



Dipartimento
di Impresa e Management

Cattedra di Informatica

Intelligenza Artificiale: Rischi e Benefici di una Rivoluzione già in Atto

Prof. Laura Luigi

RELATORE

Gabriele Giordano

Matr. 216791

CANDIDATO

Anno Accademico 2019/2020

*A mio padre e a mia madre che mi hanno sempre sostenuto,
a fratelli ed amici che mi hanno sempre sopportato,
a chi è rimasto e a chi invece è andato via, perché
se in questo giorno mi trovo a celebrare il conseguimento
di un obiettivo così importante è anche grazie a loro.*

Sommario

Introduzione.....	5
PARTE I - PASSATO	6
1.0 L'intelligenza tra ermeneutica e semantica	6
1.1 Breve cronistoria dello sviluppo dell'AI	10
1.1.1 L'inizio di una nuova era	10
1.1.2 Presupposti e condizioni per la nascita dell'AI.....	13
1.1.3 Il ventennio onirico dell'AI (1950-1969).....	17
1.1.4 "Primo grande inverno" e rinascita dell'AI (1969-1986)	19
1.1.5 "Secondo grande inverno" (1987- 1993)	24
1.2 L'AI come forma di intelligenza: ipotesi o realtà?	26
1.2.1 The Imitation Game	26
1.2.2 Stanza cinese.....	28
1.2.3 Valutazioni	29
PARTE II – PRESENTE	31
2.1 L'Intelligenza Artificiale al giorno d'oggi	31
2.1.1 Dall'analisi microeconomica.....	31
2.1.2 ... all'indagine macroeconomica	33
2.2 Il valore strategico dell'Intelligenza Artificiale	36
2.2.1 Primo caso di studio: Amazon.com.....	37
2.2.2 Secondo caso di studio: Tesla	40
2.3 Aspetti psicologici e culturali: il Rapporto BVA DOXA	44
PARTE III - FUTURO.....	54
3.1 Lavoro e ricchezza nell'era dell'AI.....	54
3.2 Principali rischi nello sviluppo dell'AI	59
3.2.1 Antagonismo e cooperazione tra l'uomo e l'AI.....	59
3.2.2 Machine Learning, Big Data e Privacy	63
3.2.3 Vigilanza e controllo dell'Intelligenza Artificiale	66
Conclusioni.....	69
Bibliografia	71

Introduzione

Con il presente elaborato di tesi si sono voluti analizzare e valutare i radicali cambiamenti generati dall'Intelligenza Artificiale in ambito lavorativo e nella produzione di ricchezza.

Discutere meramente delle potenzialità economiche collegate all'*Artificial Intelligence* senza considerare al contempo i pericoli insiti nello sviluppo dell'AI e le influenze psicologiche e culturali per chi vi interagisce, però, potrebbe risultare eccessivamente fuorviante: un'analisi così condotta offrirebbe infatti una visione ed una comprensione limitata del fenomeno, impedendo di coglierne la reale e più ampia portata.

Così, se nella prima parte dell'elaborato è stata indagata la nascita e la storia dell'Intelligenza Artificiale, arrivando a chiederci se i software ed i sistemi realizzati possano essere effettivamente definiti “intelligenti”, nella seconda parte del lavoro si è voluto analizzare il livello di sviluppo raggiunto dall'*Artificial Intelligence* al giorno d'oggi: esaminando inizialmente l'impatto generale registrato nei differenti settori industriali e gli sforzi compiuti dalle diverse Nazioni nello sviluppo dell'AI, abbiamo focalizzato la nostra attenzione sul modo in cui Amazon e Tesla utilizzano l'Intelligenza Artificiale all'interno delle rispettive società, dimostrando come questa offra ad entrambe le aziende un importante vantaggio competitivo sui propri concorrenti.

Si è considerato, poi, l'effetto determinato dall'AI con riguardo agli aspetti psicologici e culturali connessi al suo attuarsi. Indicativa, al riguardo, è stata la valutazione di quanto emerso dal rapporto BVA DOXA del 2019, che ha evidenziato l'esistenza di un differente *sentiment* suscitato dall'intelligenza artificiale.

Infine, a conclusione dell'elaborato, è stato esaminato l'impatto che si stima possa avere in futuro l'Intelligenza Artificiale su lavoro e ricchezza, mettendo in evidenza - al contempo - i principali rischi emersi dal suo sviluppo e l'urgente bisogno di risolvere le connesse criticità.

PARTE I - PASSATO

1.0 L'intelligenza tra ermeneutica e semantica

Cogito ergo sum, scriveva Cartesio nel suo “Discorso sul metodo”.

Rifacendosi alla filosofia di Agostino ma, al contempo, capovolgendone radicalmente le prospettive, Cartesio, attraverso il suo scetticismo metodologico, propone un interessante metodo d'indagine di sé e del mondo: secondo Cartesio, infatti, è solo la condizione del dubbio che può portare l'uomo a dedurre l'essere, arrivando così a discernere il falso dal vero.

Come brillantemente nota il filosofo francese, l'atto stesso di dubitare il proprio pensiero altro non è che manifestazione stessa del pensiero; anche se l'uomo non dovesse avere conoscenza precisa della propria origine, anche se non dovesse avere consapevolezza del mondo che lo circonda, avrebbe comunque certezza di sé: l'atto di supporre che io possa ingannarmi coincide infatti con l'io che verrebbe ingannato, svelando così la perfetta coincidenza che esiste tra conoscente e conosciuto. Ed è proprio l'identità di pensante e pensato che, secondo Cartesio, conduce l'uomo all'autocoscienza; se prima con il dubbio scettico si metteva in discussione che all'idea potesse corrispondere la realtà, adesso questo dubbio non ha più motivo di esistere. In breve, sarebbe quindi lo stesso *cogito* a definire ciò che l'uomo è, distinguendolo da tutto ciò che non è *res cogitans*, risultando al contempo la chiave per giungere alla consapevolezza di sé.

Ma perché vi sto annoiando con questi vecchi e soporiferi pensieri di un filosofo vissuto più di 400 anni fa? Perché, a parer mio, sarà necessario riflettere sul significato profondo di parole quali coscienza e consapevolezza per comprendere al meglio la natura della tecnologia che sarà disponibile già nel prossimo futuro. Il domani ci metterà di fronte a spinose questioni, dubbi etici; molte saranno le sfide che saremo chiamati a superare, molti saranno i dilemmi morali che emergeranno e che dovremo sciogliere: per farlo, avremo bisogno di mettere in discussione tutto ciò che fino ad ora abbiamo passivamente accettato come verità inconfutabili, dimenticando arcaici dogmi e credenze ancestrali.

Come suggerisce Cartesio, dobbiamo abbracciare il dubbio.

Forse, la prima cosa in assoluto di cui dovremmo dubitare è il linguaggio stesso: parole e concetti che fino ad oggi hanno riempito i nostri discorsi e che tanto abbiamo trovato stampigliate su ogni sorta di rivista o giornale, di colpo potrebbero svuotarsi di ogni senso e

significato originario, travolte dall'avvento inaspettato di tecnologie futuristiche e dalle loro fantascientifiche applicazioni.

Dopotutto non è niente di eclatante, si tratta solamente di evoluzione: distruggere il vecchio per far spazio a un nuovo più adatto, più aggiornato. È così per tutto e lo stesso linguaggio non è esente da questo incessante processo di rinnovamento.

Prendiamo ad esempio la definizione di intelligenza data dal dizionario Garzanti online:

*“la facoltà, propria della mente umana, di intendere, pensare, elaborare giudizi e soluzioni in base ai dati dell'esperienza anche solo intellettuale; intelletto, ingegno | qualità di chi ha particolari doti intellettuali: mostrare intelligenza; dare prova d'intelligenza | l'insieme dei processi mentali specificamente umani, che si compiono utilizzando simboli (linguistici, logici ecc.), immagini e concetti”*¹

O ancora, scorrendo leggermente la pagina verso il basso, troveremo:

*“persona dotata di capacità intellettuali: tra quei giovani ci sono delle belle intelligenze | nella teologia cattolica: l'intelligenza prima, suprema, Dio; le intelligenze angeliche, celesti, gli angeli”*²

È evidente come i sofisticati processi e le diverse facoltà a cui (consapevolmente o inconsapevolmente) ci riferiamo ogni volta che pronunciamo la parola “intelligenza” siano considerati attributi propri ed esclusivi dell'essere umano; in accordo con la convinzione comune, l'insieme di queste straordinarie capacità sarebbe tanto tipica per l'uomo quanto atipica per altri esseri viventi, come animali e piante (e *a fortiori* per gli oggetti inanimati).

Sono certo che se in passato qualche scienziato un po' visionario avesse provato a mettere in discussione una simile definizione sarebbe stato facilmente etichettato come malato di mente, legato con una camicia di forza e rinchiuso in qualche ospedale psichiatrico fino alla fine dei suoi giorni, dimenticato dal mondo e condannato alla pubblica esecuzione da parte della comunità scientifica.

Oggi però non è più così: grazie alle numerose ricerche intraprese e alle scoperte che si sono susseguite nel corso del secolo scorso abbiamo lentamente iniziato ad ammettere la possibilità

¹ Garzanti Linguistica. (2020). *Intelligenza*. <https://www.garzantilinguistica.it/ricerca/?q=intelligenza>

² Le parti sottolineate sono state così evidenziate a sostegno della tesi che si porta avanti con quanto scritto a seguire

dell'esistenza di un'altra forma di intelligenza, in un certo senso diversa da quella tipicamente umana (ma ugualmente affascinante) che invece sarebbe propria delle macchine; l'opinione diffusa per cui l'uomo sarebbe una creatura perfetta ed inimitabile, probabile retaggio di una visione antropocentrica tipica del Medioevo e, ancor di più, dell'Umanesimo e Rinascimento, si è lentamente sgretolata, rivelandosi del tutto anacronistica ed errata, come dimostrato dai recenti sviluppi tecnologici.

È del tutto evidente a qualunque nostro contemporaneo che la definizione sopra riportata sopravvaluti eccessivamente l'uomo: dopotutto, come dimostrato da Samsung agli inizi di quest'anno³ attraverso la presentazione del progetto Neon, viviamo in un'epoca in cui il confine che divide uomo e macchina si è così assottigliato che non siamo nemmeno più in grado di distinguere l'uomo dalle sue artificiose rappresentazioni.

In un contesto del genere è per noi impensabile riconoscerci quali unici depositari di capacità logiche e di ragionamento.

Ma come siamo arrivati a questo punto?

Per poter rispondere a questa domanda sarà necessario partire da un'altra definizione:

*“Disciplina che studia se e in che modo si possano riprodurre i processi mentali più complessi mediante l'uso di un computer. Tale ricerca si sviluppa secondo due percorsi complementari: da un lato l'i. artificiale cerca di avvicinare il funzionamento dei computer alle capacità dell'intelligenza umana, dall'altro usa le simulazioni informatiche per fare ipotesi sui meccanismi utilizzati dalla mente umana.”*⁴

Si tratta della definizione fornita dall'Enciclopedia Treccani per la voce *intelligenza artificiale*, un'espressione ossimorica frutto della combinazione di parole che, fino a un secolo fa, avrebbe provocato ilarità in chi l'avesse ascoltata. Quella che a prima vista può sembrare una pedante diatriba linguistica nasconde al suo interno il seme di un vero e proprio stravolgimento epocale: per la prima volta nella storia dell'umanità infatti si scinde ufficialmente l'intelligenza da quella componente umana da sempre considerata presupposto necessario, emancipandola dalla prigione in cui per millenni si è vista rinchiusa e ammettendo implicitamente che possa

³ Ansa.it. (2020). *Samsung lancia Neon, replica digitale di un essere umano*. https://www.ansa.it/sito/notizie/tecnologia/hitech/2020/01/07/samsung-lancia-neon-avatar-sudisplay_92e7602b-edeb-4364-b0e5-7151df226249.html

⁴ Enciclopedia Treccani. (2020). *Intelligenza artificiale*. <http://www.treccani.it/enciclopedia/intelligenza-artificiale>

albergare anche altrove. Una definizione del genere è preludio, a parer mio, di una rivoluzione copernicana dell'intelletto: come gli studi dell'astronomo polacco portarono al capovolgimento delle funzioni attribuite a Terra e Sole, così le ricerche future metteranno in dubbio quelle assunte da uomo e intelligenza. E in questo nostro domandarci chi ruoti intorno a chi, sarà centrale la figura e il ruolo che le macchine ricopriranno.

Il progresso avanza, a gran velocità, silenzioso e inesorabile: il futuro è già dietro l'angolo.

Dopotutto viviamo nell'avanguardista terzo millennio, dove i frigoriferi ordinano la spesa in totale autonomia, gli smartphone sono i nostri segretari personali e Alexa, che non è un'affettuosa domestica in carne ed ossa ma uno zelante assistente personale intelligente, supervisiona la nostra abitazione.

Permettetemi quindi di dirvi: benvenuti nel futuro

1.1 Breve cronistoria dello sviluppo dell'AI

1.1.1 L'inizio di una nuova era

Anche se la nascita effettiva dell'intelligenza artificiale è avvenuta solo in tempi recenti, il suo sogno ha origini molto più lontane: il desiderio dell'uomo di realizzare un soggetto in grado di emulare i propri comportamenti è un qualcosa che lo tormenta fin dalla notte dei tempi. Si tratta quasi certamente di una chimera che lo ha accompagnato in ogni luogo e periodo storico, continuamente e ovunque: dopotutto basterà ricordare che la letteratura, i miti e le leggende, da sempre considerate memoria e volontà di interi popoli, sono sovrappopolate da misteriosi esseri caratterizzati da atteggiamenti e pensieri tipicamente umani. A volte questi esseri sono stati descritti come creature antropomorfe, dall'aspetto in tutto e per tutto simile al nostro; altre volte, invece, si trattava di entità incorporee, più simili a ombre inconsistenti che a individui in carne ed ossa, contraddistinte però da capacità di pensiero e linguaggio prettamente umani. Indipendentemente da come possano apparire, ci sono stati presentati sia come mostri abominevoli, araldi di morte e distruzione, sia come spiriti benigni, portatori di prosperità e armonia.

Se poi volessimo soffermarci a discorrere dei loro padri-padroni, non basterebbe nemmeno l'intera ampiezza di questo elaborato per sviscerare adeguatamente l'argomento: autori e narratori si sono a dir poco sbizzarriti sotto questo punto di vista, spaziando da sciamani e alchimisti, scienziati pazzi, logge segrete, fanatici deliranti e sette religiose, fino ad arrivare a misticheggianti dei ed eroi.

Si potrebbe pensare che il desiderio che abbiamo sempre avuto di replicare noi stessi e, più in generale, l'essenza stessa dell'uomo, sia andato scemando, affievolito e fiaccato dai continui insuccessi che si sono susseguiti uno dopo l'altro; ma non è così.

Questo tarlo non ha fatto altro che sussurrare senza sosta nell'orecchio dell'uomo per secoli; con il passare degli anni quel sogno si è trasformato in ossessione, in una tormentosa ricerca senza fine.

Dopotutto si sa, l'uomo è una creatura così controversa: tanto meno una cosa è per lui raggiungibile e tanto più la desidera, disperandosi e struggendosi pur di poterla possedere.

La nascita della cinematografia non ha fatto altro che esasperare questo stato di frenesia, rinvigorendo quel vagheggiamento che per così tanto tempo ha abitato le nostre menti.

Come potevamo rimanere indifferenti nel vedere quelle fantasie millenarie prendere forma proprio sotto i nostri occhi, così vicino da avere l'impressione di poterle sfiorare con mano? Quelle pellicole hanno rappresentato per noi la promessa che, prima o poi, saremmo riusciti a realizzare quel sogno. L'anatra digeritrice⁵, la macchina di Anticitera⁶ e il Turco⁷, così come tutti gli altri automi che sono stati costruiti da chierici, alchimisti e ciarlatani nel corso dei secoli, non sembravano più solamente dei complessi quanto inutili giocattoli; apparivano adesso come il gattonare incerto di un bambino che prova goffamente a esplorare un mondo nuovo e sconosciuto.

La direzione era giusta, semplicemente i tempi non erano ancora maturi.

Per quanto ciascuno di loro potesse essere geniale, da Jacques de Vaucanson a Erone di Alessandria, da Al-Jazari a Wolfgang von Kempelen, passando per Muradi e Athanasius Kircher fino ad arrivare a Leonardo da Vinci, nessuno aveva le competenze né le tecnologie adeguate a infondere vita alla materia. I loro scritti, come il *Libro del Vuoto Perfetto* o il *Codice Atlantico*, alla luce delle conoscenze attuali, più che trattati scientifici ci appaiono come semplici testamenti spirituali, documenti a cui hanno affidato i loro più intimi desideri nella speranza che un giorno potessero essere esauditi.

Ma quand'è esattamente che questo sogno ha iniziato a trasformarsi in realtà?

Senza troppe difficoltà, possiamo individuare con esattezza il momento in cui è effettivamente nata l'intelligenza artificiale: era l'estate del 1956 quando al Dartmouth College di Hanover, nel New Hampshire, un manipolo di studiosi di alto livello si riunì con il solo e unico scopo di creare una macchina che fosse in grado di simulare ogni aspetto dell'apprendimento e dell'intelligenza umana. Erano giovani sognatori, ribelli, follemente geniali.

