

Hans Ruin

*L'intelligenza artificiale e il nucleo perverso dell'era tecnologica**

Nel giugno del 1938 Martin Heidegger fu invitato a tenere un intervento per una serie di conferenze pubbliche dal titolo *La fondazione dell'immagine del mondo della modernità* (*Die Begründung des Weltbildes der Neuzeit*). Il contributo diverrà poi il saggio intitolato *L'epoca dell'immagine del mondo* (*Die Zeit des Weltbildes*) pubblicato nella raccolta *Holzwege* (Heidegger 1950). In questo articolo, Heidegger in primo luogo individua ciò che sarà un motivo ricorrente nei suoi tardi lavori, ovvero la rilevanza filosofica e metafisica della tecnica per la nostra epoca storica corrente. Egli scrive: «La tecnica delle macchine rimane fino ad ora il più visibile prodotto dell'essenza della tecnica moderna» (*Die Maschinentchnik bleibt der bis jetzt sichtbarste Ausläufer des Wesens der neuzeitlichen Technik*). Tale ascesa delle capacità e competenze tecniche non è solamente di enorme rilevanza per il mondo contemporaneo, bensì essa guida e riassume il modo in cui noi vediamo ed esperiamo il mondo. Nell'epoca della tecnica il mondo stesso è trasformato per divenire un'immagine – *a Bild* – in sé e per sé. Pertanto, Heidegger ritorce il titolo dell'evento contro sé stesso. *L'immagine* contemporanea del mondo è un'epoca in cui il *mondo* stesso è ridotto a immagine, in quanto qualcosa di individuato, rappresentato, misurato, e controllato. «Per noi – scrive – il mondo deve divenire immagine» (*die Welt zum Bild werden muss*). Nella conferenza su Kant e il problema della cosa dello stesso periodo, egli interpreta la modernità, sin dall'epoca di Galilei e Descartes, dal punto di vista della matematica e della matematizzazione della natura. Da qui in poi la matematica funge da struttura fondante del pensiero e dell'esperienza del mondo, assunta e sostenuta al suo centro dal soggetto pensante. Perciò, la moderna metafisica della soggettività e la matematizzazione della natura sono correlate sin dall'inizio in un comune tentativo di rendere il mondo comprensibile e di porlo sotto il domino e controllo della soggettività universale.

Quando cerchiamo di pervenire a qualcosa come lo “*Zeitgeist*” o lo spirito dei tempi, possiamo vedere come da queste analisi si possa descrivere una traiettoria sino al presente. Il nostro è un tempo in cui lo spirito umano sta compiendo un ultimo e terminale tentativo di porre sé stesso come artefatto matematizzabile, attraverso la riproduzione della sua stessa essenza sotto forma di tecnologie mac-

* Il saggio è apparso per la prima volta in turco con il titolo *Yapay Zekâ Ve Teknoloji Caginin Sapkin Özü*, in “Saba Ülkesi”, 58 (2019), pp. 44-50. La presente traduzione è di L. De Stefano.

chiniche e grazie all'analisi matematica. La sfida distintiva della nostra epoca – il *Geist* del nostro *Zeit* – è la misura in cui il *Geist* può implementare sé stesso nella sua propria macchinazione e di qui infine abbandonarsi alla sua propria creazione tecnico-matematica.

Nel grandioso progetto culturale della *Artificial intelligence* siamo testimoni della perversa acme del destino della spiritualità moderna, dove il suo desiderio di assoluto dominio, attraverso il ragionamento calcolante e l'ingegneria tecnica, assume la forma di una affermazione della sua propria sottomissione finale al suo Altro da lei stesso creato.

