata ad un punto fermo: c'è un campo vastissimo di operatività scientifica in settori come la medicina, l'aeronautica, l'astrofisica, la fisica delle particelle, le scienze sociali, la climatologia in particolare, ed altri senza dubbio, per cui l'intelligenza artificiale pubblica è un indispensabile fattore di sviluppo a vantaggio sia della ricerca fondamentale che applicata, a vantaggio dell'umanità in ultima analisi.

Esempi di ricerca pubblica non specializzata ce ne sono: il 28 aprile 2021 è stato lanciato il più grande progetto di IA che riunisce più di 250 ricercatori facenti capo a un centinaio di laboratori e imprese di una decina di paesi⁷ (CNRS, Inria⁸, Università, Renault, Airbus, Ubisoft, Orange, Facebook, Systran⁹, e altri). Tale progetto è stato battezzato *Big Science*¹⁰, è open source, i suoi parametri e il suo codice informatico sono pubblici, è multilingue. La potenza di calcolo è assicurata dal supercalcolatore Jean Zay, a Orsay, che mette a disposizione 5 milioni di ore di calcolo. Il titolo in francese dell'articolo su *Le Monde* del 5 maggio 2021 è: *Projet géant pour faire parler une IA*, (Progetto gigantesco per far parlare una intelligenza artificiale). *Big Science* si propone di correggere i numerosi difetti dei concorrenti privati: monolingui, opachi, mal controllati e soggetti a molti rischi, come la produzione di testi fuorvianti, offensivi a volte. (Aggiungerei, volti al profitto.) Nell'articolo si sottolinea: vari gruppi studieranno le questioni etiche connesse all'impresa. In contrasto, si noti che Google ha licenziato, tra dicembre 2020 e

febbraio 2021, due ricercatrici in materia di etica dell'intelligenza artificiale, Timnit Gebru e Margaret Mitchell, che lavoravano sui rischi impliciti nei modelli linguistici. La direzione di ricerca di *Big Science* è la stessa dei concorrenti privati, per cui gli interrogativi circa costi/benefici/fini rimangono, ma le sue caratteristiche e i *caveat* sono più rassicuranti, oltre al fatto che si tratta di un progetto pubblico.

Un altro progetto, spiccatamente di indirizzo andragogico¹¹, è partito da una ricercatrice universitaria, Isabelle Guyon, franco-svizzera-americana, che con una sua piccola impresa informatica (ClopiNet) ha ideato una macchina che lancia gare pubbliche di apprendimento nel campo della IA e del *machine learning*, capace di valutare meglio gli algoritmi. La ricercatrice ha poi promosso un progetto pubblico, *Humania*, il cui scopo dichiarato è quello di democratizzare l'intelligenza ar-



Isabelle Guyon (foto Le Monde, 2018)

7. <https://bigscience.huggingface.co/blog/model-training-launched>

8. Institut de recherche en informatique et en automatique.

9. Società che sviluppa software di traduzione.

10. Larousserie D, *Projet géant pour faire parler une IA*, Le Monde Science et Médecine, 28-4-2021

<https://bigscience.huggingface.co/blog/model-training-launched>

11. Andragogia è la scienza sociale che si occupa dell'educazione degli adulti, il cosiddetto *life-long learning*, l'educazione continua, così poco attuato in Italia.

tificiale, finanziato dalla Agenzia Nazionale della Ricerca francese, ANR¹². Sulla pagina di presentazione del sito ci sono due video per principianti, e poi una lunga lista degli esercizi delle gare già fatte. Tutto liberamente fruibile e anche divertente. Ma sembra che simili iniziative siano rare.