Guidati dalla convinzione per cui "un avanzamento significativo può essere fatto in uno o più ambiti se un gruppo attentamente selezionato di scienziati ci lavora assieme per un'estate"⁸, questo gruppo di arditi intellettuali si era posto l'ambizioso obiettivo di:

⁵ Wikipedia. (2020). *Anatra Digeritrice*. https://it.wikipedia.org/wiki/Anatra_digeritrice

⁶ Wikipedia. (2020). *Macchina di Anticitera*. https://it.wikipedia.org/wiki/Macchina_di_Anticitera

⁷ Wikipedia. (2020). *Il Turco*. https://it.wikipedia.org/wiki/Il_Turco

⁸ Jerry Kaplan. (2016). *Le Persone non Servono – Lavoro e Ricchezza nell'Epoca dell'Intelligenza Artificiale*-pag. 28 http://www.club-cmmc.it/attivita/Libro%20KAPLAN_le_persone_non_servono__Lavoro_e.pdf

“esaminare la congettura che ogni aspetto dell'apprendimento od ogni altra caratteristica dell'intelligenza possa essere, in linea di principio, descritto in modo tanto preciso da poter far sì che una macchina lo simuli”.

Avevano appena due mesi per riuscirci: la loro era una disperata corsa contro il tempo.

Marvin Minsky, junior fellow di matematica e neurologia ad Harvard, Nathaniel Rochester, direttore e figura di spicco dei Laboratori di Ricerca Watson dell'IBM, Claude Elwood Shannon, padre della teoria dell'informazione e affermato ricercatore dei prestigiosi Bell Laboratories di proprietà del colosso americano delle telecomunicazioni AT&T, questi sono solo alcuni nomi dei dieci ricercatori che presero parte a quel celebre seminario che avrebbe conferito alla disciplina una natura propria e a sé stante, autonoma ma complementare ad altri campi di ricerca quali, ad esempio, l'elettronica e l'informatica. Dietro suggerimento di John McCarthy, assistant professor del Dartmouth College nonché organizzatore dell'evento, il seminario, attraverso l'utilizzo di tecniche e metodologie tipiche del brainstorming quali ad esempio discussioni comuni e dibattiti aperti e destrutturati, avrebbe dovuto far emergere un nuovo approccio teoretico alla materia.

Fu proprio in quell'occasione, peraltro, che si utilizzò per la prima volta il termine *intelligenza artificiale*: coniato dallo stesso Zio John McCarthy, l'espressione servì a mettere un po' di ordine e chiarezza in quel mare magnum linguistico che mescolava concetti e nozioni del neonato campo di ricerca e che, inevitabilmente, ne distorceva gli obiettivi.

È importante sottolineare che il convegno non venne pensato come momento di astratta riflessione da parte di un'élite di pensatori; si trattava, piuttosto, di un incubatore di idee che potesse condurre ad applicazioni e risultati squisitamente pratici, tangibili con mano: è chiaro quindi come l'attenzione di tutti i partecipanti fu catturata dal Logic Theorist di Allen Newell ed Herbert Simon, un programma per computer in grado di dimostrare fino a 38 dei 52 teoremi contenuti nel *Principia Mathematica* di Russel e Whitehead, arrivando perfino a trovare per alcuni di loro nuove e più eleganti dimostrazioni.

Era stato indiscutibilmente provato che la realizzazione di un'intelligenza artificiale (o comunque di *automated reasoning*) era possibile. In una spirale crescente di euforia ed ottimismo, quelli che seguiranno al seminario di Dartmouth saranno anni pieni di aspettative e di grandi promesse, alimentate anche dalla crescita vertiginosa dei supporti informatici utilizzabili.

Forse, per comprendere al meglio l'ebbrezza che si respirava in quel periodo, risulterà illuminante riportare la dichiarazione rilasciata da Simon e Newell in seguito al successo incontrato dal loro Logic Theorist:

“Abbiamo inventato un programma capace di pensare in modo non numerico, e quindi di risolvere il venerabile problema corpo-mente”.⁹

1.1.2 Presupposti e condizioni per la nascita dell'AI

Prima di proseguire oltre nell'esposizione del nostro argomento, sarà necessario sciogliere alcuni importanti nodi e chiarire diversi punti della tematica. Facciamo quindi un passo indietro. Sebbene il seminario tenutosi presso il Dartmouth College nell'estate del '56 sia considerato convenzionalmente (e all'unanimità della comunità scientifica) il momento esatto in cui è nata l'intelligenza artificiale e la sua disciplina, in realtà questa data non sancisce solo un inizio, ma anche e soprattutto la conclusione di un percorso durato decenni, se non addirittura secoli. Infatti, né il Logic Theorist né l'intelligenza artificiale in generale sarebbero potuti venire alla luce senza i numerosi contributi e studi che li hanno preceduti e che hanno rappresentato il presupposto necessario per lo sviluppo di intelletti sintetici. Potremmo quindi dire (senza incontrare troppe obiezioni) che se la nascita dell'AI è rintracciabile nell'estate del '56 dopo una lunga gestazione, il suo concepimento ha però origini ancora più lontane: c'è chi individua i primi embrioni di quell'idea che porterà alla nascita dell'IA già nella *macchina analitica* di Charles Babbage (siamo tra il 1834 e 1837) o, ancora prima, negli studi di Gottfried Wilhelm von Leibniz sulla meccanizzazione della ragione umana, da cui deriverà la realizzazione di una macchina in grado di effettuare somme, differenze e moltiplicazioni in maniera ricorsiva (1674). C'è perfino chi ritrova nelle pagine de *L'Ars generalis* di Raimondo Lullo (1308) concetti e nozioni fondamentali per il successivo sviluppo del calcolo computazionale, retrodatando così il suddetto concepimento a un momento ancora anteriore, fino a quasi mille anni prima della nascita dell'intelligenza artificiale. È necessario a questo punto fare una precisazione: più che parlare di precursori dell'AI, sarebbe più corretto parlare di tradizione di ricerca. Risulterebbe infatti fuorviante considerare i personaggi e le innovazioni

⁹ Silvio Henin. (2019). *AI -Intelligenza Artificiale tra Incubo e Sogno-*. p. 34.

appena esaminate (e in generale tutto ciò che precede il 1956) come appartenenti alla categoria dell'IA, sia perché considerare Lullo o Leibniz precursori dell'AI significherebbe attualizzare impropriamente tali figure, sia perché è la sola ricerca informatica degli anni Cinquanta dello scorso secolo a rappresentare (a buon diritto) l'aspetto essenziale e fondante dell'IA. Casomai, come già magistralmente illustrato da Silvio Henin nel suo romanzo *"AI – Intelligenza artificiale tra incubo e sogno"*¹⁰, furono quattro brillanti scienziati i veri ed unici precursori di quel miracolo tecnologico che noi tutti conosciamo con il nome di AI: grazie alla loro arguzia e alla loro sagacia, coltivarono intuizioni senza cui non sarebbe stata possibile la nascita di tale disciplina.

Il primo di loro è probabilmente il più noto al grande pubblico, grazie soprattutto alle numerose opere letterarie, cinematografiche e teatrali che gli sono state dedicate sul finire del XX secolo: stiamo parlando di Alan Turing, matematico, logico, crittografico e filosofo britannico, la cui storia peraltro è recentemente tornata alla ribalta per mezzo della pellicola *"The Imitation Game"* (2014).

Considerato il padre dell'informatica e dell'intelligenza artificiale, nonché uno dei più grandi matematici del XX secolo, la sua figura è ricordata principalmente per il *progetto Enigma* e per il ruolo centrale che ebbe nella Seconda guerra mondiale nella decifrazione dei messaggi che venivano scambiati tra diplomatici e militari delle potenze dell'Asse. In pochi, invece, sono a conoscenza del fatto che nel 1936 Turing pubblicò un articolo intitolato *"On computable numbers, with an application to the Entscheidungsproblem"* in cui teorizzava la realizzazione di una macchina ideale dotata di un nastro di lunghezza potenzialmente infinita, su cui sarebbe stata in grado di leggere e scrivere simboli secondo un insieme prefissato di regole ben definite. Questa macchina, conosciuta come *Macchina universale di Turing*, rappresenterà il modello di riferimento per la successiva progettazione dei moderni computer. Quattordici anni più tardi, nel 1950, ritroviamo nell'articolo *"Computing Machinery and Intelligence"*, pubblicato sulla rivista *Mind*, suggerimenti e riflessioni la cui importanza riecheggia ancora oggi: tra tutte, la più lungimirante intuizione risulta essere senza ombra di dubbio quella secondo cui per ottenere un'intelligenza sintetica che fosse il più possibile fedele a quella umana sarebbe stato conveniente coltivare tale intelletto fin dalla sua creazione, allevandolo e crescendolo come se fosse un bambino. Come si legge dal seguente estratto, Turing riconobbe per primo i benefici

¹⁰ Silvio Henin. (2019). *AI -Intelligenza Artificiale tra Incubo e Sogno-*. p. 25-33.

offerti dall'implementazione di un sistema *bottom-up* rispetto a una programmazione di tipo *top-down*:

“invece di tentare di scrivere un programma per simulare una mente adulta, perché non scriverne uno che simuli una mente infantile?”¹¹

Per incontrare John von Neumann, il secondo eroe di questa vera e propria epopea scientifica, dobbiamo spostarci nella parte opposta del globo, dalla tenuta di Bletchley Park fino a Princeton, negli Stati Uniti. Il più grande contributo dato alla nascita dell'IA dal matematico e fisico ungherese è facilmente rintracciabile nel suo scritto *“First Draft of a Report on the Edvac”*, distribuito da Herman Goldstine, responsabile alla sicurezza del progetto ENIAC, nel giugno del 1945. In questo documento incompleto di 101 pagine troviamo per la prima volta la descrizione di quell'architettura hardware a programma memorizzato che prenderà il nome di *architettura di von Neumann*: questa importante invenzione, però, altro non è che il frutto di un'intuizione avuta dallo stesso Alan Turing una decina di anni prima, contenuta nell'articolo *“On computable numbers, with an application to the Entscheidungsproblem”*; l'EDVAC (la prima macchina digitale programmabile tramite software basata sull'architettura di von Neumann), in sostanza, non sarebbe nient'altro che la realizzazione della macchina universale inventata da Turing nel 1936.

Iniziamo quindi a intravedere l'esistenza di un filo invisibile che lega tra loro questi acuti scienziati, i loro lavori, le loro teorie; ed è solo seguendo questo percorso che saremo finalmente in grado di uscire dal dedalo di invenzioni, cogliendo il reale ordine che si cela dietro i vari, apparentemente caotici, articoli e le svariate pubblicazioni.

Giungeremo così a Claude Shannon, matematico, crittografo e ingegnere elettronico statunitense il cui lavoro più importante *“A Mathematical Theory of Communication”* si riallaccia sia ai lavori di Turing che alle teorie probabilistiche di Norbert Wiener. L'articolo, pubblicato in due parti sul giornale scientifico *Bell Labs Technical Journal* nel luglio e nell'ottobre del 1948, aprì la strada alla progettazione dei moderni sistemi di codificazione, elaborazione e trasmissione dell'informazione: fu in quell'occasione che venne coniata la parola *bit*, utilizzata per designare l'unità di misura elementare d'informazione.

¹¹ Silvio Henin. (2019). *AI -Intelligenza Artificiale tra Incubo e Sogno-*. p. 29.

L'informazione, da concetto astratto e puramente intuitivo, si trasformò così in una grandezza perfettamente misurabile, matematicamente calcolabile.

Per comprendere pienamente la portata rivoluzionaria dell'entropia dell'informazione sviluppata da Shannon nel suo saggio, basti pensare che, senza questa teoria, ad oggi non esisterebbero CD, DVD, Internet e telefoni cellulari, senza considerare che lo studio della linguistica, così come l'esplorazione spaziale, sarebbero ancora a livelli arretrati.

L'ultimo elemento della nostra ricostruzione è rappresentato dai contributi provenienti dal campo della cibernetica. Questa disciplina, nata durante la Seconda guerra mondiale dalle ricerche pionieristiche di Norbert Wiener, studia meccanismi quali il calcolo automatico, la comunicazione e l'autoregolazione, focalizzandosi in particolar modo sulla modifica adattiva dei comportamenti generata dalle sollecitazioni esterne dell'ambiente circostante (retroazione). Dopo la pubblicazione di *"Behavior, Purpose and Teleology"* nel 1943 da parte di Wiener (che come già anticipato in precedenza sancirà la nascita della disciplina), tra i più importanti eventi che segneranno la cibernetica possiamo annoverare senza ombra di dubbio la formalizzazione dell'equivalenza logica tra il sistema nervoso e il calcolatore elettronico, come dimostrato da Warren McCulloch e Walter Pitts ne *"A Logical Calculus of Ideas Immanent in Nervous Activity"* (1943) attraverso l'uso dell'algebra booleana.

Come noterà von Neumann:

"Si è sostenuto spesso che le attività e le funzioni del sistema nervoso umano siano così complesse da non poter essere eseguite da nessun meccanismo ... Il risultato di McCulloch e Pitts mette fine a tutto questo e prova che tutto ciò che può essere descritto completamente e senza ambiguità a parole, può essere ipso facto realizzato con una rete neurale finita"

L'articolo di McCulloch e Walter Pitts, che si basava fortemente sulla conoscenza fisiologica dei neuroni e sulla teoria della computabilità di Alan Turing, dimostrò non solo che ogni funzione calcolabile poteva essere elaborata da una qualche rete di neuroni interconnessi, ma soprattutto negò il dogma dell'inimitabilità delle funzioni intellettive dell'uomo, annullando così la distanza che per millenni aveva separato l'uomo dalla macchina. Dicotomia che è entrata ancora più in crisi nel 1949 in seguito alla pubblicazione de *"L'organizzazione del comportamento"*, in cui lo psicologo canadese Donald Olding Hebb scandagliava il nesso esistente tra sistema nervoso e comportamento, giungendo così alla formulazione di una nuova teoria neuropsicologica dell'apprendimento associativo.

Ora che abbiamo dipanato il filo sottile che lega discipline e scienziati in quella fitta rete di influenze reciproche che condurrà alla nascita dell'intelligenza artificiale, possiamo riprendere la trattazione dell'argomento centrale del presente elaborato.

1.1.3 Il ventennio onirico dell'AI (1950-1969)

Gli anni che seguirono la conferenza di Dartmouth furono pieni di entusiasmo e ottimismo: ogni giorno nuovi e ambiziosi progetti fiorivano in tutto il paese, poli come il MIT, l'Università di Stanford e IBM Corporation rappresentavano veri e propri centri di ricerca all'avanguardia e la DARPA (*Defense Advanced Research Projects Agency*), di proprietà del Ministero della Difesa statunitense, investiva cospicuamente nel settore con la speranza di poter utilizzare, un giorno, una tecnologia così rivoluzionaria per i propri scopi.

In quel periodo, sembrava come se ogni cosa fosse possibile: l'unico ostacolo che si sarebbe potuto interporre tra l'uomo e quel futuro così promettente era la limitatezza della prospettiva futura, l'incapacità di comprendere le infinite possibilità che l'AI avrebbe potuto offrire.

Il susseguirsi convulso dei numerosi successi incoraggiò i pionieri dell'intelligenza artificiale: Newell e Simon svilupparono un programma in linguaggio IPL chiamato *General Problem Solver*, in grado di risolvere problemi generali formalizzati grazie alla sottopartizione in più obiettivi; McCarthy ideò un nuovo linguaggio di programmazione, il LISP, e teorizzò *Advice Taker*, un ipotetico programma per computer che si sarebbe dovuto basare (per la prima volta) sulla capacità di ragionamento sul senso comune¹²; nello stesso periodo, Minsky si dedicò allo studio dei *micromondi*¹³, Rochester produsse i primi programmi di AI mentre Herbert Gelernter, nel 1959, programmò il *Geometry Theorem Prover*, capace di dimostrare alcuni dei teoremi di Euclide.

E ancora, potremmo ricordare: il Perceptron di Rosenblatt, che spalancò le porte a un nuovo approccio, chiamato *connessionismo* (1958); il SAINT di James Slagle, primo sistema esperto

¹² ScienceDirect.com (2020). *The Advice Taker/Inquirer, a System for High-Level Acquisition of Expert Knowledge*. <https://www.sciencedirect.com/science/article/abs/pii/S0736585388800327>

¹³ Con il termine micro-mondo si fa riferimento a un problema limitato che richiede l'utilizzo della logica per la sua risoluzione.

in assoluto (1961); il programma ANALOGY di T.G. Evans (1962); il robot SHAKEY di Charles Rosen (1966); il software SHRDLU di Terry Winograd (1968).

Il lancio dello Sputnik 1 e l'inasprirsi della Guerra fredda tra Unione sovietica e Stati Uniti provocarono una brusca accelerazione nello sviluppo dell'AI: per Washington, infatti, era inaccettabile che l'URSS trionfasse nella corsa allo spazio e, più in generale, nella lotta per la supremazia tecnologica.

I finanziamenti governativi aumentarono e vennero canalizzati verso specifici ambiti e applicazioni: in particolar modo, al fine di ridurre il gap tecnologico che separava i due paesi, gli USA si concentrarono nello sviluppo di un'intelligenza artificiale in grado di tradurre testi dal russo all'inglese; per comprendere i progressi tecnologici compiuti dall'Unione Sovietica, infatti, era indispensabile per gli Stati Uniti essere in grado di decodificare un'enorme quantità di articoli scientifici russi il più rapidamente possibile. Un compito del genere, tanto delicato quanto importante, poteva essere portato a termine solo dalle macchine. In realtà, però, i risultati non furono quelli auspicati: sebbene i testi tradotti presentassero una struttura grammaticale e sintattica impeccabile, il loro significato risultava del tutto travisato. La leggenda vuole che la frase *"The spirit is willing, but the flesh is weak"* (Lo spirito è forte ma la carne è debole), tradotta prima in russo e poi nuovamente in inglese applicando la stessa metodologia utilizzata per i saggi sovietici, si trasformasse in *"The vodka is good, but meat is rotten"* (La vodka è buona ma la carne è marcia). Il motivo di una traduzione così grottesca è facilmente rinvenibile nell'approccio adottato per la risoluzione del problema: molto ingenuamente, infatti, venne digitalizzato un dizionario russo-inglese e venne implementato un programma che sostituisse automaticamente, parola per parola, ciascun termine russo con il proprio corrispettivo anglosassone, senza curarsi minimamente di contesto, senso logico e significati nascosti del testo originario.