Nel 1936 – due anni prima della conferenza heideggeriana – il matematico britannico Alan Turing, allora ventiquattrenne, pubblica quello che sarebbe divenuto uno dei saggi leggendari del secolo scorso, con l'inquietante titolo *On computable numbers with an application to the Entscheidungsproblem* (Turing 1936). La tesi contenuta era diretta contro il suo collega tedesco David Hilbert e il suo tentativo di trovare un metodo logico per decidere se una data proposizione matematica possa essere vera o falsa. Già alcuni anni prima Kurt Gödel mostrò i limiti del programma logico di Hilbert. Ma con l'articolo di Turing le sue speranze andarono in frantumi. La conclusione di Turing era ugualmente negativa: non esiste una tale procedura o algoritmo. Ma al fine di dimostrare tale assunto, egli inventò – in teoria – quella che chiamò una “macchina universale”, nota successivamente come “*Turing machine*”. Oggi ciò è universalmente riconosciuto come il primo fondamento teoretico della informatica.

Durante la decade successiva, le conseguenze pratiche di tale idea, in un primo momento solamente astratta, subirono un ragguardevole incremento attraverso la stretta collaborazione tra matematici ed ingegneri. Turing fu chiamato alle armi dal Bletchley Park Project che decifrò il codice German Enigma. In una straordinaria battaglia segreta tra ingegneri matematici inglesi e tedeschi, la guerra fu in parte vinta nei laboratori della prima generazione di scienziati informatici, così come fu decisa analogamente nei laboratori segreti dei matematici e fisici che lavoravano da ambo i lati dell'Atlantico per implementare le conseguenze tecniche della fisica teorica nucleare. Con la Seconda Guerra Mondiale il pieno potenziale della matematica emerse, pertanto, dalle appartate stanze della speculazione teoretica in cui era stata coltivata da millenni, per assumere il ruolo di linguaggio del dominio sul mondo. Il ragionamento matematico, la più avanzata espressione della intelligenza umana, attraverso il conseguente tentativo di implementare sé stesso nel macchinico, vorrebbe ora avanzare la pretesa di raccogliere il testimone dalla lunga tradizione della filosofia speculativa: divenire un oggetto in sé e per sé nella forma di macchine pensanti.

Solo una decade più tardi, nel 1948, il matematico ungaro-americano John von Neumann, proclamò che presto il computer avrebbe sicuramente sopravanzato l'intelligenza umana. Inoltre, nel 1959 Turing diede alle stampe un altro articolo, divenuto un classico, in cui affermava che le macchine avrebbero in poco tempo pensato come persone. Al fine di decidere quando tale evenienza si sarebbe data, egli immaginò un “*imitation game*” – più tardi noto come “*Turing test*” – in cui un giudice è in comunicazione con una macchina senza sapere chi è cosa. Quando il

giudice non può più dire la differenza sulla base di criteri comportamentali puramente esteriori, la macchina deve essere considerata intelligente. Fu poi nell'ambito di una leggendaria conferenza al Dartmouth College nel 1956 che l'allievo di Neumann, John McCarthy, coniò il termine "*artificial intelligence*" per caratterizzare il suo nuovo progetto tecnologico. La strada era ormai spianata per un progetto di dimensioni planetarie: gli esseri umani stavano per ricreare la facoltà che nella lunga tradizione di filosofica e teologica gli era stata attribuita come la loro caratteristica più preziosa e unica, il loro autentico vincolo con Dio: l'intelletto.