Sappiamo che il progresso tecnologico evolve in sistemi capitalistici, e oggi, in un periodo di neoliberalismo trionfante in cui il fattore profitto è centrale, come la concorrenza a livello commerciale, industriale, scientifico, militare e politico, la cooperazione è spesso ignorata. L'insorgere del finanzia-capitalismo, come l'ha giustamente chiamato il sociologo Luciano Gallino¹³, esaspera l'estrazione di valore da ogni brandello e angolo della terra e da ogni essere umano, tanto più da ogni conquista tecnologica e scientifica per di più piegata a fini che con il miglioramento della condizione umana hanno ben poco a che vedere. Così abbiamo da anni droni e bombe intelligenti che cadono su feste di matrimonio o su funerali, teleguidate da migliaia di chilometri di distanza. Un documento recente del SIPRI¹⁴, Istituto svedese di ricerca sulla pace con sede a Stoccolma, inizia constatando che l'IA è al centro della concorrenza in campo militare tra le grandi potenze mondiali, in particolare tra USA e Cina, e per quanto concerne l'Europa sollecita l'UE a proporre comuni principi da rispettare nell'uso militare dell'IA. A che punto siamo ora? Il contesto attuale di belligeranza, di accanimento spionistico e di rivalità tra potenze con le rispettive sfere di influenza non promette niente di buono. Le spese nazionali in armamenti stanno crescendo a dismisura.

In campo medico l'intervento di affidabili robot-chirurghi controllati remotamente da spe-

cialisti può da un lato moltiplicare la capacità operativa ospedaliera¹⁵ (un robot di questo tipo costa minimo 1,4 milioni di euro), aumentando l'efficienza, e rendere disponibili un numero maggiore di letti ai malati, si dice. Ma in caso di un qualsiasi errore, chi sarà responsabile? E che dire del rapporto di fiducia che deve esistere tra malato e chirurgo? Lo scopo principale è liberare letti più rapidamente e risanare il bilancio degli ospedali o migliorare il benessere dei pazienti e il numero degli interventi chirurgici con esito soddisfacente? In un'ottica di sanità pubblica globale, in termini di patologie che causano maggiore morbidità e mortalità, è più consigliabile investire in robot chirurghi o in formazione di medici dove occorre, in fornitura di acqua potabile¹⁶, servizi igienici e in zanzariere imbevute di insetticida¹⁷ dove si tratta di vita o di morte? E la prevenzione buttata alle ortiche?



Un robot chirurgo

Se è chiaro lo scopo in ricerca fondamentale nel campo della fisica del Large Hadron Collider al CERN di Ginevra, il più potente acceleratore di

12. Larousserie D, *Isabelle Guyon veut démocratiser L'intelligence atyficielle*, Le Monde Science et Medicine, 11-4-2018 <https://guyon.chalearn.org/projects/humania>

13. Gallino L, *Finanzcapitalismo. La civiltà del denaro in crisi*, Torino, Einaudi, 2011.

14. Boulanin V, Goussac N, Brun H, Richards L, *Responsible military use of artificial intelligence*, SIPRI, November 2020

15. Le Saint R, *Des hôpitaux misent sur des robots chirurgicaux à 1,4 milion d'euros pour liberer des lits*, Mediapart 9-4-2023. <https://www.mediapart.fr/journal/france/090423/des-hopitaux-misent-sur-des-robots-chirurgicaux-14-million-d-euros-pour-liberer-des-lits>

chirurgicaux-14-million-d-euros-pour-liberer-des-lits

16. Circa 2 miliardi di persone non hanno accesso ad acqua potabile e 3,6 miliardi sono senza latrine e canalizzazione di acque usate. Murugesu JA, *Around two milion people don't have access to clean drinking water*, NewScientist, 21-3-2023. <https://www.newscientist.com/article/2365541-around-two-million-people-dont-have-access-to-clean-drinking-water>

www.newscientist.com/article/2365541-around-two-million-people-dont-have-access-to-clean-drinking-water