Nonostante questo trascurabile insuccesso, le riviste scientifiche e i giornali diffondevano ottimismo, certi del benessere che gli irrefrenabili sviluppi dell'AI avrebbero portato nel prossimo futuro.

Giornali e riviste di ogni tipo traboccavano di frasi come quelle riportate di seguito:

"Entro dieci anni un calcolatore digitale farà scoprire e dimostrare un nuovo importante teorema matematico"

Allen Newell

“Le macchine saranno in grado, entro vent’anni, di fare qualsiasi lavoro che un umano possa fare”

Herbert Simon

“In futuro, tra i prossimi tre e otto anni, avremo una macchina con l’intelligenza generale di un essere umano medio”

Marvin Minsky

Venivamo continuamente preparati all’imminente arrivo di una nuova Età dell’oro.

1.1.4 “Primo grande inverno” e rinascita dell’AI (1969-1986)

D’improvviso, però, queste speranze e questi sogni vennero traditi.

Le grandi aspettative che tutto il mondo aveva riposto nei confronti dell’AI iniziarono a incrinarsi già nel 1966, quando il Rapporto ALPAC¹⁴ criticò aspramente i risultati che fino a quel momento erano stati raggiunti nel campo della traduzione automatica, determinando l’interruzione dell’erogazione di sostegni economici da parte del Consiglio Nazionale delle Ricerche (NRC).

Ad aggravare la situazione che si era venuta a creare in quel periodo, furono le critiche contenute nel libro *“Perceptrons: an introduction to computational geometry”*, edito nel 1969. Nel loro scritto, Minsky e Papert, attraverso la meticolosa raccolta di prove matematiche, evidenziarono i rilevanti limiti che incontravano i perceptron di Rosenblatt nel calcolo di alcuni predicati di connessione, e in particolar modo della funzione logica XOR.

La rottura definitiva tra AI e istituzioni governative/accademiche avvenne poco più tardi: nel 1969 il parlamento americano, in applicazione di quanto stabilito nell’Emendamento Mansfield, chiese alla DARPA di finanziare unicamente progetti volti allo sviluppo di tecnologie militari funzionali, escludendo così da ogni tipo di sovvenzionamento la ricerca generale e non-orientata; nel Regno Unito, invece, le ricerche sull’AI si arrestarono nel 1973

¹⁴ Comitato scientifico istituito nel 1964 negli USA. Guidato da John R. Pierce, aveva il compito di valutare i progressi raggiunti in linguistica computazionale e, più in generale, nella traduzione automatica.

come conseguenza del Rapporto Lightill, una relazione redatta dal professor Sir James Lightill per il *Science Research Council*¹⁵. Il documento valutava pessimisticamente i frutti prodotti dagli studi sull'intelligenza artificiale, soffermandosi sulle preoccupanti perplessità sollevate dall'esplosione combinatoria¹⁶ e affermando che, sebbene le tecniche utilizzate funzionassero bene nell'ambito di piccoli domini matematici, mal si adattavano alla risoluzione di problemi caratterizzati da domini più ampi, più vicini alle casistiche del mondo reale.

Come sentenzia duramente il rapporto nel trarre le proprie conclusioni:

"In nessuna parte del campo le scoperte fatte finora hanno prodotto il grande impatto che è stato poi promesso"

A eccezione di poche università di secondo livello (tra cui University of Edinburgh e University of Sussex), si assistette al completo smantellamento di interi laboratori destinati alla ricerca dell'intelligenza artificiale.

Iniziò così per l'AI un periodo di stagnazione caratterizzato dalla riduzione di finanziamenti e da un dilagante pessimismo: era il Primo grande inverno dell'AI, che durerà fino agli inizi degli anni Ottanta.

Ma quali furono le cause che portarono a un arresto così improvviso delle ricerche?

Come emerse più tardi, le problematiche che avevano impedito ulteriori avanzamenti nello sviluppo dell'AI erano riconducibili a tre categorie differenti.

Innanzitutto, il primo grande problema che preoccupò gli studiosi riguardava lo stesso *modus operandi* del computer: questo, infatti, esaminava ogni singolo caso possibile, analizzandone condizioni, esiti, probabilità, valutando ogni variabile e combinazione. Adottare un simile metodo (noto come *ricerca esaustiva della soluzione* o, più comunemente, *metodo forza bruta*) per giungere alla soluzione ottimale esistente non solo risulta in alcuni casi eccessivamente dispendioso, ma poteva rivelarsi un sistema troppo lento, perfino per i più moderni supercomputer di oggi. Si consideri inoltre che all'epoca la velocità di elaborazione dei dati e la memoria disponibile erano a dir poco inadeguati: basti pensare che il Semantic Memory System pensato da Ross Quillian nel 1966 poteva supportare al massimo un vocabolario composto solamente da venti parole.

¹⁵ Agenzia britannica incaricata del finanziamento di attività di ricerca scientifica e ingegneristica tramite fondi pubblici

¹⁶ Wikipedia. (2020). *Combinatorial Explosion*. https://en.wikipedia.org/wiki/Combinatorial_explosion

La seconda criticità incontrata concerneva l'abilità stessa del computer nel trovare soluzioni al problema che gli veniva posto: il Logic Theorist era in grado di dimostrare teoremi matematici, ma non quelli geometrici; viceversa, il Geometry Theorem Prover poteva dimostrare anche complessi teoremi geometrici, ma nulla poteva di fronte anche al più semplice quesito matematico. Come chiarisce esaurientemente l'esempio esposto, i sistemi fino ad allora sviluppati potevano essere infallibili per determinati ambiti o circostanze, ma risultavano del tutto inutili se traslati in altre situazioni; bastava che il problema cambiasse anche solo di poco o che fosse posto in maniera in parte diversa, ed ecco che il programma diventava completamente inutilizzabile. Non si riusciva a creare un sistema che fosse abbastanza generale e versatile da adattarsi a molteplici contesti variabili: gli unici problemi che potevano essere affrontati da quelle primitive forme di AI dovevano necessariamente essere definiti da regole semplici e da un linguaggio elementare, non ambiguo; in una parola, dovevano essere facilmente formalizzabili. Un escamotage per ovviare a questo inconveniente poteva essere quello di aumentare regole e dati, estendendo di fatto il dominio del programma. In questo modo, però, si sarebbe facilmente ricaduti in un problema di esplosione combinatoria¹⁷, e conseguentemente, di intrattabilità¹⁸ del quesito.

Infine, l'ultimo e insormontabile ostacolo: l'incapacità di dotare la macchina di quel tipo di conoscenza che è nota come "conoscenza comune", l'insieme di informazioni e nozioni che apprendiamo, quasi inconsciamente, nei primissimi anni di vita. Anche se questo tipo di conoscenza può sembrarci banale o addirittura scontata, in realtà si tratta di una mole impressionante di dati difficilmente codificabili (in quanto legati a intuizioni e sensazioni), che molto spesso presuppongono a loro volta altre conoscenze implicite.

Come possiamo spiegare a una macchina che non bisogna mettere le mani sul fuoco se non vogliamo bruciarci e provare dolore? E soprattutto, come possiamo spiegargli cosa sia il dolore?

Come chiarirà Moravec nel suo famoso paradosso, contrariamente a quanto si crede, le capacità senso-motorie di basso livello richiedono enormi risorse computazionali, molto più del ragionamento di alto livello; in parole povere, è più facile che un software dimostri l'ultimo teorema di Fermat, che un robot impari come versare correttamente il vino nel calice.

¹⁷ In matematica, l'incremento esponenziale del numero di operazioni da effettuare dovuto all'aumento del numero di dati

¹⁸ Un problema si definisce intrattabile se, nonostante possa avere o meno una soluzione, non esiste un algoritmo con complessità polinomiale in grado di risolverlo

Sebbene “il primo grande inverno” scosse l’ambito dell’AI fin dalle sue fondamenta, arrivando perfino a far dubitare della stessa validità della disciplina, al contempo permise agli studi sull’intelligenza artificiale di riorganizzarsi, imparare dai propri errori e ridefinire i propri obbiettivi: si configurò così un nuovo ordine.

Da quel momento iniziarono a delinearsi due scuole di pensiero diverse e contrapposte, contraddistinte per visione e obiettivi: una mirava a replicare l’insieme delle facoltà intellettive e cognitive dell’uomo, emulando il modo in cui si risponde agli stimoli provenienti dal modo esterno: era la cosiddetta “*intelligenza artificiale forte*” (*strong AI*); l’altra, invece, si focalizzò esclusivamente sull’esecuzione di singoli e limitati lavori (ad esempio il gioco degli scacchi o della dama) appartenenti a un determinato dominio. Per contrapporla e distinguerla dalla prima, venne chiamata “*intelligenza artificiale debole*” (*weak AI* o anche *narrow AI*).

È importante notare come tale distinzione, però, è tutt’altro che netta e assoluta. Anzi, si potrebbe dire che il confine esistente tra *strong AI* e *weak AI* sia fumoso e incerto, così come la loro stessa definizione: ad esempio, un’intelligenza debole potrebbe essere considerata forte nell’ambito ristretto in cui opera; al contrario, un’intelligenza forte potrebbe risultare debole in un campo più ampio del suo normale dominio. Si tratta quindi di una classificazione relativa e mutevole: come in qualsiasi gioco prospettico, dipende tutto dal punto da cui osserviamo il fenomeno.

Tra i tanti elementi che il grande inverno sovvertì completamente, i più rilevanti furono senza dubbio l’identità dei finanziatori e le competenze richieste alle IA; se prima degli anni ’80 i principali attori erano stati enti governativi e mondo accademico, dopo quella battuta d’arresto, furono le aziende e i venture capitalist ad assumere il ruolo di investitori e mecenati, che offrirono nuove prospettazioni all’intelligenza artificiale.

Il capitale privato sostituì quello pubblico. Cambiarono pure le prospettive: non si trattava più del puro e astratto desiderio di conoscenza, né tantomeno del sogno di imperialismo americano che aveva portato la DARPA ad elargire cospicue somme di denaro; la logica del profitto divenne il fine ultimo di quegli studi.

Si passò così dall’ambiente accademico e governativo all’industria privata.

Una schiera di giovani imprenditori hi-tech fiutò le potenzialità offerte dall’AI debole e puntò (con grande lungimiranza) su di questo tipo di approccio, accantonando il sogno di realizzare un’intelligenza artificiale generale (*Artificial General Intelligence*): troppe avversità tecniche ne impedivano ancora la creazione, mentre la *weak AI* era lì, a portata di mano, facile da implementare e altamente remunerativa. Si assistette così all’esplosione dei sistemi esperti, programmi in grado di riprodurre prestazioni di individui esperti per risolvere problemi

appartenenti a specifici campi di attività. Nel 1969 vide la luce il sistema esperto DENDRAL, capace di identificare molecole organiche sconosciute grazie all'analisi spettroscopica¹⁹. Nel 1972 nacque invece MYCN, progettato come sostegno ai medici per il riconoscimento delle infezioni batteriche dei pazienti, per la prescrizione dell'antibiotico e perfino per il dosaggio consigliato.

A queste prime creazioni ne seguiranno molte altre: MOLGEN, PROSPECTOR, XCON e STREAMER sono solo alcuni dei tanti nomi che circolarono nell'ambiente di professori, scienziati e giornalisti dell'epoca. L'importanza dei sistemi esperti, nel tempo, crebbe a tal punto che sul finire degli anni Ottanta quasi ogni azienda statunitense aveva attivo un proprio sistema esperto. Gli stessi Governi riconobbero il ruolo strategico che questi avrebbero potuto assumere e li posero al centro di numerosi progetti: basterà ricordare il *Fifth Generation* implementato dal Giappone, o il consorzio americano *Microelectronics and Computer Technology Corporation (MCC)*, pensato come risposta al minaccioso programma nipponico. Nel frattempo, si assistette anche al ritorno delle reti neurali: la *rete neurale associativa* di John Joseph Hopfield (1982) e, soprattutto, l'algoritmo *back-propagation* di Geoffrey Hinton e David Rumelhart (1986) risultarono fondamentali per la riabilitazione del connessionismo.

Gli anni Ottanta del secolo scorso, in generale, si caratterizzarono per il rifiorire delle ricerche sull'intelligenza artificiale, rappresentando per l'AI una vera e propria primavera: la potenza di calcolo aumentò vertiginosamente, così come la capacità di memoria disponibile; i personal computer iniziarono a proliferare e la rete Internet (1969) accorciò ogni distanza, rendendo l'intero globo un luogo percorribile

La *Legge di Moore*, che iniziò a verificarsi in quel periodo, ci promise che ogni due anni il numero dei transistor presenti sulla piastrina di silicio dei microprocessori sarebbe raddoppiata, contrariamente al loro costo unitario, che invece si sarebbe dimezzato.

Non avevamo mai disposto, prima di allora e in tutta la storia dell'umanità, di strumenti così potenti a un prezzo così basso.

Il futuro sembrava nuovamente promettente e quel primo grande inverno dell'intelligenza artificiale appariva ormai lontano.

¹⁹ Wikipedia. (2020). *Spettroscopia*. <https://it.wikipedia.org/wiki/Spettroscopia>

1.1.5 “Secondo grande inverno” (1987- 1993)

La stessa fiducia che aveva permesso al settore di uscire dall'impasse di qualche anno prima, fu la causa della rovina che si abbatté per la seconda volta sull'intelligenza artificiale: l'ondata di estremo ottimismo offuscò la capacità di giudizio dei trader e l'oggettività lasciò il posto a irrazionali convinzioni, troppo spesso lontane dalla realtà dei fatti.

Il mercato degli hardware specializzati iniziava a mostrare le stesse perverse fluttuazioni che scuotono i mercati nella fase precedente lo scoppio di una bolla speculativa, arrivando a valere la stratosferica cifra di mezzo miliardo di dollari. Nonostante le preoccupazioni mostrate dagli esperti del settore circa l'imminente crollo del mercato, la bolla continuò a gonfiarsi senza sosta finché nel 1987, tre anni dopo le allarmanti dichiarazioni rilasciate da Schank e Minsky, esplose con tutta la sua drammaticità: in una sola notte venne spazzato via l'intero settore, con conseguente fallimento di oltre trecento aziende operanti nel ramo AI entro la fine del 1993.

Era calato il Secondo grande inverno, che recava con sé la fine della prima ondata commerciale dell'intelligenza artificiale.

Fattore scatenante di quello stravolgimento finanziario fu il surclassamento delle macchine Lisp da parte di nuovi e più performanti prodotti. L'egemonia delle macchine LISP, commercializzate da Symbolics e affini, infatti, venne messa a dura prova già dall'ingresso nel mercato di workstation supportati da *Lisp IDE (Integrated Design Environment)*²⁰, risultato di partnership vincenti tra aziende come la Lucid Inc. e Sun Microsystem. Ma il completo e irreversibile tracollo si ebbe quando i computer desktop commercializzati da imprese come Apple e IBM divennero più potenti delle costose macchine Lisp. Nonostante il disperato tentativo di offrire versioni sempre più aggiornate che potessero rivaleggiare con i nuovi prodotti, la Symbolics (così come le altre imprese simili) non era più in grado di arginare la dirompente ascesa dei propri competitors, vedendo sottrarsi in breve tempo quelle quote di mercato che a gran fatica si era assicurata nel corso degli anni.

Tra i tanti eventi che si registrarono in quel sessennio di sconvolgimenti, è doveroso ricordarne due in particolare, che gettarono ancora più nel caos il settore dell'AI: il primo fu la caduta dei sistemi esperti, rivelatisi estremamente fragili, incapaci di apprendere e difficili da aggiornare, oltre che troppo costosi da mantenere; le poche società produttrici di sistemi esperti sopravvissute alla bufera finanziaria furono costrette a ridimensionarsi e ad esplorare nuovi

²⁰ Software che durante la fase di programmazione supporta i programmatori nello sviluppo del codice sorgente di un programma

paradigmi, come il *case-based reasoning*²¹ o l'*universal database acces*²². Il secondo avvenimento degno di nota fu il fallimento della *Fifth generation*: alla fine del 1991 soltanto pochi dei numerosi obbiettivi che il Giappone si era posto risultarono effettivamente realizzati. Ancora una volta, le aspettative si erano mostrate troppo ambiziose rispetto al reale livello di sviluppo tecnologico che era stato raggiunto.

²¹ Wikipedia. (2020). *Case-Based Reasoning*. https://en.wikipedia.org/wiki/Case-based_reasoning

²² ProgSoft.Net (2020). *Universal Data Access Components*. <https://progsoft.net/it/software/universal-data-access-components>

1.2 L'AI come forma di intelligenza: ipotesi o realtà?

Abbiamo visto in che modo l'intelligenza artificiale sia nata e si sia evoluta nel corso del tempo, quali ostacoli abbia dovuto superare, come abbia dovuto reinventarsi per poter sopravvivere fino ai nostri giorni. Abbiamo disquisito di scienziati, articoli, pubblicazioni, sempre con la sensazione, però, di omettere qualcosa di importante. Infatti, più ci addentriamo nelle tematiche del nostro elaborato, più sentiamo crescere dentro di noi il bisogno di rispondere ad una semplice domanda: le macchine possono davvero essere considerate "intelligenti"? O si limitano semplicemente a eseguire ciò che l'uomo impone loro di fare, prive di qualsiasi forma di autonomia e totalmente ignare di sé stesse e delle proprie azioni, incapaci di comprendere il senso e lo scopo di ogni loro operare?