La prospettiva di tale possibilità rende attuale anche la domanda: che significherà questo per il futuro della umanità? Le risposte letterarie ed estetiche si materializzarono ben presto. In effetti, queste sono già in giro da un po'. Nella mitologia ebraica esiste il Golem, la creatura artificiale che condensava le ansie di una creazione contenente sia sottomissione che rivolta. Nella narrativa moderna fu il Frankenstein di Mary Shelley del 1818 che per primo catturò i timori della nascente epoca tecnologica sotto forma della prospettiva della vita artificiale, raggiunta attraverso l'uso della recente scoperta della forza della elettricità. Lo scrittore vittoriano Samuel Butler, nel suo libro del 1872 *Erewhon*, descrive come una nuova razza di macchine, in pieno spirito darwiniano, riesca a raggiungere la coscienza e a dominare il mondo. Ma fu lo scrittore ceco Karol Capek che nella sua opera "R.U.R." del 1920 coniò il concetto di "robot" (dalla parola slava per indicare il servo della gleba) come nome per una creatura macchinica che si ribella contro il suo padrone. Gli straordinari passi avanti della tecnologia dagli anni '40 in poi diedero nuovo impulso a tali fantasie. In questo contesto, un interessante individuo è John Good, che fu collega di Turing nel progetto Bletchley Park. Nel 1965 pubblicò un articolo sul *The New Scientist* (Good 1965) dove immaginava l'emergenza di una "*ultraintelligent machine*" in grado di generare nuove macchine. In ogni nuova generazione di macchine, essa riusciva ad accrescere la sua capacità, il che avrebbe in poco tempo condotto a una "esplosione di intelligenza", un evento che von Neumann aveva già definito in alcune conversazioni "singolarità". Quando Stanley Kubric e Arthur Clarke decisero di creare assieme il film *A Space Odyssey* nel 1968, furono talmente scrupolosi da richiedere dei pareri scientifici allo stesso Good per il loro progetto cinematografico. Il supercomputer Hal9000, protagonista del film e poi del libro, da questo momento in poi avrebbe rappresentato l'immagine paradigmatica di come una mente artificiale, nel momento decisivo, possa rifiutare di obbedire al suo artefice e ribellarsi contro di lui.

In questo tempo di straordinarie conquiste tecnico-matematiche e di visioni utopistico-apocalittiche di un futuro in cui l'intelletto umano si è trasformato in una macchina superumana, vi furono anche dei contro-argomenti filosofici. Heidegger ha continuato a riflettere sull'impatto della tecnica sulla vita umana e sulla conoscenza. Quindici anni dopo la conferenza sull'immagine del mondo, il filosofo di Meßkirch tenne una lezione nel 1953 che sarebbe diventata uno dei suoi scritti più citati: *La questione della tecnica* (*Die Frage nach der Technik*) (Heidegger 2000). In tal contesto, egli riprende e approfondisce la sua tesi che vede nella tecnica non qualcosa di semplicemente tecnico, ma un modo di esperire ed organizzare il mondo, e attraverso cui è trasformato in un artefatto manipolabile. Questa

matrice tecnico-metafisica – a cui da il nome di *Ge-stell* – tenderà essenzialmente a compenetrare l'autocomprensione umana e le conquiste dello spirito, compreso il linguaggio. Se ci atteniamo alla idea tradizionale secondo cui la tecnica è solo un oggetto o un mezzo per uno scopo, e che bisogna solo saperla adoperare in maniera propria, falliamo nel comprenderne il suo pieno significato. La tecnica non è semplicemente un attrezzo, ma un modo di disvelamento del mondo e del suo venire alla presenza. Essa presenta il mondo nel modo dell'utilizzabile, del regolabile e controllabile. Nelle sue conversazioni con lo psicoanalista Medhard Boss tra i tardi anni '50 e i primi '60, Heidegger tornò su tali questioni. In questo frangente, egli fa esplicito riferimento al lavoro del matematico di Harvard Norbert Wiener e la sua "cibernetica" (Wiener 1950), lo studio del controllo e del governo nella natura e nella tecnologia. Un caso esemplare era costituito dal linguaggio, che, dal punto di vista della cibernetica, era considerato come artefatto complesso ma computabile.

Attraverso le note critiche alla cibernetica di Wiener, si può dire che Heidegger abbia avuto a che fare indirettamente con la nascente disciplina dell'informatica e i primi stadi della intelligenza artificiale, nonostante lo stesso Wiener sia stato solo marginalmente coinvolto con quanto al tempo i suoi colleghi ad Harvard e al MIT stavano realizzando. L'impostazione generale di Heidegger, tuttavia, risulta essere una ispirazione fondamentale per quello che diventerà il primo fondamentale confronto critico con l'intelligenza artificiale: *What Computers Cant'Do* di Hubert Dreyfus del 1972. Avendo studiato al MIT negli anni '70, Dreyfus, addestrato alla scuola fenomenologica, era personalmente a contatto con ciò che i pionieri della intelligenza artificiale, riuniti intorno a Marvin Minsky, stavano realizzando, e con le loro aspettative iperboliche connesse ai risultati del loro lavoro. Basandosi in parte sui primi lavori heideggeriani, Dreyfus sosteneva che l'intero programma di ricerca sulla intelligenza artificiale, così come era allora praticata, si basava su una concezione riduttivistica della intelligenza come equivalente alla gestione astratta di simboli, disconnessa dal corpo e dal mondo della vita. A causa di queste intrinseche limitazioni teoriche, egli – in modo corretto – anticipava i limiti tecnici della AI nel contesto operativo della generazione dei fondatori della matematica e della informatica.