17. La malaria causa ancora circa 600 milioni di morti annue. *World Malaria Report 2022* <https://www.who.int/publications/i/item/9789240064898>

particelle esistente, che consuma quantità immense di energia, ed egualmente chiaro lo scopo dell'ultimo miracolo di ingegneria astrofisica, il James Webb Telescope, la cui trentennale costruzione e il cui funzionamento sono sommamente energivori, di risalire alle origini dell'universo, meno chiaro è lo scopo di creare macchine di intelligenza artificiale energivore. Energia che è invece necessario economizzare nell'ottica di una sobrietà nell'uso delle risorse sempre più oculato, se vogliamo ancora cercare di non far bollire il nostro pianeta. L'unico che abbiamo, almeno nel presente e nel prevedibile futuro. Ma, guarda caso, coloro che sono dietro alla creazione di queste fantastiche macchin/azioni intelligenti e gridano al lupo dopo averlo allevato arricchendosi, pensano che si possa andare a vivere su altri pianeti. Almeno loro. Lo scorso marzo è stata pubblicata una lettera aperta¹⁸ firmata da un migliaio tra imprenditori, ricercatori informatici, professori e varie celebrità made in USA ambisesso (ma in maggioranza maschi) che perora la causa di una battuta d'arresto (di sei mesi¹⁹) nello sviluppo della IA, adducendo un fondato timore che si possa in breve arrivare alla creazione involontaria di una super-intelligenza, superiore a quella umana, che “uscendo dalla scatola” possa prevaricare e distruggere il genere umano²⁰! Lo scenario è degno di film (americani!) del terrore tipo *Terminator*.

È stupefacente che tale prospettiva possa sembrare un pericolo concreto, mentre innumerevoli altri pericoli e svantaggi ben più immediati e concreti per il genere umano non vengono assolutamente menzionati nella lettera. Roberto

Ciccarelli, su *Il Manifesto*²¹, riporta una intervista al prof. Antonio Casilli che rivela la quantità di forza lavoro umana dietro l'addestramento di questi pappagalli stocastici. Ad esempio, i curiosi che si sono affrettati a bersagliare di domande ChatGPT hanno fornito dati gratis per migliorarlo nella totale assenza di un quadro normativo. Ma c'è di peggio: la “pulitura” di tutto il bagaglio di nozioni divorato dal pappagallo ChatGPT4 in cui era presente la spazzatura raccattata dal web, non filtrata, è stata eseguita in Kenya da manovalanza pagata con pochi spiccioli, tra 1,32 e 2 dollari all'ora, persone che inoltre ne hanno sofferto psicologicamente, come denuncia Billy Perrigo su *Time*²², fornendo dati e fatti. OpenAI ha ingaggiato una società informatica di Nairobi (Sama) che a sua volta ha arruolato decine di lavoratori sottopagati, uno dei quali descrive come “tortura” il compito di visionare descrizioni o scene disgustose per espellerle dal bagaglio dello strumento informatico. Il “miracolo” virtuale gronda sudore e sangue, uno sfruttamento bestiale, come sappiamo già dalle denunce di chi ha lavorato per analoghi fini con Facebook. Il virtuale cresce su una materialità disumana.

Inoltre bisogna considerare i cosiddetti *opportunity costs*, consistenti nell'investire miliardi in queste macchine super-energivore invece che in infrastrutture sociali e in sistemi pubblici di IA volti a favorire e accelerare sempre più l'urgente transizione ecologica, industriale ed energetica. I risvolti negativi della IA sono, oggi e non domani, ben presenti: violazione della privacy, ottundimento dello spirito critico specie nei

18. Open Letter, *Pause Giant AI Experiments: an open letter*, Future of Life Institute, 22 March 2023. <https://futureoflife.org/open-letter/pause-giant-ai-experiments/>

19. Se veramente si vuole un ripensamento profondo, a che servono solo sei mesi?

20. Si legga a questo proposito sul “test della scatola”, ideato da uno dei firmatari della lettera, Eliezer Yudkowsky, <https://towardsdatascience.com/the-ai-box-experiment-18b139899936>