Senza ombra di dubbio, il quesito che si prospetta noi è spinoso, ingannevole, difficile da sciogliere; per poterci districare tra questi ingarbugliati pensieri, pertanto, avremo bisogno di tutto l'aiuto possibile.

A tal proposito sarà utile conoscere i giudizi e le opposte conclusioni a cui giunsero due grandi pensatori del passato: questo non ci renderà certo esperti della materia né fornirà maggiore autorevolezza alle nostre valutazioni; quasi sicuramente, però, le loro riflessioni potrebbero offrire nuove e sorprendenti prospettive da cui osservare il problema.

1.2.1 The Imitation Game

Immaginate di ritrovarvi rinchiusi in una stanza vuota, completamente isolata dal resto del mondo. Non potete né chiamare né chiedere aiuto, non c'è alcun condotto di aerazione che vi possa garantire una fuga, per quanto rocambolesca possa essere.

In mezzo ad un ambiente tanto spoglio, esattamente al centro del locale, campeggia un tavolo. Accanto vi è una sedia; è rivolta verso di voi, sembra quasi che esorti a sedervi. Decidete quindi di accettare quel tacito invito e vi dirigete verso quel misterioso mobilio. Siete un po' titubanti. Vi sedete.

Sul tavolo trovate un cellulare e un foglio di carta: c'è scritto che l'unico modo che avete per evadere da quella prigione è vincere a un gioco, un gioco di imitazione. Le regole sono molto semplici: usando il cellulare che vi è stato fornito dovete conversare con due individui, A e B; obiettivo del gioco è quello di capire, solo in base alle risposte che riceverete ai vostri messaggi, se l'interlocutore con cui state colloquiando sia un uomo, una donna o una macchina.

Sareste in grado di capire chi si cela realmente dalla parte opposta del telefono?

Anche se in maniera eccessivamente romanzata e teatrale, quanto scritto sopra illustra fedelmente il celebre Test di Turing, apparso per la prima volta sulla rivista Mind nel 1950.

Nell'articolo intitolato "*Computing machinery and intelligence*", Alan Turing espone dettagliatamente il criterio che, a parer suo, meglio permetteva di rispondere al dubbio amletico (le macchine pensano o semplicemente eseguono?). Il matematico britannico si era a lungo interrogato sulla questione senza mai giungere a conclusioni soddisfacenti finché, improvvisamente, un popolare gioco di società non gli diede l'illuminazione: si trattava del famoso *gioco dell'imitazione*. Stando a quanto ipotizzato da Turing, sarebbe bastato far partecipare una macchina all'*Imitation game* per poter valutare adeguatamente la capacità di pensiero di questa e, più in generale, la sua intelligenza: se il computer avesse elaborato risposte così spontanee da risultare indistinguibili da quelle fornite dal concorrente umano, allora si sarebbe dovuto necessariamente concludere che, visto il grado di esattezza con cui il computer simulava il pensiero umano, dovesse necessariamente considerarsi intelligente.

Nello stesso articolo, Turing spiegava anche come si dovesse intendere correttamente il termine intelligenza: a suo avviso, poteva definirsi intelligente una macchina che fosse in grado di concatenare tra loro una serie di idee, elaborando di conseguenza risposte coerenti e non prive di significato.

Come si legge dal seguente estratto di "*Computing machinery and intelligence*", riferendosi al cogito cartesiano,

«Secondo la forma più estrema di questa opinione, il solo modo per cui si potrebbe essere sicuri che una macchina pensa è quello di essere la macchina stessa e sentire se si stesse pensando. [...] Allo stesso modo, la sola via per sapere che un uomo pensa è quello di essere quell'uomo in particolare. [...] Invece di discutere in continuazione su questo punto, è normale attenersi alla educata convenzione che ognuno pensi.»

È evidente come Turing non si ponesse nemmeno il problema di come la macchina giungesse alla formulazione di una determinata risposta. Poco importava se la macchina fosse stata guidata da una propria scelta cosciente o se, più semplicemente, si fosse limitata a eseguire quanto le era stato chiesto di fare, totalmente ignara di ciò che accadeva e che sarebbe accaduto: gli unici aspetti che secondo Turing dovevano essere considerati rilevanti erano il comportamento adottato e il risultato a cui questo conduceva, indipendentemente dalle cause sottostanti che avevano portato la macchina ad attuare esattamente quelle specifiche azioni.

Dopotutto, come osservava cinicamente il genio britannico, non potremmo nemmeno essere certi delle abilità logiche di chi ci circonda: secondo Turing ogni volta che istintivamente giudichiamo il grado di intelligenza di un individuo, in realtà, non stiamo valutando le sue capacità logiche ma i gesti che egli compie; le conclusioni successive a cui giungiamo per metodo induttivo, più che la conclusione inconfutabile di un teorema, non sarebbero altro che semplici postulati, principi indimostrati di cui a priori ammettiamo la validità per pura convenzione.

Incapaci di verificare empiricamente le capacità di ragionamento degli individui, non abbiamo altra scelta se non studiare la mente umana come fosse una misteriosa *black box*, una scatola nera che comprendiamo nelle sue correlazioni input-output ma di cui ignoriamo il funzionamento interno.

Se le conclusioni a cui giungiamo dipendono unicamente da output certi e risultati tangibili, estranei a meccanismi interni e processi mentali, perché non utilizzare gli stessi criteri e lo stesso metodo anche per l'Intelligenza Artificiale?

1.2.2 Stanza cinese

Numerose furono le critiche che vennero mosse ad Alan Turing e sul suo lavoro: in molti infatti misero in dubbio la stessa validità del test, accusando il matematico inglese di esser stato eccessivamente influenzato da visioni e dottrine tipiche del *comportamentismo*²³.

Nonostante il test di Turing avesse incontrato copiose obiezioni e un nutrito dissenso, prive del supporto di qualsiasi argomentazione logica o evidenza empirica le riserve avanzate rimasero per lungo tempo delle semplici considerazioni sporadiche, che mai si articolarono in una vera confutazione; bisognerà aspettare il 1980 per far riferimento ad un documento equipollente al “*Computing machinery and intelligence*” e contestare scientificamente quanto teorizzato da Turing nel suo scritto. In quell’anno, infatti, venne pubblicato sulla rivista “*The Behavioral and Brain Sciences*” il saggio intitolato “*Minds, Brains and Programs*”, in cui il filosofo John Searle negava la possibilità che le macchine fossero dotate di capacità di pensiero.

Searle, ricalcando in parte le condizioni prospettate da Turing nel suo *gioco delle imitazioni*, propose un esperimento mentale: provò a immaginare cosa sarebbe successo se, chiuso in una

²³ Il *comportamentismo* (o anche *behaviorismo*) è un approccio alla psicologia sviluppato agli inizi del Novecento da John Watson; fondamento della corrente di pensiero era il rifiuto dell’introspezione quale oggetto di analisi scientifica

stanza, avesse dovuto conversare con una persona esterna tramite messaggi scritti su fogli di carta. Prima ipotesi fondamentale per il proprio ragionamento era che l'interlocutore esterno (totalmente ignaro di quanto avvenisse all'interno della stanza) fosse in grado di comunicare unicamente in cinese. Serale, a sua volta, non aveva alcuna dimestichezza con gli *hànzì*²⁴: data la reciproca incapacità di comprendersi, per i due dialogare sarebbe stato normalmente impossibile. Scrivo normalmente perché Searle, come seconda condizione necessaria alla propria dimostrazione, suppose che all'interno della stanza vi fosse un manuale d'istruzione (scritto in inglese) che gli indicasse, attraverso un complesso sistema di regole, come rispondere correttamente ai messaggi che gli venivano recapitati dal proprio corrispondente cinese. Sebbene Searle ignorasse completamente il contenuto dei messaggi era, ciononostante, ugualmente in grado di portare avanti la conversazione in maniera così brillante da indurre il proprio interlocutore a credere di discorrere con un vero madrelingua, quando in realtà egli si era limitato ad eseguire, senza alcuna cognizione di causa, un insieme di regole predefinite.

Grazie a questa acuta allegoria, Searle dimostrò quanto fosse inesatto considerare intelligenti le macchine: per quanto le attività svolte dai computer siano impeccabili, per quanto i risultati a cui giungono possano essere simili a quelli di un essere umano, le loro azioni mancheranno sempre di intenzionalità. La semplice e inconsapevole manipolazione di simboli non basta per poterli definire intelligenti: privi di reali funzioni cognitive, incapaci di autodeterminarsi e di essere consapevoli, gli intelletti sintetici risulterebbero soltanto la mimesi grottesca e caricaturale di quelli umani.

Come emerse dalla *Stanza cinese*, il divario abissale esistente tra sintassi e semantica non permetteva di assimilare la passiva esecuzione di compiti e mansioni al concetto di intelligenza. Era necessario qualcosa di più, una peculiarità tipica dell'uomo che le macchine non erano ancora riuscite ad emulare: la consapevolezza.

1.2.3 Valutazioni

Chi tra i due intellettuali aveva ragione? Alan Turing, che ammetteva la possibilità di riconoscere le macchine quali soggetti intelligenti, o John Searle, fervido sostenitore del parere contrario, secondo cui i computer, per quanto abili nell'imitare i comportamenti dell'uomo, mai ne avrebbero acquisito le capacità intellettive?

²⁴ Trascrizione in lingua latina dei logogrammi cinesi con cui ci si riferisce al sistema di scrittura cinese

Per quanto possa sembrare assurdo e contrario ad ogni intuizione, le opposte conclusioni a cui giunsero i due studiosi potrebbero essere considerate, rispettivamente, corrette ed errate; tutto dipende dal significato che si intende attribuire al termine “intelligenza”.

Infatti, qualora si ritenga necessaria la presenza di consapevolezza di sé e del proprio operato nonché di autodeterminazione per il sostanzarsi di una qualsivoglia forma di intelligenza, allora, in linea con il pensiero espresso da John Searle nel saggio “*Minds, Brains and Programs*”, si deve ritenere che le attuali forme di intelligenza artificiale non possano essere definite “intelligenti”. Qualora invece si ritenga sufficiente la risoluzione di specifici problemi in domini limitati per qualificare i computer “intelligenti”, allora - in accordo con quanto sostenuto da Alan Turing in “*Computing machinery and intelligence*” – bisognerà necessariamente riconoscere una certa forma di intelligenza negli attuali sistemi di AI.

C'è da dire comunque che oggi, come dimostrato da Searle, i sistemi di AI sono privi di una propria consapevolezza, tanto di loro stessi quanto delle loro azioni.

Bisogna tuttavia e non di meno chiedersi, a questo punto, se si possa affermare e dimostrare con certezza che l'intelligenza artificiale mai in assoluto possa essere in grado, un giorno, di acquisire simili capacità.

Al riguardo si potrebbe considerare che se le macchine non sono ancora in grado di pensare, di ragionare e di avere, come noi, autodeterminazione questo deriva dal non essere mai state indotte a ciò dall'operato dell'uomo e dalla mancanza di un'istruzione loro impartita, di uno sviluppo loro assicurato.

Proprio come aveva già suggerito Turing, sarebbe al riguardo proficuo -per il raggiungimento del risultato- assimilarle ad una mente infantile, a cui va insegnato come se fosse un bambino. L'intuizione di Turing potrebbe risultare, tuttavia, incompleta. Come non considerare, infatti, che risultati maggiori si raggiungerebbero qualora si concretizzasse anche la possibilità di dotare le macchine di un corpo senziente?

In una simile prospettiva, condizione più che utile sarebbe l'esistenza di un corpo fisico. Si pensi: in tal modo le macchine avrebbero un corpo, capace di cogliere molteplici sensazioni. Concetti ed informazioni sarebbero a loro più facilmente accessibili e comprensibili, consentendo loro di scoprire ed imparare maggiormente.

Sarebbero di certo più assimilabili all'essere umano, proiettate a ragionare quasi al nostro pari.

PARTE II – PRESENTE

2.1 L'Intelligenza Artificiale al giorno d'oggi

Anche se non ce ne rendiamo conto, l'intelligenza artificiale è già in mezzo a noi e sta silenziosamente trasformando le nostre vite ad una velocità impressionante.

La maggior parte delle persone crede erroneamente che l'AI sia un qualcosa di così complesso e avanzato da essere distante anni luce dalla realtà e dalla vita di tutti i giorni, ma non è così; quando cerchiamo qualcosa su Google o scorriamo il catalogo dei film Netflix, ogni volta che eliminiamo le mail di spam con cui Amazon ci consiglia esattamente il prodotto di cui abbiamo bisogno, senza saperlo, stiamo interagendo con un'intelligenza artificiale.

Per quanto il campo pubblicitario rappresenti senza dubbio l'ambito in cui l'*artificial intelligence* brilla maggiormente, non si tratta dell'unico settore in cui invisibili intelletti sintetici, febbrili e instancabili, lavorerebbero senza sosta: l'industria automobilistica, la medicina, l'industria videoludica, la finanza, la ricerca scientifica, l'industria bellica e aereospaziale sarebbero solo alcuni dei numerosi settori in cui l'intelligenza artificiale si sarebbe insinuata, trasformando e sovvertendo dall'interno il modo stesso di operare e fare business.

Ma quali sono le nazioni e i settori economici maggiormente coinvolti in questa silenziosa rivoluzione?

2.1.1 Dall'analisi microeconomica...

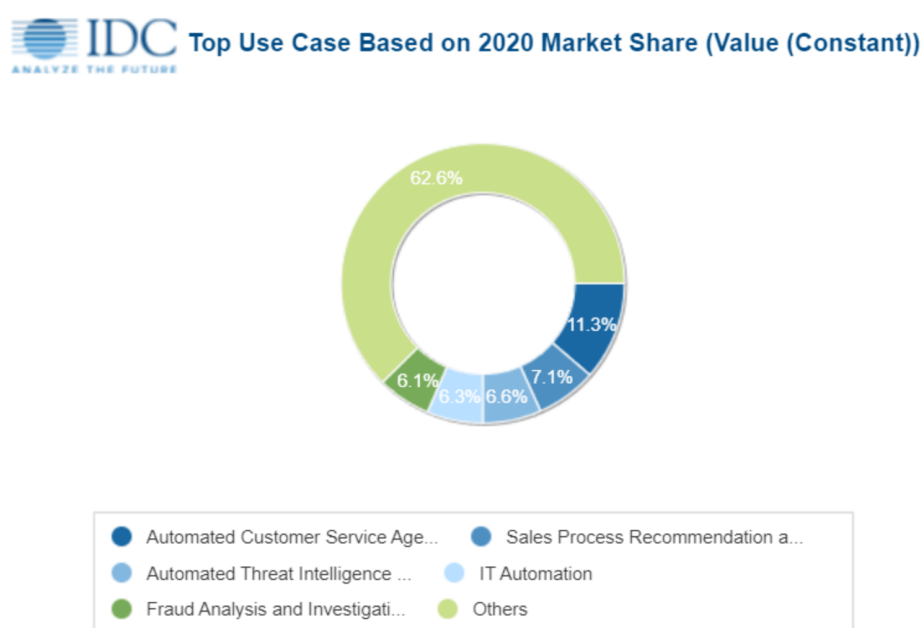
Secondo quanto riportato dall'International Data Corporation (IDC) nella "*Worldwide Artificial Intelligence Spending Guide*" i livelli globali di spesa in relazione all'intelligenza artificiale per i prossimi quattro anni sembrerebbero andare incontro ad un notevole incremento. Le cifre riportate dall'IDC sono sconvolgenti: da 50,1 miliardi di dollari nel 2020 a oltre 110 miliardi di dollari nel 2024; in pochi anni dunque numeri duplicati, e ciò – a detta dell'IDC - a causa dell'impiego dell'AI nella trasformazione del digitale, in un contesto di competizione economica mondiale sempre più aggressiva.

"Le aziende adotteranno l'intelligenza artificiale, non solo perché possono, ma perché devono"
e *"avranno la capacità di sintetizzare dati, di imparare, di fornire informazioni su larga scala"*

ha affermato Ritu Jyoti, vicepresidente del programma “Intelligenza artificiale” dell’IDC, un programma pensato non solo con lo scopo di fornire al cliente un miglior risultato, ma un concreto piano di sostegno per la qualità del lavoro: agenti di servizio clienti automatizzati, raccomandazioni e automazione dei processi di vendita, intelligence e prevenzione automatizzate delle minacce e automazione IT.

Secondo le previsioni, i due settori che investiranno di più nell’AI saranno il *Retail* e il *Banking*. Il settore della vendita al dettaglio concentrerà gli investimenti sul miglioramento dell'esperienza del cliente, tramite chatbot e motori di raccomandazione, mentre il settore bancario orienterà la spesa per analisi e indagini sulle frodi, consulenti di programmi e sistemi di raccomandazione. A completare i primi 5 settori di spesa per l’AI nell’anno 2020 saranno la Produzione discreta, la Produzione di processo e la Sanità. In accordo alle previsioni 2020-2024 invece, le sezioni con un maggior livello di spesa in relazione all'AI saranno i Media, il Governo e i Servizi professionali.