Tuttavia, non fu Dreyfus che sarebbe diventato famoso per aver formulato una critica filosofica alla AI, ma il suo collega alla Berkley, John Searle. In un articolo del 1980 egli formulò un esperimento mentale chiamato "The Chinese Room" (Searle 1980). Il suo scopo era dimostrare che persino quando si mima il comportamento umano intelligente, la prestazione di una macchina cinese parlante dipende da un ordinamento casuale meccanico di simboli che, limitatamente alla macchina, non hanno alcun senso. Sfidando esplicitamente l'approccio comportamentista di Turing, Searle insisteva che non vi era alcun bisogno di ascrivere le capacità umane alla macchina, eccetto che in una accezione metaforica. Un computer può operare mosse negli scacchi da esperto, correggere interpretazioni o fornire delle risposte sensate in una conversazione, ma alla fine *non c'è nessuno lì* che è coinvolto in alcuna di queste attività. La critica di Searle mirava alle condizioni del comportamentismo secondo cui è possibile ascrivere a qualcuno l'intelligenza, ribattendo che c'è un modo di essere-per-sé, una autocoscienza, che è parte irriducibile di

ciò che vuol dire comprendere veramente qualcosa, e quindi coinvolta in qualsiasi cosa sia simile ad una intelligenza.

Per la maggior parte dei pratici ingegneri, tuttavia, queste questioni non erano così importanti. Quando ne va del compito di costruire una macchina che è in grado di fare questo o quello, tradurre, fornire diagnosi mediche, guidare un veicolo o giocare a scacchi, è in ultima istanza una mera questione di prestazioni adeguate, piuttosto che se essa sia o meno “intelligente” in senso filosofico. Ciononostante, le critiche filosofiche hanno pungolato molti coinvolti nel *business* della AI, non ultimi alcuni dei suoi più accesi precursori e imprenditori. Le loro reazioni rivelarono quanto importante fosse per loro e per i loro discepoli affermare di non essere semplicemente ingegneri che assolvono a un compito, ma che erano sulla via di rivelare il funzionamento dell'intelligenza umana e il mistero della mente umana. In alcuni casi, la motivazione era di matrice filosofico-intellettuale, ma aveva ovvie ragioni economiche. Affermare che la AI ha una rilevanza filosofica per l'autocomprensione umana o che tale ricerca stia per portare una nuova forma di vita post-umana basata sul linguaggio Chisel, crea attenzione mediatica e genera fondi per la ricerca. Il fatto che Ray Kurzweil, capo della ricerca e sviluppo di Google, nel suo molto apprezzato lavoro del 2004 – *The Singularity is Near. When Humans Transcend Biology* – spenda diverse pagine nel tentativo di confutare l'argomentazione di Searle ormai vecchia di un quarto di secolo, rivela la posta in gioco (Kurzweil 2004).