21. Ciccarelli R, *ChatGPT, Antonio Casilli: il lato oscuro dell'algoritmo è la forza lavoro*, *Il Manifesto*, /4/2023.

<https://ilmanifesto.it/chatgpt-il-lato-oscuro-dellalgoritmo-e-la-forza-lavoro-che-addestra-la-macchina-digitale>

22. Perrigo B, *Exclusive: Open AI used Kenyan workers at less than 2\$ per hour to make ChatGPT less toxic*, 18-1-2023. <https://time.com/6247678/openai-chatgpt-kenya-workers/>

giovanissimi, computer-dipendenza, creazione di falsità che appaiono plausibili, concorrenza sleale, distorsioni del mercato del lavoro, sfruttamento che non dice il suo nome (vedi Kenya), automazione eccessiva e nociva, sorveglianza poliziesca e controllo antidemocratico. Sono le attuali direzioni di ricerca nel campo della IA il pericolo maggiore, e lo sfrenato sfruttamento in senso commerciale e militare.

Ma c'è dell'altro. Infatti, tornando alla lettera che richiede la moratoria summenzionata, con stupore e alla fine qualcosa che assomiglia al rac-capriccio, mi si è rivelato un mondo, una pseudo-filosofia di tipo paranoico dei suoi firmatari principali, questa sì pericolosissima dato che tra i protagonisti ci sono potenti tecno-miliardari, come li chiama Andrea Signorelli (su Wired), che maneggiano capitali come spiccioli per il caffè. Questi individui sono *tycoons* tipo Elon Musk o a capo di società di IA, inseriti in Università prestigiose, in organizzazioni delle Nazioni Unite, sono stati consiglieri di responsabili politici, e aderiscono ad una corrente di pensiero che è stata definita *longtermism*, lungotermismo in italiano, sviluppatasi nella Silicon Valley. E probabilmente anche Geoffrey Hinton, canadese, che dopo 10 anni ha appena lasciato Google, è un lungotermista. A sua volta il lungotermismo scaturisce da un movimento sociale targato U.S.A., l'*effective altruism* (altruismo efficace).

Di primo acchito il concetto sembra condivisibile e ovvio, il principio che tutti gli umani hanno eguale valore e che occorre preoccuparsi di salvaguardarne l'integrità e il benessere, ovunque essi siano. Il promotore di tale movimento è stato il filosofo scozzese Will MacAskill²³, che con varie iniziative ha raccolto negli anni molti fondi per scopi filantropici, (filantropia, si torna indietro di 200 anni), egli stesso filantropo. Poi leggiamo che per massimizzare i benefici delle nostre azioni, secondo i dettami dell'altruismo effica-

ce, bisogna cercare di arricchirsi il più possibile, raggiungere posizioni di potere e guadagnare a palate, per poter fare più beneficenza. Magari arricchirsi come l'imprenditore di criptovalute Sam Bankman-Fried, un sostenitore del movimento, fondatore di FTX, incappato in grossi problemi giudiziari. Arricchirsi ad ogni costo per poter fare beneficenza, ben strano principio.

Leggiamo anche che Will MacAskill pensa che per servire meglio l'umanità è bene guardare molto oltre il presente, non nel futuro che incombe ora e nei decenni successivi, ma nei millenni a venire. Sul sito Vox²⁴ si legge che l'altruismo efficace poggia su tre premesse: 1, coloro che nasceranno nel futuro saranno tanti; 2, la loro sorte è importante per noi; 3, è in nostro potere rendere le loro vite felici o infelici. L'aumento esponenziale dei gas serra, gli stravolgimenti climatici attuali e futuri anche peggiori non turbano gli altruisti. Né l'estinzione di migliaia di specie viventi in corso. Chi conta (loro, oggi) non ne sarà presumibilmente toccato. E dato che in questo futuro ipotizzato e lontano, lontanissimo, la specie umana si sarà moltiplicata presumibilmente su vari pianeti di ricambio (il bravo Will ipotizza 80 bilioni di esseri umani, tombola), è giusto preoccuparci ora di salvarli: si legga l'intervista su *The Guardian* del 2 agosto 2022²⁵. È meglio rendere felici bilioni di persone invece che qualche miliardino. Così la pensano i lungotermisti: l'altruismo efficace trasformato in lungotermismo è il credo cui aderiscono molti santoni della Silicon Valley e i vari "scienziati" che firmano la lettera in cui chiedono che la ricerca in IA faccia una pausa perché il pericolo che possa sfociare nella creazione di una super-intelligenza (ASI, *artificial super intelligence*) è reale e la nostra (!) democrazia può essere messa in pericolo. Anche quella dietro a Guantanamo.