Software e servizi rappresenteranno ciascuno poco più di un terzo di tutta la spesa per il 2020, mentre l'hardware fornirà il resto. La quota maggiore della spesa per il software sarà riservata alle applicazioni AI (14,1 miliardi di dollari), mentre la categoria più ampia di spesa per i servizi sarà quella dei servizi IT (14,5 miliardi di dollari). I server (con un costo pari a \$ 11,2 miliardi) copriranno la spesa per l'hardware. Sempre per il periodo 2020/2024, il software vedrà una crescita di spesa più rapida, con un CAGR quinquennale del 22,5%.



Source: IDC Worldwide Artificial Intelligence Spending Guide - Forecast 2020 | Aug (V2 2020)

Figura 1: Principale Impatto dell'AI nei Vari Settori. (Fonte: IDC - Worldwide Artificial Intelligence Spending Guide -)

2.1.2 ... all'indagine macroeconomica

A livello macroeconomico risulta evidente come le nazioni tecnologicamente (ed economicamente) più avanzate si siano lanciate in una sfrenata corsa allo sviluppo dell'AI.

Come ha ammesso qualche anno fa il Presidente russo Vladimir Putin al National Knowledge Day di fronte a migliaia di studenti attoniti “*L'intelligenza artificiale non riguarda unicamente il futuro della Russia, ma quello di tutto il mondo [...] Chi diventerà leader in questo settore sarà il padrone del mondo*”; una verità che numerosi Paesi sembrano aver compreso, come dimostrano gli ambiziosi progetti ed i crescenti investimenti che i Governi di tutto il mondo stanno freneticamente approvando.

A guidare la corsa allo sviluppo dell'intelligenza artificiale sono Cina e USA, entrambi animati dal desiderio di assumere entro il 2030 la leadership mondiale nel settore: mentre il Consiglio di Stato della Repubblica Popolare Cinese pubblica il “*New Generation Artificial Intelligence Development Plan*” – un piano dettagliato che si compone di tre fasi quinquennali e che delinea le strategie da seguire e gli obiettivi da raggiungere - e investe 13,8 miliardi di yuan nella costruzione di un parco tecnologico dell'AI di 54,87 ettari a ovest di Pechino, nel quartiere periferico di Mentougou, nel 2018 la *Defense Advanced Research Projects Agency* dichiara l'AI priorità nazionale e annuncia, con il programma *AI Next*, un investimento di 2 miliardi di dollari a sostegno dell'intelligenza artificiale per i cinque anni successivi.

Anche la Commissione Europea spinge sull'acceleratore: nel biennio 2018-2020 investe fino a 1,5 miliardi di euro nell'AI, portando gli investimenti complessivi ad almeno 20 miliardi di euro entro la fine del 2020 – in accordo con quanto previsto dal programma quadro per la ricerca e l'innovazione *Horizon 2020* – e prevedendo un'ulteriore mobilitazione di 2,5 miliardi di euro da parte di partenariati pubblico-privati; a questi, andrebbero ad aggiungersi i 500 milioni di euro che si stima vengano liberati dal *Fondo europeo per gli investimenti strategici* entro il 2020. Attraverso il potenziamento dei centri di ricerca esistenti e combinando lo sviluppo mirato dell'AI in determinati settori chiave ad una politica di maggiore collaborazione tra i vari Paesi dell'Unione, l'Europa intende migliorare la propria competitività nel settore.

All'interno dell'Unione Europea Francia e Germania dominano la scena: il Governo Federale tedesco lancia nel 2018 il progetto “*AI made in Germany*” e si impegna a fornire circa 3 miliardi di euro per il periodo 2019-2025, preparandosi al contempo alla creazione di 100 nuove cattedre nel campo dell'AI e all'attuazione di un vasto programma di riqualificazione professionale; nello stesso anno Parigi presenta la strategia francese “*AI for humanity*”: investendo 1,5 miliardi di euro – di cui 700 milioni solamente nella ricerca – nello sviluppo

dell'intelligenza artificiale e promuovendo la ricerca interdisciplinare dell'AI a livello nazionale attraverso la formazione di una rete di istituti pubblici di istruzione superiore (Università di Parigi, Tolosa, Grenoble e Nizza) coordinata dall'*Institut national de recherche en informatique et en automatique (INRIA)*, la Francia ambisce a consolidare la propria posizione, per diventare leader nel settore.

Al di fuori dell'Eurozona, poi, degni di nota sono gli sforzi compiuti dal Regno Unito nello sviluppo dell'intelligenza artificiale: in accordo con quanto contenuto nell'*AI Sector Deal*, il Governo britannico avrebbe stanziato quasi 1 miliardo di sterline per l'attuazione dell'Accordo settoriale sull'Intelligenza Artificiale, che andrebbe a sommarsi a 1,7 miliardi di sterline provenienti dall'Industrial Strategy Challenge Fund. Nel programma britannico c'è l'obiettivo di creare, da qui a cinque anni, 16 centri di formazione dottorale ed istituire 1000 nuovi dottorati; così come l'intento di aumentare il numero di visti per talenti eccezionali (LIV 1) – fino a 2000 l'anno – per attirare le menti più brillanti, nonché l'istituzione di nuove regole sull'immigrazione, offrendo a scienziati e ricercatori stranieri dotati di visto per talenti eccezionali la possibilità di richiedere l'insediamento accelerato dopo 3 anni.

A sua volta la Corea del Sud dichiara di essere determinata a investire 187 milioni di dollari tra il 2020 e il 2024, impegnandosi a formare più di 5.000 nuovi ingegneri di alta qualità e concentrando i finanziamenti in AI in settori quali sanità, sicurezza pubblica e difesa nazionale. Passi innovativi sono stati compiuti anche dagli Emirati Arabi Uniti, che annunciano la "*UAE Strategy for Artificial Intelligence*": istituendo per la prima volta nella storia un Ministero per l'Intelligenza Artificiale, canalizzando i fondi verso 9 settori strategici (sanità, trasporti, acqua, energia rinnovabile, spazio, educazione, tecnologia, traffico e ambiente) e investendo ingenti somme nell'AI, il Consiglio Federale Supremo punta a realizzare ambiziosi obiettivi; è infatti prevista entro il 2030 l'introduzione di robot poliziotti e sistemi di trasporto completamente automatizzati.

Ancora - per quanto riguarda il Governo russo - con il fondo sovrano *Russian Direct Investment Fund* (RDIF) Mosca ha costituito importanti partnership con società leader del settore, arrivando a raccogliere più di 2 miliardi di dollari a sostegno di aziende e start-up russe operanti nel campo dell'intelligenza artificiale; come scopo, entro il 2030, quello di fornire una copertura completa a fronte di un'importante carenza di specialisti nel settore, attirando perfino massimi esperti stranieri.

Come si colloca invece l'Italia nella feroce competizione per il dominio dell'AI?

Da quanto si legge dal documento pubblicato lo scorso 2 luglio dal Ministero dello Sviluppo Economico - *“Proposte per una strategia italiana per l’Intelligenza artificiale”* - l’Italia prevede di investire nei prossimi 5 anni 800 milioni di euro in settori specifici, indicati dal documento stesso, quali IoT, manifattura e robotica; servizi, trasporti, agrifood ed energia; aereospazio e difesa; servizi pubblici, cultura e *digital humanities*.

Tra le proposte di spicco avanzate dalle 82 raccomandazioni contenute nella strategia nazionale - elaborata da un team di 30 esperti indipendenti nominati dal MiSE – la creazione di un Istituto Italiano per l’Intelligenza Artificiale (I3A), l’up-skilling e re-skilling della forza lavoro e la formazione di figure altamente specializzate, l’assunzione di un ruolo centrale e direzionale da parte della Pubblica Amministrazione e l’istituzione di una Cabina di Regia interministeriale dotata di poteri di vigilanza e controllo. Proposte e raccomandazioni che ben concorrono nel quadro della competitività internazionale, rivolte alla creazione di un’intelligenza artificiale affidabile e antropocentrica, nel rispetto di linee guida etiche nazionali.

2.2 Il valore strategico dell'Intelligenza Artificiale

Gli statisti delle diverse nazioni precedentemente esaminate non sarebbero gli unici ad aver compreso l'importanza che già oggi – e ancora di più nel futuro – assume il controllo e l'uso dell'intelligenza artificiale; ad aver intuito il valore strategico e le potenzialità offerte dall'AI sono anche manager visionari e colossi multinazionali: mentre Google annuncia l'acquisto di DeepMind– società britannica esperta di machine learning e reti neurali - per la stratosferica cifra di 500 milioni di dollari agli inizi del 2014, Tesla rileva nell'ottobre del 2019 DeepScale - una start-up americana specializzata nello sviluppo di soluzioni per la visione artificiale - avanzando così sempre più verso la realizzazione di autoveicoli dotati di sistemi di guida autonoma di livello 4; intanto, Yahoo! acquista la start-up di photo-analysis LookFlow, Facebook nomina Yann LeCun – professore di Informatica e Neuroscienze della New York University – primo direttore del Facebook AI Research ed Amazon festeggia la vendita di oltre 100 milioni di dispositivi Alexa in tutto il mondo.

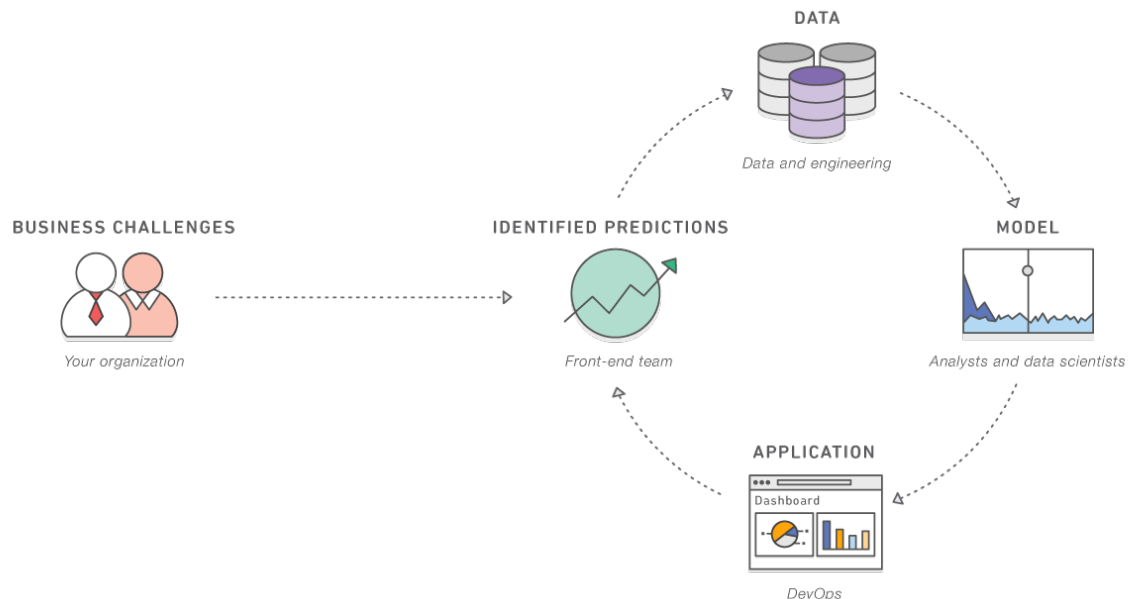
Ma in che modo l'intelligenza artificiale offre un vantaggio competitivo ineguagliabile alle aziende che la utilizzano?

Si riportano di seguito due casi di studio, per osservare come l'intelligenza artificiale venga impiegata in due differenti contesti/settori e comprendere in che modo l'AI crei (in entrambi i casi) valore per l'impresa che la utilizza, rivelandosi così un fattore strategico determinante per il successo di entrambe le aziende analizzate.

2.2.1 Primo caso di studio: Amazon.com

Amazon.com Inc. è la più grande azienda di commercio elettronico e la più vasta internet company al mondo; fondata da Jeff Bezos nel Luglio 1994, ad oggi Amazon vanta un fatturato di 280,5 miliardi di dollari e 840.000 dipendenti sparsi in tutto il mondo. Da sempre impegnata a stravolgere i processi aziendali e l'intero mercato del lavoro, tutt'oggi Amazon risulta essere una delle principali multinazionali che hanno fatto dell'intelligenza artificiale e dell'automazione l'elemento strategico determinante nella costruzione del proprio successo.

In particolare, Amazon utilizza il machine learning per leggere ed elaborare tempestivamente i dati raccolti nello svolgimento delle proprie attività, creare previsioni e implementare successivamente processi aziendali efficaci²⁵, ottimizzando in tal modo funzioni ed eventi incerti: partendo dall'analisi dei dati storici (come ad esempio vendite passate, quantità di prodotti venduti o rimanenze di magazzino) e sviluppando modelli di apprendimento coerenti ai suddetti dati Amazon è in grado di elaborare modelli automatici di analisi di scenari futuri e applicazioni aziendali.²⁶



²⁵ Amazon.com (2020). *Cos'è L'Intelligenza Artificiale*. <https://aws.amazon.com/it/machine-learning/what-is-ai/>

²⁶ Amazon.com (2020). *Cos'è L'Intelligenza Artificiale - Apprendimento Automatico e Apprendimento Approfondito* -. <https://aws.amazon.com/it/machine-learning/what-is-ai/>

È importante evidenziare come in Amazon l'intelligenza artificiale non venga utilizzata unicamente per analizzare scenari futuri, ridurre i costi operativi e, più in generale, per rendere più efficienti i processi aziendali; l'internet company infatti si servirebbe dell'intelligenza artificiale (e in particolare del machine learning) per:

- *Rilevare, gestire ed analizzare* anomalie e/o eventi non conformi a specifici pattern predeterminati;
- *Osservare e valutare analisi fraudolente*: sofisticati algoritmi utilizzano i dati raccolti per sviluppare modelli predittivi attraverso cui è possibile identificare potenziali venditori fraudolenti e/o recensioni non veritiere;
- *Prevedere ed analizzare il tasso di abbandono dei clienti*: l'intelligenza artificiale identifica i clienti potenzialmente più inclini all'abbandono del servizio e mette in atto azioni e strategie volte alla fidelizzazione degli stessi (promozioni, offerte ecc.);
- *Personalizzare i contenuti mostrati*: l'analisi predittiva ed il machine learning permettono ad Amazon di individuare i prodotti che potrebbero interessare al cliente e ne suggeriscono l'acquisto.

Particolarmente rilevante è il contributo apportato dall'apprendimento approfondito all'interno della compagnia; grazie al machine learning approfondito Amazon è infatti in grado di stimare, interpretare e fornire previsioni ed interpretazioni di dati estremamente complessi quali:

- *Dati relativi alla classificazione e categorizzazione di immagini/video*: le reti neurali permettono di identificare e categorizzare pattern ed oggetti in immagini e video. Esempi di tali applicazioni di AI si osservano nei servizi Amazon Rekognition, Amazon Prime Photos e Firefly;
- *Dati relativi al riconoscimento vocale*: l'apprendimento approfondito e il machine learning permettono a dispositivi come Alexa di riconoscere e comprendere il linguaggio umano e rispondere efficacemente alle richieste avanzate dal proprio interlocutore;
- *Dati riguardanti il riconoscimento del linguaggio naturale*: registrando e analizzando le diverse tonalità di voce, gli stessi device sono in grado di comprendere le differenti emozioni del proprio interlocutore e/o concetti come il sarcasmo, l'umorismo ecc..

L'introduzione dell'intelligenza artificiale all'interno delle varie business units ha permesso ad Amazon sia di risparmiare tempo e costi di produzione che di massimizzare gli output e migliorare il controllo di gestione dei vari processi aziendali, offrendo senza dubbio all'azienda un importante vantaggio competitivo sui propri competitors²⁷.

Un altro importante ambito di applicazione dell'intelligenza artificiale in Amazon è certamente la logistica: analizzando i dati d'acquisto e sfruttando le potenzialità del machine learning, Amazon è in grado di prevedere e stimare la futura domanda di beni da parte dei clienti ed organizzare efficientemente la catena di produzione e distribuzione, permettendo alla compagnia di ridurre sensibilmente i propri costi di magazzino e logistica.



Figura 2: Un umano ed un Robot al lavoro in Centro Amazon nel New Jersey (Fonte: Amazon)

²⁷ AI4Business.it. (2020). *Amazon e l'intelligenza artificiale, «Così risolviamo i problemi meglio e prima del tempo», la promessa di Bezos*. <https://www.ai4business.it/intelligenza-artificiale/amazon-intelligenza-artificiale-machine-learning/>

2.2.2 Secondo caso di studio: Tesla

Fondata a San Carlos nel 2003 da Martin Eberhard e Marc Tarpenning, Tesla Inc. è un'azienda statunitense specializzata nella produzione di pannelli fotovoltaici, auto elettriche e sistemi di stoccaggio energetico con sede a Palo Alto, California. Da piccola nicchia di mercato qual era, ad oggi Tesla Inc rappresenta un'importante realtà industriale, vantando un fatturato di 24,58 miliardi di dollari e 45.000 dipendenti distribuiti in ecologiche *gigafactory* ultramoderne; nel Marzo del 2020 la società ha raggiunto l'importante traguardo di 1 milione di auto elettriche prodotte, confermandosi indiscutibilmente come prima casa automobilista produttrice di veicoli elettrici.

Ma in che modo l'Intelligenza artificiale ha contribuito alla costruzione dell'incredibile successo di Tesla Inc.?

Le continue novità e l'alto tasso di innovazione che si ritrovano nei prodotti Tesla hanno permesso all'azienda di Palo Alto di costruire nel tempo un marchio forte e competitivo: nell'Agosto del 2015 Forbes eleggeva Tesla Inc l'azienda più innovativa al mondo e solamente cinque anni più tardi, nel luglio di quest'anno - esattamente dieci anno dopo la sua prima quotazione in Borsa – il valore di mercato delle azioni della società schizzava a 1.100 dollari, portando così la capitalizzazione di mercato dell'azienda a 207,2 miliardi di dollari e rendendo Tesla la casa automobilistica con il più alto valore al mondo.