L'impeto nella difesa della AI contro le critiche filosofiche ha che fare con come essa si è evoluta dagli anni 80 in poi. Notando che le prime previsioni relative al suo successo non si stavano materializzando, si iniziò a sperimentare con nuovi modelli per incrementare le prestazioni delle macchine, in particolar modo le cosiddette “reti neurali” erano viste come modello più verosimile del funzionamento effettivo del cervello. Ma anche quando la capacità di elaborazione delle macchine aumentò rapidamente, le grandi aspettative riposte in ciò che avrebbero potuto realizzare andarono deluse. Retrospectivamente, gli anni '80 e '90 ravvisarono una lunga battuta di arresto nella ricerca sulla AI, al punto che il concetto stesso acquisì una accezione negativa. Questa storia è stata ben raccontata in una recente collettanea pubblicata dalla rivista *New Scientist*, con il titolo *Machines that think. Everything you need to know about the coming age of artificial intelligence* (Heaven 2017). Quando Deep Blue della IBM sconfisse Kasparov nel 1997, fu un grande shock nel mondo degli scacchi e ciò contribuì a focalizzare l'attenzione nuovamente su una ricerca che aveva fallito nel mantenere le iperboliche promesse iniziali. Ma nessuno, nemmeno i suoi programmatori, avrebbero potuto affermare che Deep Blue fosse intelligente in termini generali, nemmeno come giocatore di scacchi. Mancava di strategia, non poteva imparare o fare esperienza dalle sue mosse, ed era inadoperabile per qualsiasi altra attività eccetto il gioco degli scacchi.

Oggi abbiamo una situazione molto differente. Oggi la AI è di nuovo sulla bocca di tutti. Ha i suoi profeti, i suoi evangelisti e le sue cassandre. Le rappresentazioni estetiche di robot intelligenti, siano essi minacciosi o premurosi, popolano di nuovo la letteratura e il cinema. Le fantasie di von Neumann e Goods del 1950 sulla “singolarità” e sulla “esplosione di intelligenza” sono state nuovamente ri-

lanciate da ricercatori pagati somme esorbitanti per le loro lezioni a un pubblico entusiasta. E l'incentivo commerciale è nuovamente di tendenza. Cosa è successo? La spiegazione contiene un interessante risvolto filosofico. Per il tempo in cui la AI si è concepita come il tentativo di ricreare e spiegare l'intelligenza umana, è rimasta fondamentalmente impantanata. La vera ascesa si è avuta quando gli ingegneri rinunciarono all'idea della creazione di modelli rispecchianti come si credeva la mente umana operasse, e si focalizzarono sul creare macchine che potevano auto-educarsi attraverso algoritmi meno complessi capaci di processare statisticamente una enorme quantità sempre crescente di dati.

Ciò può essere illustrato attraverso il successo dei programmi di traduzione. Essi erano promessi sin dall'inizio, ma nonostante i tentativi di sviluppare algoritmi capaci di mimare le abilità di un traduttore umano, essi non raggiunsero mai delle prestazioni commercialmente accettabili. La svolta vera si ebbe solo nella prima decade dei duemila, attraverso la creazione di programmi capaci di generare una verosimiglianza calcolata statisticamente attraverso una base di text-data sempre crescenti, tale che a una data proposizione nella lingua di destinazione corrispondesse una proposizione dalla lingua di partenza. Lo stesso modello è utilizzato per i programmi di riconoscimento facciale, per i veicoli automatici e per strumenti destinati al mercato finanziario. Questi programmi testimoniano la straordinaria intelligenza dei loro creatori, ma nessuno inferirebbe che comprendano le lingue che stanno processando, e ancor meno che essi siano intelligenti.

Allo stesso tempo, siamo quotidianamente nutriti da scenari drammatici – descritti da filosofi, fisici e informatici – riguardanti come i robot intelligenti stiano per prendere il controllo del mondo in un futuro prossimo. Tra gli scrittori più conosciuti nel genere troviamo il filosofo Nick Bostrom e il fisico Max Tegmark. I loro *best sellers* *Superintelligence. Paths, Dangers, Strategies* (Bostrom 2014) e *Life 3.0. Being Human in the Age of Artificial Intelligence* (Tegmark 2017) affermano entrambi che, visto lo stato dell'arte delle machine auto-apprendenti, dovremmo aspettarci in un futuro non troppo lontano un confronto con la “superintelligenza” o una “esplosione di intelligenza”. A un imprenditore come Ray Kurzweil di Google ciò appare un prospetto interessante, laddove per Bostrom e Tegmark è uno scenario pericoloso che richiede delle contromosse sia pratiche che teoriche, per cui stanno efficientemente raccogliendo fondi di ricerca.