Un altro filosofo di riferimento lungotermista è Nick Bostrom, fondatore nel 2005 del *Future of*

23. https://it.wikipedia.org/wiki/William_MacAskill

24. <https://www.vox.com/future-perfect/23658383/longter>

Humanity Institute, parte del dipartimento di Filosofia all'Università di Oxford, che condivide i suoi locali con il "Centro per l'altruismo efficace" con il quale collabora e di cui condivide i principi: fare sì che l'umanità a venire possa beneficiare il più possibile delle scelte che si compiono qui e ora. Ottica lungotermista condivisa da un altro filosofo di Oxford, Toby Ord, che è stato consigliere dell'OMS, della Banca Mondiale e del *World Economic Forum*. Per gli ultra-presbiteri i disastri già provocati dal cambiamento climatico in corso, le epidemie di colera, siccità e inondazioni a ripetizione, l'alta mortalità materno-infantile e simili in paesi poveri, la corsa agli armamenti non sono problemi seri, è una frazione tutto sommato piccola dell'umanità che è coinvolta. Per la crisi climatica, l'opzione migliore è cercare un altro pianeta per emigrare come propugna l'altro noto lungotermista, Elon Musk, "per realizzare la nostra vocazione di stirpe multiplanetaria". Certo che occorreranno parecchi pianeti per tutti quei bilioni di umani.

Anche Jaan Tallin, fondatore di Skype e co-fondatore del *Future of Life Institute*, legato al *Future of Humanity Institute*, pensa che la crisi climatica non rappresenti un rischio esistenziale per una specie destinata agli spazi interplanetari (Andrea Signorelli su *Wired*²⁶). Per Nick Bostrom tragedie come Chernobyl o l'epidemia di AIDS sono "semplici increspature sul grande mare della vita"²⁷. Insomma, il collasso climatico possibile e probabile non è un problema grave per la sopravvivenza dell'umanità, lo spettro di una super-intelligenza tiranna invece lo è. E Nick Beckstead, membro del *Future of Humanity Institute*, arriva a sostenere che "salvare vite umane nelle nazioni povere potrebbe essere meno utile che salvare vite nelle nazioni ricche"²⁸. Si arriva a conclusioni che non sarebbero dispiaciute affatto a un tal Signor Goebbels.

Spaventa pensare che questi "filosofi", amministratori delegati, influencer, questa rete di istituti, fondazioni, società informatiche e i loro iniziatori e aderenti possano allargare la loro influenza e acquistare sempre più potere e magari autorevolezza. Colossi commerciali come Amazon Google Microsoft sono dei vivai pericolosi di simili esemplari umani, e spiace che siano presenti anche in Istituti Universitari a Oxford. O all'ONU. Credo che occorra informare, informarsi, mobilitarsi per far sì che l'IA sia sin da ora al servizio dell'umanità, della ricerca utile ad allontanare e sventare i reali pericoli presenti e prevedibili che incombono su questo pianeta ora e non in un futuro immaginario e fantasmatico. E che occorra smascherare questi profeti paranoici. ●



Fantascienza (un artista immagina astronauti su Marte)...



... e realtà (alluvioni assassine)

27. Signorelli AD, *Cos'è il lungotermismo, la nuova fanatica utopia della Silicon Valley*, 24-3-2023. <https://www.iltascabile.com/scienze/lungotermismo/>

28. *Ibidem*.