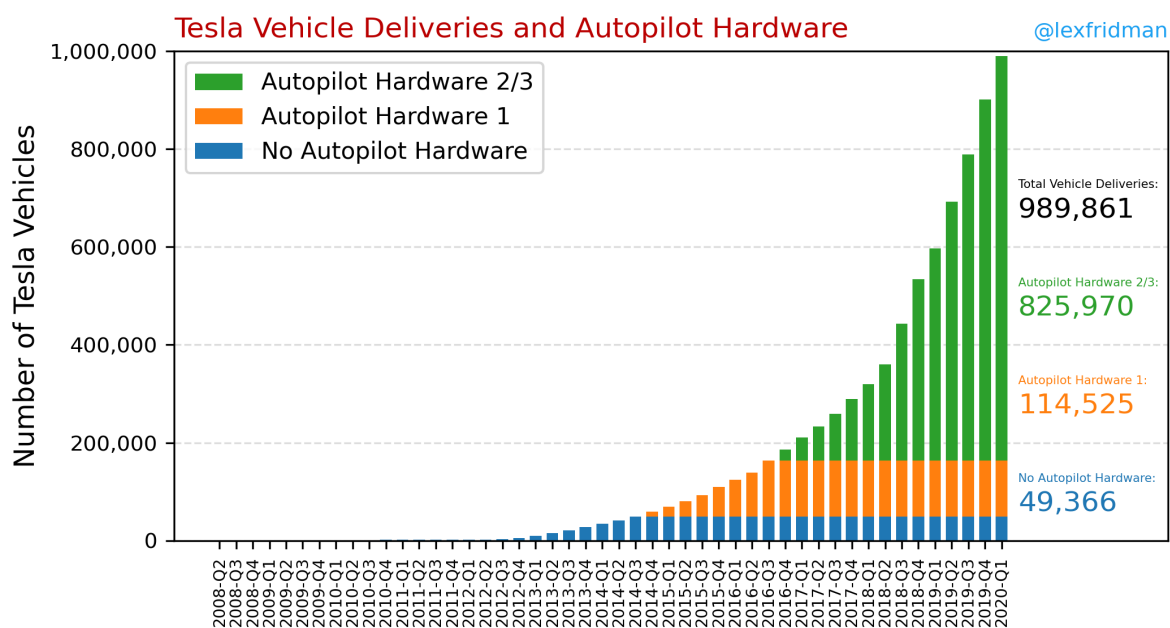


Figura 3: Numero di Autovetture Tesla Acquistate in Rapporto allo Sviluppo dell'Autopilot. (Fonte: @lexfridman, Twitter)

Senza dubbio l'intelligenza artificiale ha rappresentato per la società californiana un fattore strategico insostituibile nella costruzione del proprio successo: come dimostra il report *“Tesla Vehicle Deliveries and Autopilot Mileage Statistics”* pubblicato da Lex Friedman - ricercatore dell'AI presso il Massachusetts Institute of Technology – la richiesta dei veicoli Tesla è stata trainata dal crescente impiego dell'intelligenza artificiale e dai progressi traggurati dai sistemi di guida autonoma dei propri prodotti; se nei primi 6 anni di attività Tesla aveva venduto appena 49.366 vetture prive di autopilota, negli ultimi 4 anni l'azienda di Palo Alto ha registrato 825.970 richieste d'acquisto di modelli dotati di hardware ultima generazione, pari all'83,44% dell'intero autoparco prodotto da Tesla nel corso della sua storia.

L'annuncio del potenziamento dell'Autopilot 3 e la presentazione del sistema Full Self-Driving (FSD) in occasione dell'Autonmy Day ha certamente rappresentato un forte stimolo all'acquisto dei veicoli Tesla; implementato su tutte le Model S, Model X e Model 3 prodotte a partire da aprile 2019, il sistema Full Self-Driving rappresenta un ulteriore passo avanti compiuto da Tesla verso la commercializzazione di veicoli a totale guida autonoma.

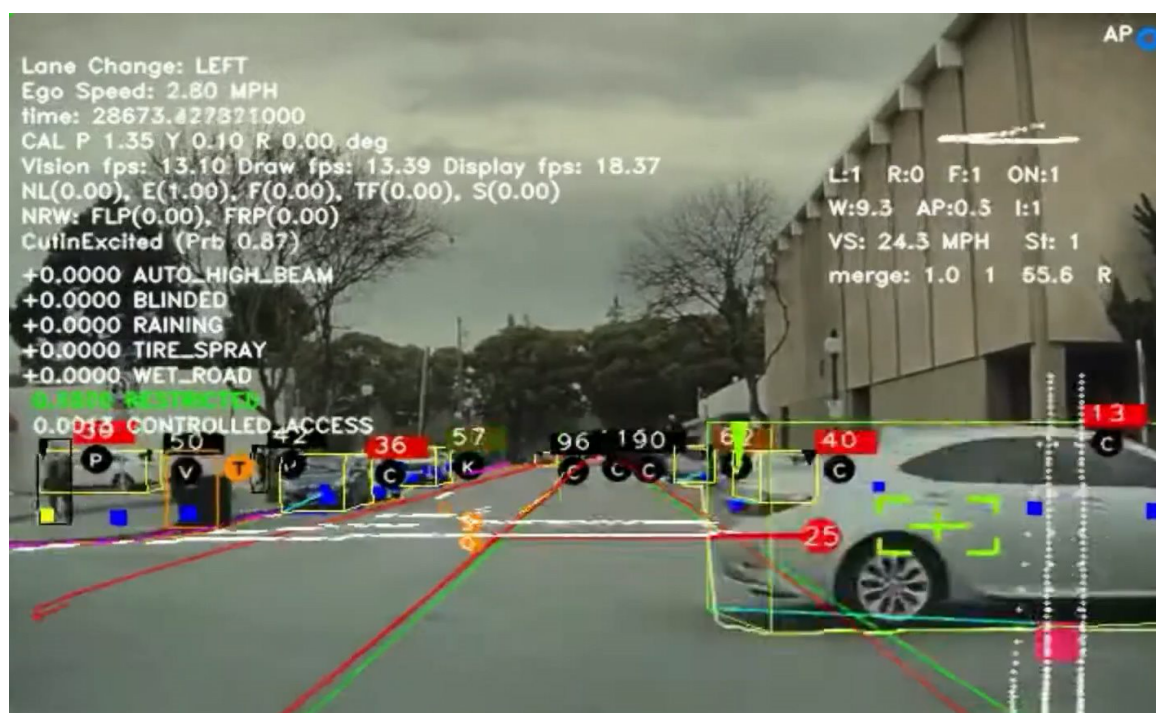


Figura 4: Come una Tesla vede il mondo: L'AI permette al veicolo di apprendere e catalogare i fenomeni circostanti (Fonte: Tesla)

Basato su un nuovo microprocessore - internamente sviluppato da Tesla e prodotto da Samsung - dotato di 6 miliardi di transistor e 21 volte più potente del precedente NVIDIA Drive Xavier, FSD opera attraverso due distinti system-on-a-chip²⁸ che gestiscono due reti neurali separate secondo la logica della ridondanza: elaborando separatamente due volte le informazioni raccolte da 8 telecamere a 360°, radar frontale e sensori ultrasonici a corto raggio, il sistema è in grado di gestire l'auto negli svincoli autostradali e cambiare corsia in totale autonomia, imparando e perfezionandosi dopo ogni viaggio.

Lo sviluppo del Full Self Driving system, però, non sarebbe altro che il punto di partenza per la realizzazione di sistemi di guida autonomi ancora più sofisticati; in accordo con quanto dichiarato da Elon Musk, infatti, Tesla si pone l'ambizioso obiettivo di lanciare entro la fine del 2020 un servizio di taxi completamente autonomi. La vision di Tesla sarebbe a dir poco futuristica: la creazione di una flotta di taxi a guida autonoma (composta da Model S, Model X e Model 3) integrata e facilmente prenotabile tramite app, che permetta uno spostamento sicuro ed efficiente all'interno della città.

Tesla Autopilot Safety

Million miles between accidents.

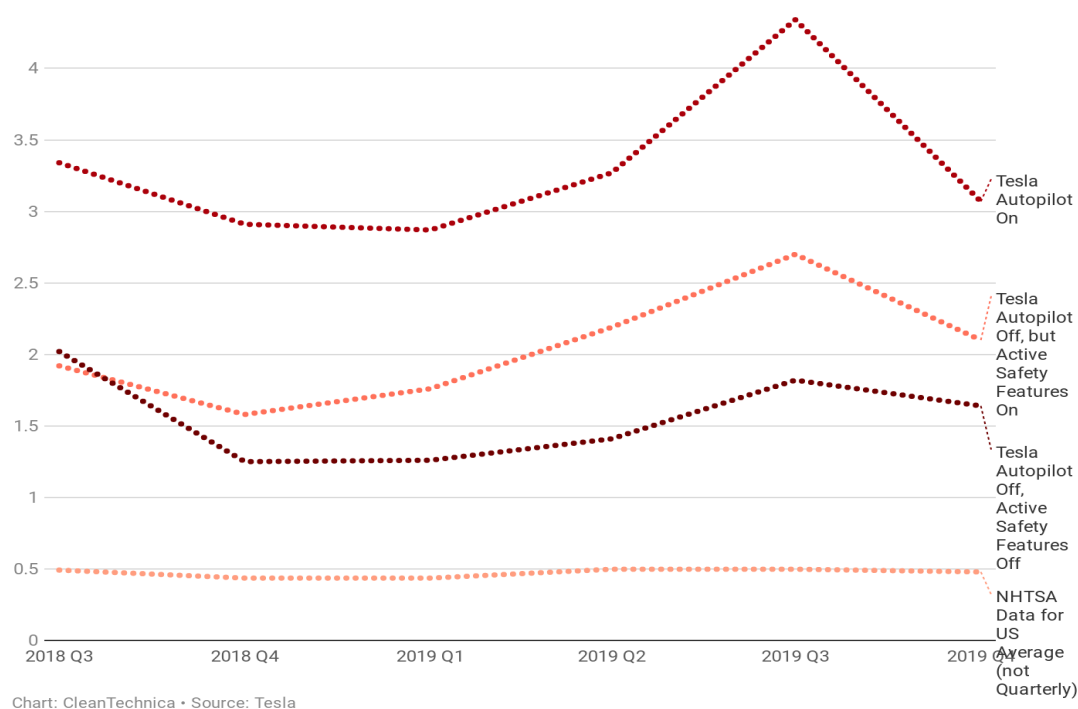


Figura 5: Sicurezza dell'Autopilot Tesla. (Fonte: Tesla)

²⁸ Wikipedia. (2020). *System-on-a-chip*. <https://it.wikipedia.org/wiki/System-on-a-chip>

Insomma, come dimostra il grafico sopra riportato, l'efficienza e la sicurezza generate dall'utilizzo dell'AI alla guida rappresentano i cardini che hanno permesso a Tesla di realizzare il proprio successo, costruendo in breve tempo un marchio forte e competitivo e diventando così un'azienda leader nel settore dell'industria automobilistica.

2.3 Aspetti psicologici e culturali: il Rapporto BVA DOXA

Come emerso chiaramente a Shanghai durante la World Artificial Intelligence Conference 2019, affascinati e respinti dalle infinite possibilità che ci promette, l'intelligenza artificiale accende in noi sentimenti contrastanti, generando così una vera e propria scissione all'interno del mondo politico, accademico e imprenditoriale: se da un lato uomini d'affari del calibro di Jack Ma, fondatore e presidente di Alibaba Group, sono convinti che l'intelligenza artificiale non farà altro che *“aprire un nuovo capitolo del mondo in cui le persone saranno in grado di comprendere meglio loro stesse”*, dall'altro lato c'è chi, come Elon Musk, teme che l'AI possa sfuggire al nostro stesso controllo, rappresentando una minaccia per l'intero genere umano.

Il rapporto BVA DOXA del 2019 evidenzia chiaramente l'esistenza di questa spaccatura: intervistando un campione di 30.890 individui provenienti da oltre 40 paesi il sondaggio realizzato da BVA DOXA tra l'ottobre del 2018 e il gennaio del 2019 - in collaborazione con WIN- fotografa inequivocabilmente l'ambiguo e differente *sentiment* che l'intelligenza artificiale suscita.

Il 52% del campione intervistato ha infatti dichiarato di non temere che l'intelligenza artificiale possa causare nei prossimi 10 anni la perdita del proprio posto di lavoro, contro il 48% che invece ha ammesso di sentirsi minacciato dallo sviluppo dell'AI.

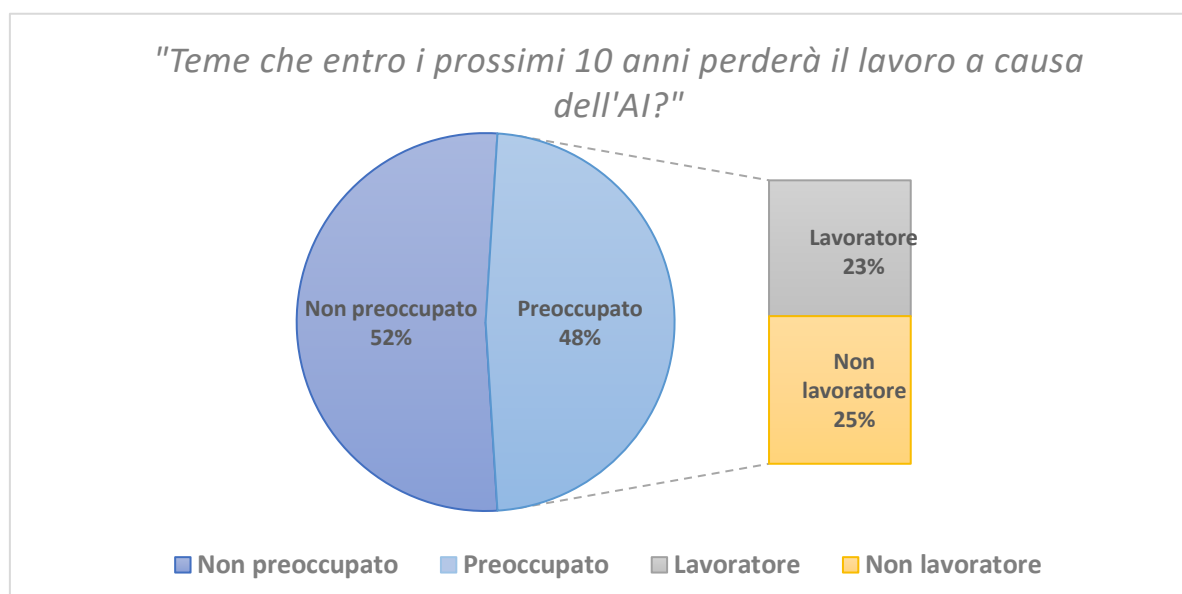


Figura 6: Questionario al Pubblico circa l'AI e l'automazione

Particolarmente rilevante sarebbe poi la differente sensazione manifestata dalle varie categorie di soggetti intervistati: in accordo con quanto emerso dalla ricerca, a temere maggiormente l'intelligenza artificiale sono le donne (48%) ed i soggetti di età inferiore a 35 anni (51%) e superiore a 54 anni (45%); gli uomini (57%), i lavoratori altamente qualificati (67%), gli occupati a tempo pieno (70%) ed i soggetti di età compresa tra i 35 e i 54 anni (59%) rappresenterebbero invece le categorie di individui meno spaventate dal progresso tecnologico e dallo sviluppo dell'intelligenza artificiale.