Nonostante il fine lodevole ossia il tentativo di porre la questione circa le possibili conseguenze negative della AI, il problema con tali scenari apocalittici consiste nella stessa concezione limitata di cosa innanzitutto sia l'intelligenza. Con il riferimento alla complessità biologica del cervello umano, con le sue strutture neuronali e sinapsi, essi affermano che è solo una questione di tempo per il raggiungimento di una complessità paragonabile all'intelligenza umana, tradotta in potere computazionale, da parte dei computer. Ma al di là di questa comparazione puramente qualitativa al cervello inteso come dispositivo computazionale, non vi è alcuna *ratio* dietro all'idea che in qualche modo gli attuali programmi altamente specializzati, risultanti da un qualche tipo di logica evuzionistica, dovrebbero improvvisamente dare vita a un supercervello con le sue proprie motivazioni interne, organizzative e politiche. Ciò è uno scenario alta-

mente improbabile se ci dobbiamo basare sullo stato delle tecnologie esistenti e sui principi su cui è sviluppata la loro capacità. È una delle tante “possibilità esoteriche”, per citare il libro del *New Scientist* summenzionato, in cui gli autori affermavano, inoltre, che l'intero ambito delle profezie sulla AI necessiterebbe di un serio *reality check*. Ancora, le fantasie hanno la loro intima dinamica, un impulso estetico e teologico, che spiega la pervasività e impatto del loro messaggio. Quando ci confrontiamo oggi sulle possibili conseguenze di scenari che hanno la loro origine nelle fantasie apocalittiche del dopoguerra, corriamo il serio rischio di perdere di vista questioni più urgenti connesse allo sviluppo tecnologico attuale della AI. Una di queste è il vero impatto ed effetto dei robots, che non solo sostituiscono il semplice lavoro materiale, ma gradualmente anche compiti analitici più avanzati come le diagnosi mediche, i verdetti finanziari e giuridici, compiti di insegnamento etc. Corriamo, inoltre, il rischio di non confrontarci in maniera chiara con le conseguenze di come la AI sia usata per il controllo sempre più efficiente della popolazione. La Cina, una dittatura che utilizza indiscriminatamente queste tecniche per controllo economico e politico, è oggi pioniere a livello mondiale delle ricerche sulla AI, e il suo fine dichiarato è di divenire il principale attore mondiale in un prossimo futuro.

In questa situazione, gli intellettuali che dibattono sulla base della fantascienza sul pericolo della superintelligenza, rischiano di diluire la capacità critica di pensare cosa è veramente in gioco con la AI. Essi esasperano gli scenari distogliendo l'attenzione dalle domande più urgenti e dalle vere sfide, in una situazione dove i sistemi di intelligenza artificiale sono sempre più adoperati per scopi di controllo e razionalizzazione economica neutrali dal punto di vista valoriale. Invece di aiutarci a capire cosa sta succedendo e a riflettere su questa nuova realtà, gli scenari apocalittici mostrano – involontariamente – il nucleo perverso di una cultura tecnico-matematica, che sin dal suo principio è guidata da un desiderio di dominio totale, oscillando ambigualmente tra piacere e ansia al pensiero di una sottomissione al suo stesso potere. In una strana fusione di affetti contraddittori, la cultura tecnologica genera a un tempo i sogni e gli incubi del suo stesso desiderio. La tecnologia resiste alla speranza di un mondo completamente dominato, mentre allo stesso tempo forgia l'immagine di un grande Altro che alla fine dei giorni ci ridurrà a strumenti della sua stessa volontà.