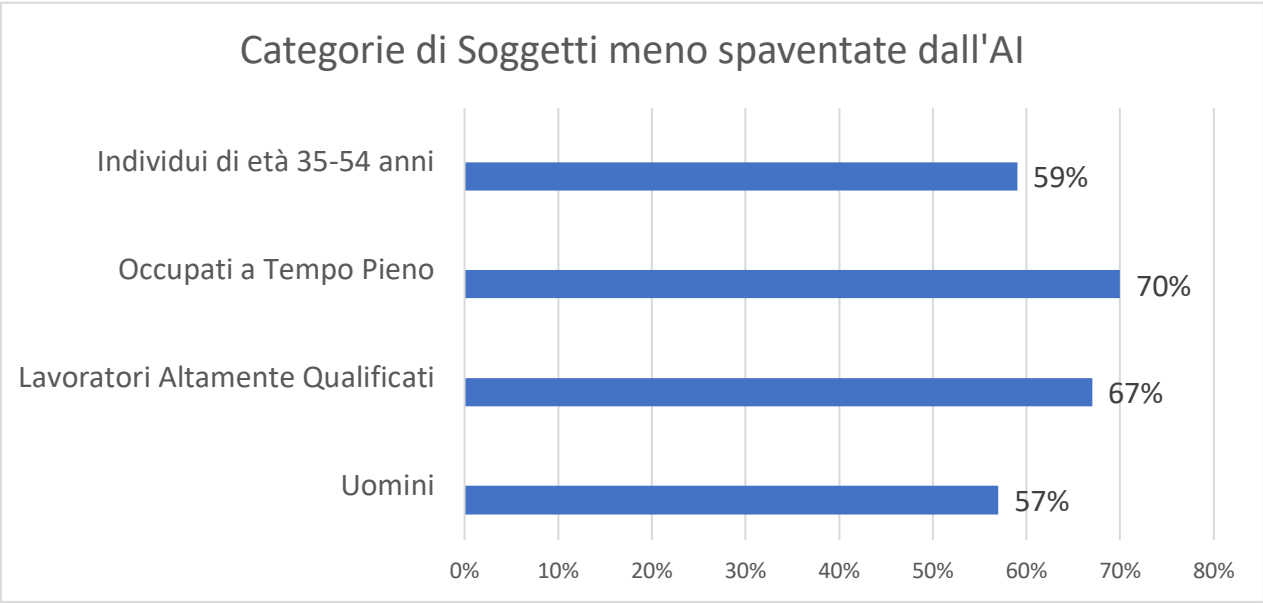


Figura 7: Categorie di Soggetti meno spaventate dallo sviluppo dell'Intelligenza Artificiale

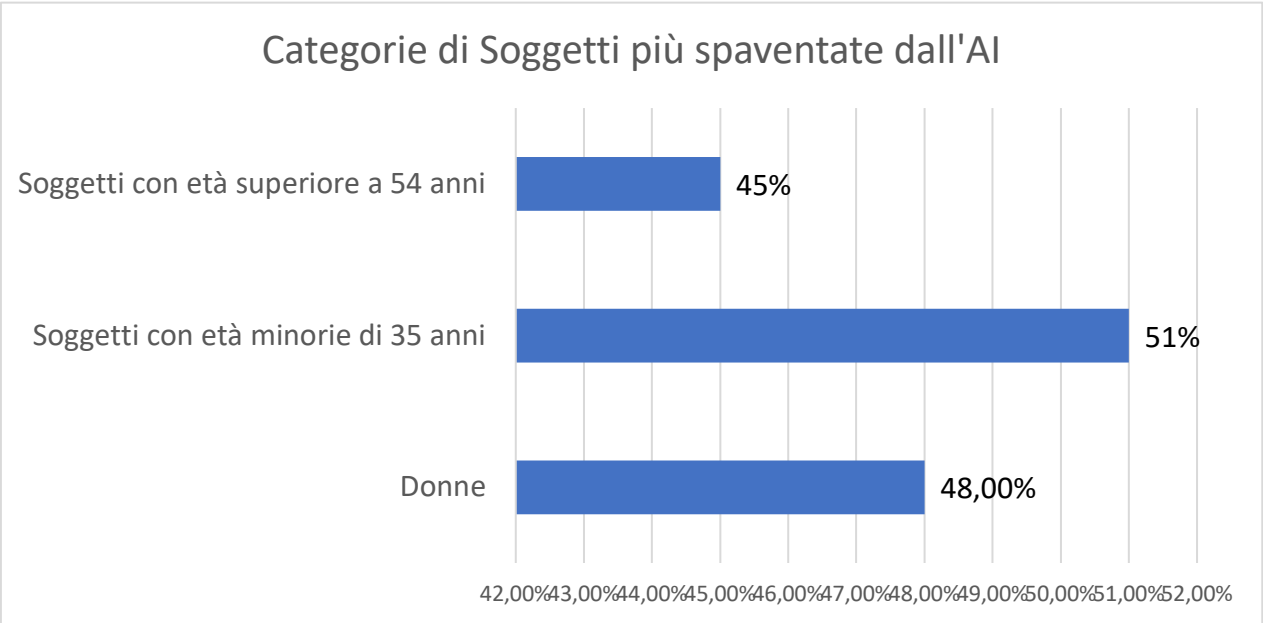


Figura 8: Categorie di Soggetti più Spaventate dallo Sviluppo dell'AI

Degno di nota è anche il divario registrato tra Paesi industrializzati e quelli in via di sviluppo: se in Germania quasi otto lavoratori su dieci dichiarano di non temere tagli e licenziamenti conseguenti all'uso dell'AI, nelle economie meno avanzate il numero di lavoratori fiduciosi diminuisce sensibilmente, registrando minimi in paesi come India (3,7 lavoratori su 10), Argentina (3,5 lavoratori su 10) e Malesia (4 lavoratori su 10).

Come mostra chiaramente il Rapporto DOXA, a dubitare maggiormente dell'intelligenza artificiale sarebbero dunque i lavoratori appartenenti a paesi meno sviluppati.

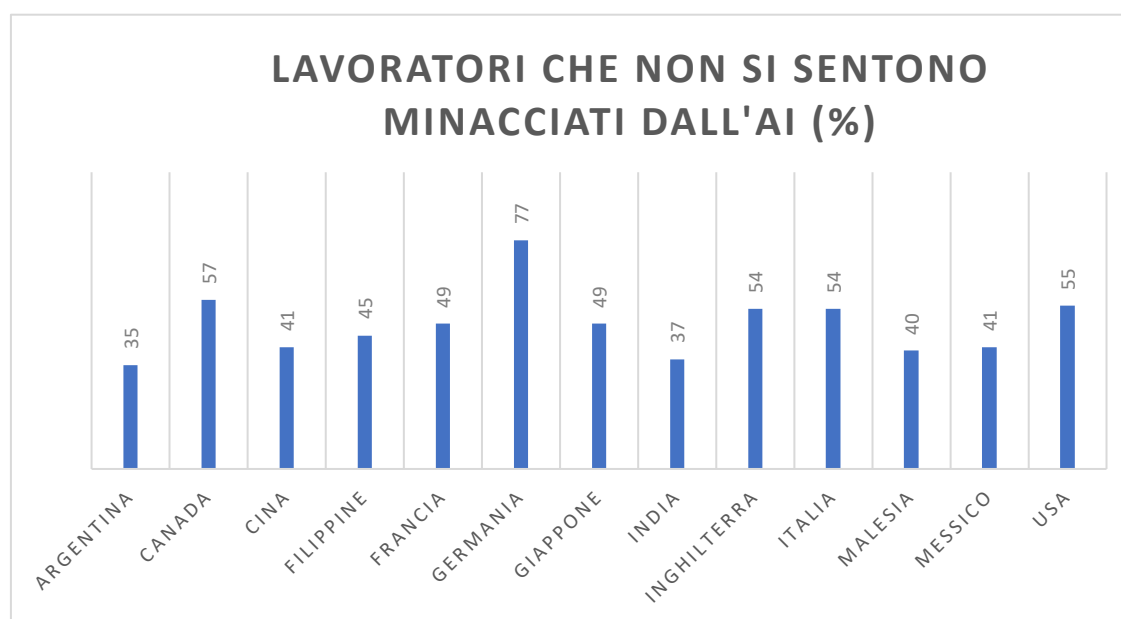


Figura 9: Differenti Sentiments di Lavoratori su base Geografica

Contrariamente agli altri quesiti sottoposti, la domanda “*Quanto sarebbe favorevole al fatto che l’IA sostituisca le attività svolte dal personale medico?*” ha registrato risposte concordi (e negative) da parte delle differenti categorie di soggetti intervistate da BVA DOXA: solo il 7% degli intervistati si è dichiarato favorevole a una totale sostituzione del personale medico da parte dell’AI. Forti resistenze sono state registrate in Italia e Germania e, in generale, nei paesi membri del G7 – con l’eccezione del Giappone, in cui solo il 12% degli intervistati ha manifestato un netto rifiuto alla possibile sostituzione del personale medico da parte dell’intelligenza artificiale – mentre in paesi quali India, Libano e Cina si è registrata una controtendenza alla diffusa opinione negativa, dovuta probabilmente alla sfiducia che i cittadini nutrono nei confronti del sistema sanitario nazionale.

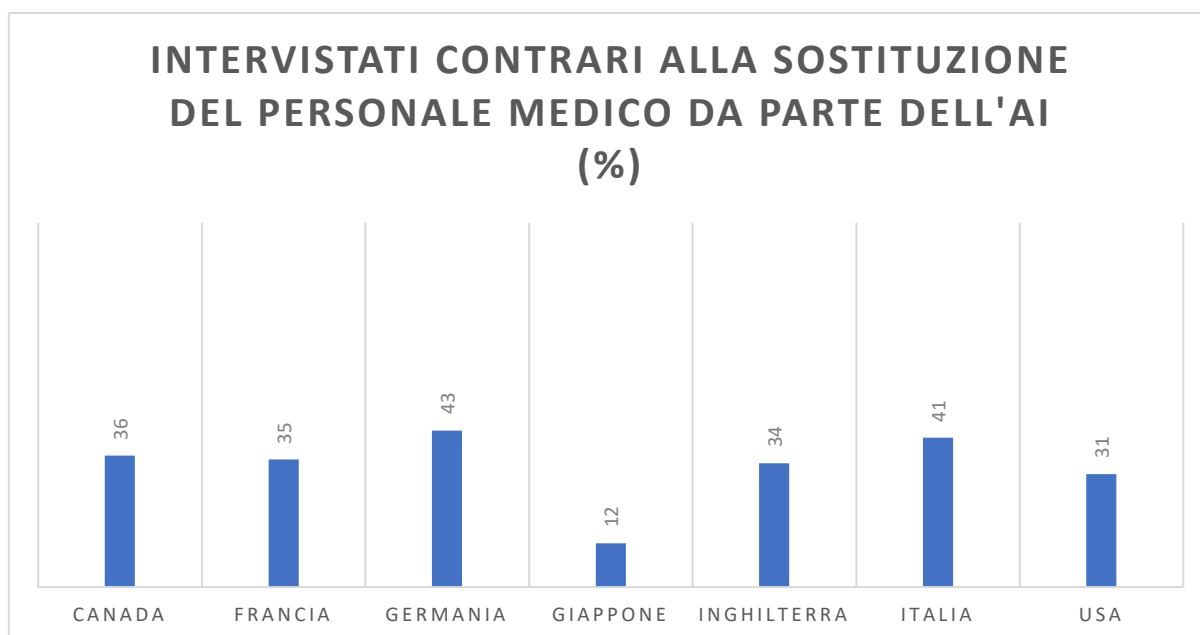


Figura 10: AI e Sanità -Percentuale del Campione avverso all'automazione in ambito healthcare-

Ma qual è l'idea che gli italiani hanno dell'intelligenza artificiale?

Stando a quanto evidenziato dal Rapporto DOXA, l'87% degli italiani avrebbe un'opinione positiva di robot e intelligenza artificiale - di cui il 12% molto positiva – mentre solamente l'8% avrebbe manifestato un giudizio negativo su AI e robot; il 5% degli italiani, invece, non avrebbe alcuna particolare opinione.

Secondo l'89% del campione, poi, robot e intelligenza artificiale non saranno mai in grado di sostituire l'uomo ma ciononostante il 70% degli italiani teme che il progresso tecnologico possa provocare nei prossimi anni la perdita di migliaia di posti di lavoro: in particolare, secondo 50% degli italiani vi sarebbe il concreto rischio che in futuro le macchine conquistino il predominio sull'uomo.

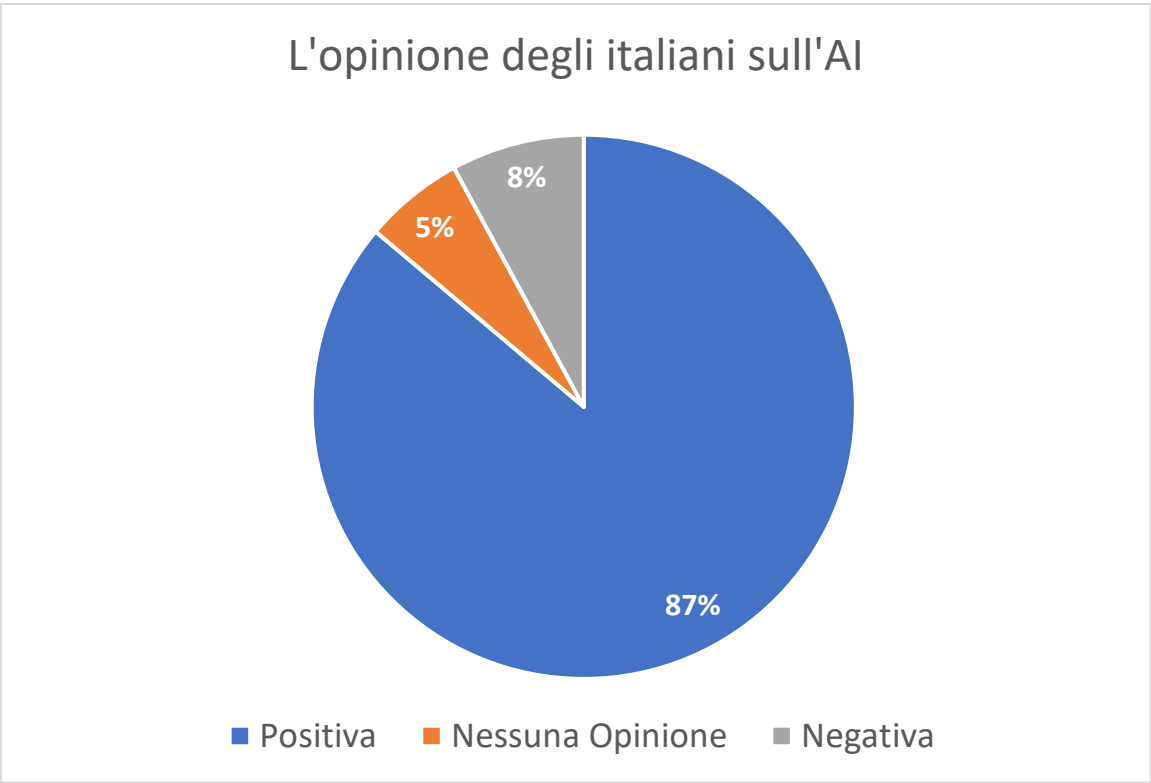


Figura 11: Gli Italiani e le loro opinioni sull'AI.

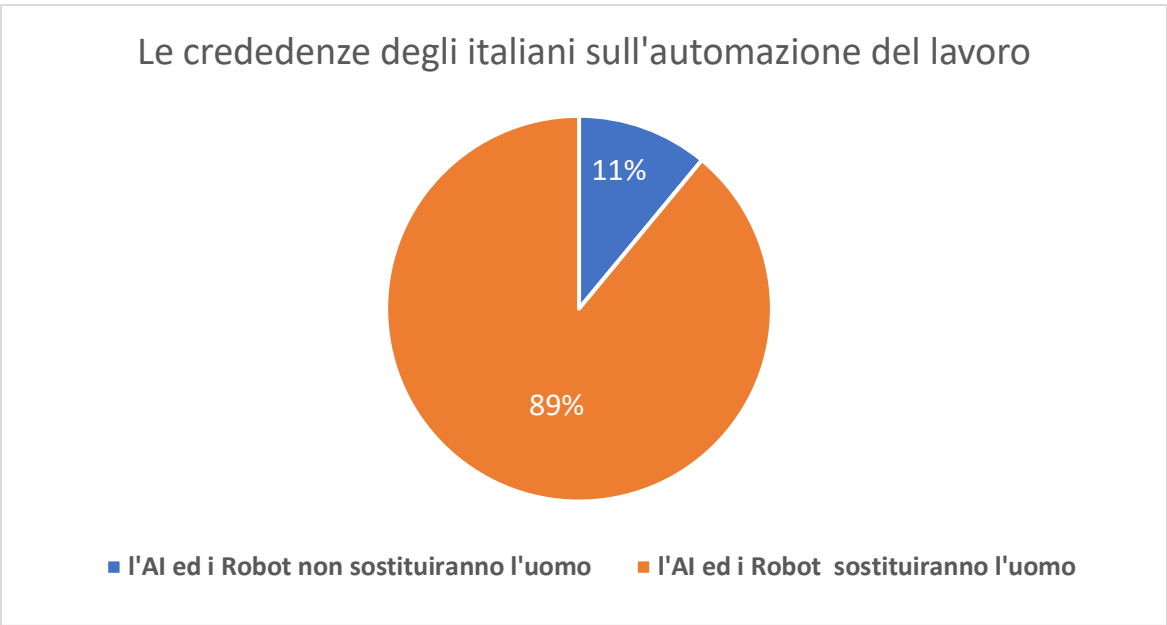


Figura 12: Gli Italiani e le loro credenze sul lavoro futuro.

Degne di nota sono le considerazioni manifestate dai soggetti intervistati: secondo l'87% degli italiani robot e intelligenza artificiale contribuiscono a migliorare il benessere e la qualità della vita; per il 94% del campione robot e intelligenza artificiale hanno permesso all'uomo di trarre in breve tempo risultati e scoperte impensabili, mentre per l'89% degli italiani l'uso delle nuove tecnologie si è rivelato essenziale per lo svolgimento di attività pericolose e/o eccessivamente faticose.

Secondo il 92% del campione, però, è necessario adeguare la normativa ed emanare nuovi regolamenti per cogliere le opportunità offerte dall'AI senza che vi siano rischi o pericoli.