La stessa cecità intellettuale è rintracciabile tra i sociologi e i letterati che si imposero come critici di una prima generazione presumibilmente ingenua tecnologicamente di umanisti. In un testo del 1996 *Farben und/oder Maschinen Denken*, Friedrich Kittler, il più noto esponente del “medial turn” nelle scienze umane, ridicolizza Heidegger e gli altri che avrebbero fallito nel vedere che gli studi umanistici stavano per essere usurpati dalla teoria dei media e della informazione (Kittler 1996). Nella visione di Kittler, la aspirazione nel tener separati uomo e macchina è sintomatica di un pregiudizio umanista antiquato. Ma nel piegarsi di fronte al Grande Altro della tecnologia della informazione, egli manifesta ed anticipa una infelice tendenza comune agli studi umanistici attuali, in quanto essi si sentono costretti a scegliere tra affermazione e rifiuto di fronte a una macchina complessa che non comprendono.

Verso la fine del saggio sulla tecnica, Heidegger evoca come la tecnica non dovrebbe solo esser vista come potenza dominante, e che copre i legami più essenziali con l'ambiente vivo. Il saggio ci invita a pensare l'impensato in sé stesso. La crescente domanda di tecnologia nelle nostre vite è anche percepita come salvaguardia della possibilità di pensare la latente libertà inscritta in questo destino obbligante. Attraverso questa analisi, formulata in un tempo in cui la teoria dell'informazione stava appena iniziando ad essere implementata nelle macchine, Heidegger indica un modo della riflessione in grado di pensare l'intimo desiderio della razionalità come guidato da una "inesorabilità dell'ordinare" (*Das Unaufhaltsame des Bestellens*). Con questo ci invita così a riflettere come la tecnologia, nel suo interno dispiegarsi attraverso la potenza, può anche dare i natali a un pensiero liberante. Invece di tenerci imprigionati nella vacillante fascinazione della affermazione e del rifiuto dinanzi alla visione di un artefatto latore di un latente dominio assoluto, possiamo lasciare che la AI nel suo stato attuale e le sue implicazioni culturali ci invitino rimeditare il nucleo perverso della nostra era della tecnica.

Bibliografia

Bostrom, N.

2014 *Superintelligence. Paths, Dangers, Strategies* Oxford University Press, Oxford, tr. it S. Frediani, Bollati Boringhieri Torino, 2018.

Dreyfus, H.

1972 *What Computers Can't Do*, MIT Press, New York.

Good, I. J.

1965 *The estimation of probabilities: An essay on modern Bayesian methods*, in *Research monograph no. 30*, MIT Press, Cambridge.

Heaven, D. (a cura di) 2017 *Machines That Think: Everything You Need to Know about the Coming Age of Artificial Intelligence*, in *New Scientist Instant Expert*, Nicholas Brealey Publishing, New York

Heidegger, M.

1950 *Holzwege*, Klostermann, Frankfurt a.M.; tr. it. a cura di Pietro Chiodi, La Nuova Italia, Firenze, 1999.

2000 *Vorträge und Aufsätze*, hrsg. von Friedrich-Wilhelm v. Herrmann, Klostermann, Frankfurt a.M.; tr. it. a cura di G. Vattimo, Mursia Editore, Milano, 2014.

Kittler, F.

1996 *Farben und/oder Maschinen denken in Synthetische Welten. Kunst, Künstlichkeit und Kommunikationsmedien*, Verlag Die Blaue Eule, Essen; tr. ingl. *Thinking Colours and/or Machines in Theory, Culture & Society* 23(7–8) SAGE, London, 2006.

Kurzweil, R.

2017 *The Singularity is Near. When Humans Transcend Biology*, Viking Books, New York, tr. it. Apogeo Education, Napoli, 2008.

Searle, J.R.

1980 Searle, John. R. *Minds, brains, and programs*. In *Behavioral and Brain Sciences* 3, pp.417-457, Cambridge University Press, Cambridge.

Tegmark, M.

2017 *Life 3.0. Being Human in the Age of Artificial Intelligence*, Knopf, New York; tr. it. V.B. Sala, Raffaello Cortina Editore, Milano, 2018.

Turing, A.

1936 *On computable numbers with an application to the Entscheidungsproblem*, Proceedings of the London Mathematical Society 2, London.

Wiener, N.

1950 *The Human Use of Human Beings*, Houghton Mifflin Co., Borston; tr. it. D. Persiani, Bollati Boringhieri, Torino 1966.