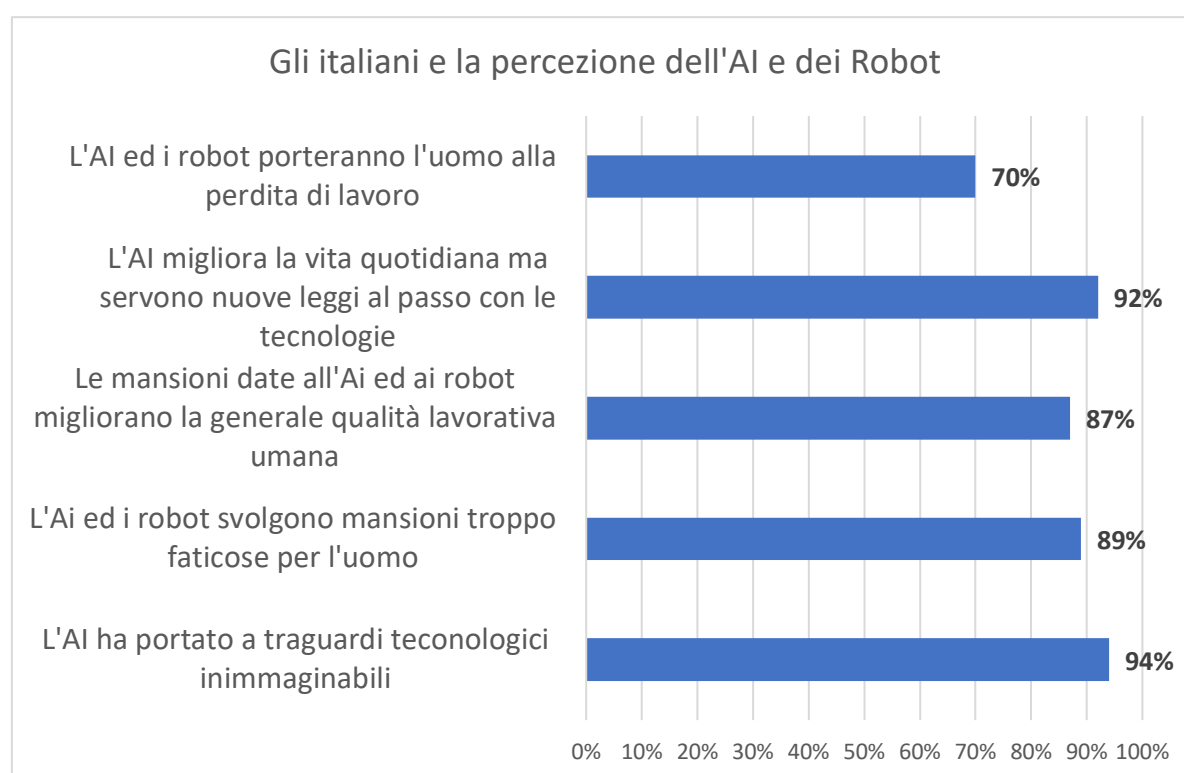


Figura 13: Gli Italiani e la loro Percezione sull'Automazione ed AI.

Per quanto in Italia le nuove tecnologie (il 65%) e l'intelligenza artificiale (89%) – sia pure con in maniera differente - rappresentino un interesse diffuso e trasversale a tutte le categorie di soggetti intervistati, a mostrare maggiore conoscenza e attenzione per la materia sono i giovani maschi di età inferiore a 35 anni, con elevato titolo di studio e residenti nel Sud e nelle Isole.

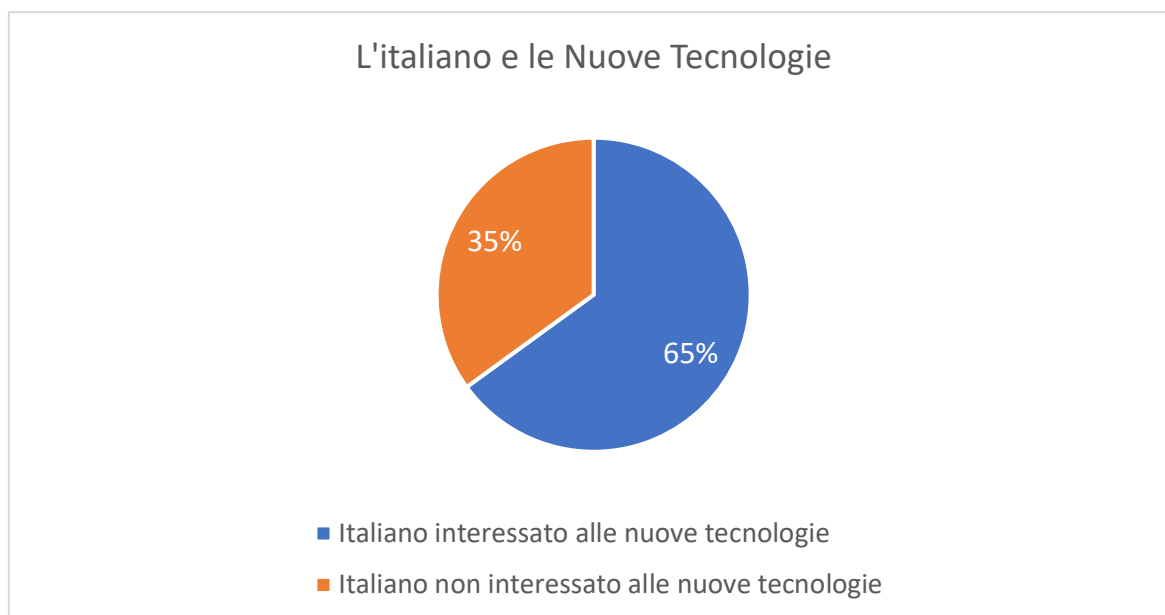


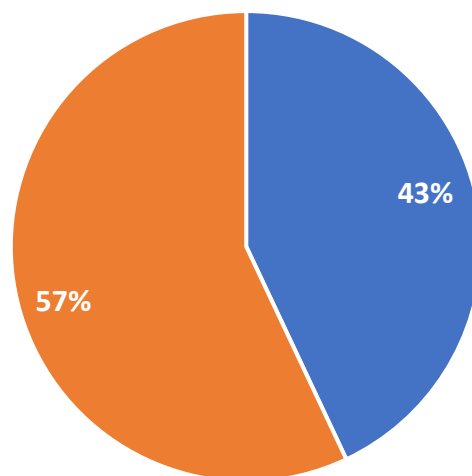
Figura 14: Gli Italiani e l'interesse per le Nuove Tecnologie

Malgrado ciò, l'interazione tra uomo ed AI avrebbe già interessato una buona parte degli individui intervistati, a prescindere dalla categoria di appartenenza dei soggetti coinvolti: il 43% degli italiani ha infatti dichiarato di aver già utilizzato sistemi di robot e intelligenza artificiale al lavoro e/o a casa, mentre il 47% del campione dichiara di aver ricorso a piattaforme e soluzioni basate su robot e IA per l'acquisto di beni e servizi.

Stando a quanto rivelerebbe la ricerca, gli italiani si affiderebbero a robot e intelligenza artificiale principalmente per la pulizia della casa / svolgimento di lavori domestici (21%) e per la logistica e spedizioni (12%); altrettanto importante ma meno diffuso sarebbe poi l'utilizzo dell'AI nell'assistenza virtuale e nelle chatbot (es. call center, acquisti online, colloqui di lavoro, prenotazioni viaggi o visite mediche) e nell'uso di apparecchiature quali computer, smartphone e smart tv (7%).

In ogni caso, come evidenziato da BVA DOXA, ad utilizzare maggiormente le nuove tecnologie sono soprattutto uomini giovani con elevata scolarizzazione e white-collar.

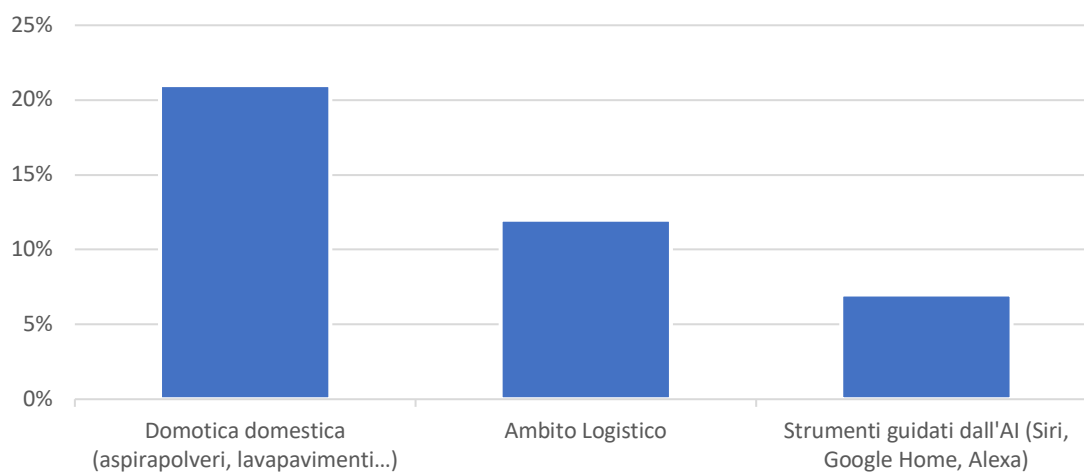
Gli Italiani e l'approccio all'AI ed ai Robot



- Ha già utilizzato sistemi di AI e/o Robot
- Non ha ancora utilizzato sistemi di AI eo/robot

Figura 15: Gli Italiani e L'approccio all'Automazione.

Ambiti dove si utilizza l'AI



- Domotica domestica (aspirapolveri, lavapavimenti...)
- Ambito Logistico
- Strumenti guidati dall'AI (Siri, Google Home, Alexa)

Figura 16: Ambiti di Utilizzo dell'AI.

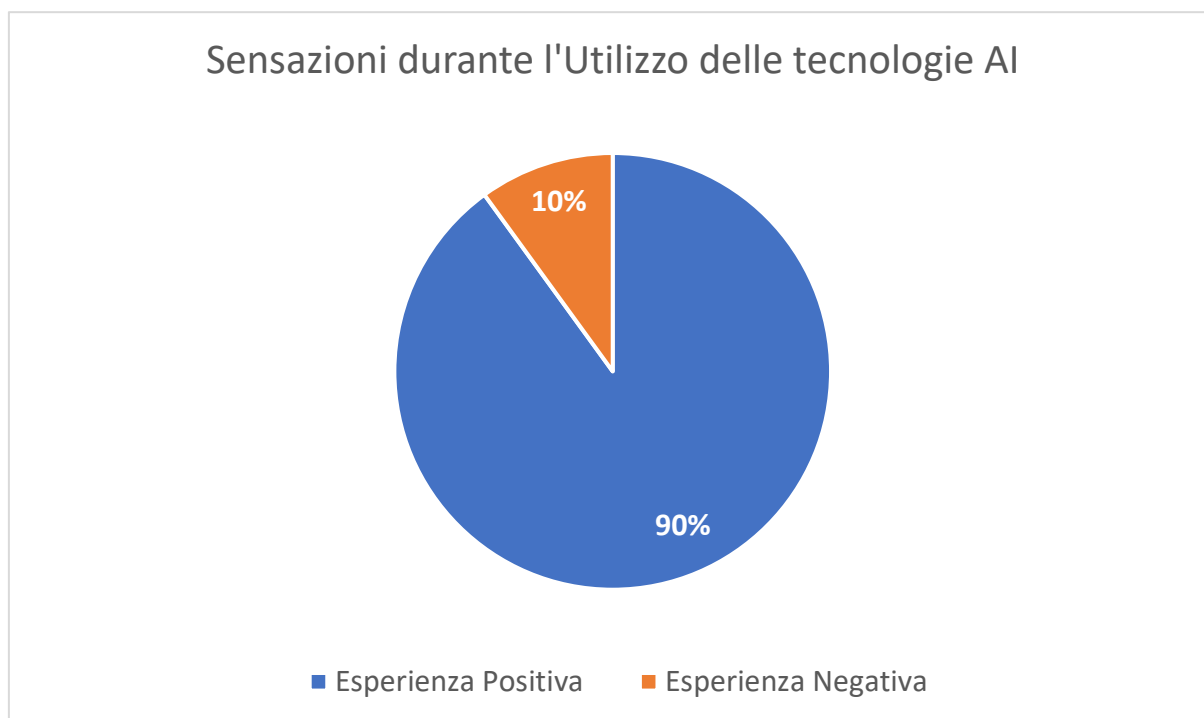


Figura 17: Sentiments provati durante l'utilizzo dell'AI.

Ma cosa ne pensano gli italiani dell'utilizzo di robot e intelligenza artificiale in riferimento ai differenti ambiti di applicazioni e ai vari settori in cui vengono impiegati?

Secondo il 53% del campione intervistato il più importante contributo apportato da robot e intelligenza artificiale si osserverebbe nella logistica e nei trasporti, nel settore manifatturiero e nell'industria (51%) e nella medicina e nei servizi sanitari (50%). Altrettanto importante sarebbe poi, secondo gli italiani, l'impatto positivo che l'utilizzo di robot e AI avrebbe sul settore automobilistico, sul settore militare (48%) e nella sicurezza; benefici più modesti, invece, si registrerebbero nel lavoro domestico (40%) e nei servizi alla persona (32%).

Particolarmente significativo risulta essere il generale dissenso incontrato dall'utilizzo di robot ed intelligenza artificiale per scopi educativi: il 43% degli italiani avrebbe infatti indicato l'istruzione e la scuola come il più importante ambito a cui non andrebbero applicati i robot e l'IA.

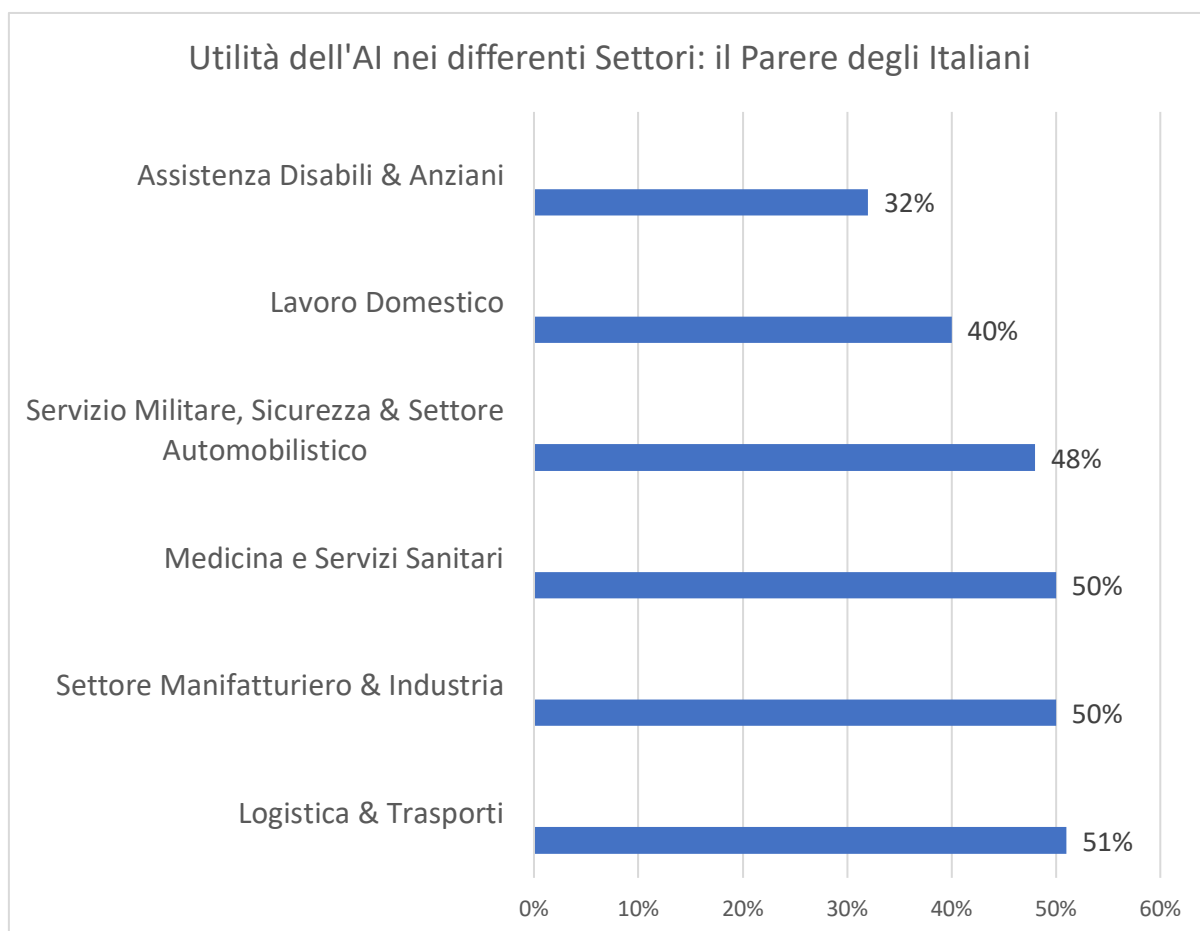


Figura 18: L'AI e la sua Utilità nei differenti Settori secondo gli Italiani.



Figura 19: L'AI ed il Settore dell'Istruzione -il parere degli Italiani-

PARTE III - FUTURO

3.1 Lavoro e ricchezza nell'era dell'AI

“Ciò che sappiamo è una goccia, ciò che ignoriamo è un oceano”, ripeteva continuamente Isaac Newton, con un pizzico di malinconia, circa tre secoli fa.

Con ogni probabilità, nessun'altra frase sarebbe in grado di descrivere con altrettanta efficacia lo scenario di incertezza che si prospetta nell'immaginare in futuro l'impatto dell'AI sulle nostre vite; per quanto tentiamo disperatamente di interpretare il domani, la verità è che brancoliamo nel buio: se, tendenzialmente, è difficile intuire cosa ci riservi il domani, ancor più complesso è prender posizioni nette, radicali, nel discutere dell'AI e delle sue possibili conseguenze; non possiamo far altro che affidarci a stime incerte.

Tanto McKinsey & Company, multinazionale americana di consulenza strategica, quanto PricewaterhouseCoopers, azienda leader nella fornitura di servizi di consulenza direzionale e revisione contabile, prefigurano per l'intelligenza artificiale un impatto economico più che travolgente: si è infatti stimato che il crescente impiego dell'*artificial intelligence* nel mondo del lavoro provocherebbe, entro il 2030, un aumento del 14% del Prodotto Interno Lordo mondiale, corrispondente alla stratosferica cifra di 15.700 miliardi di dollari.

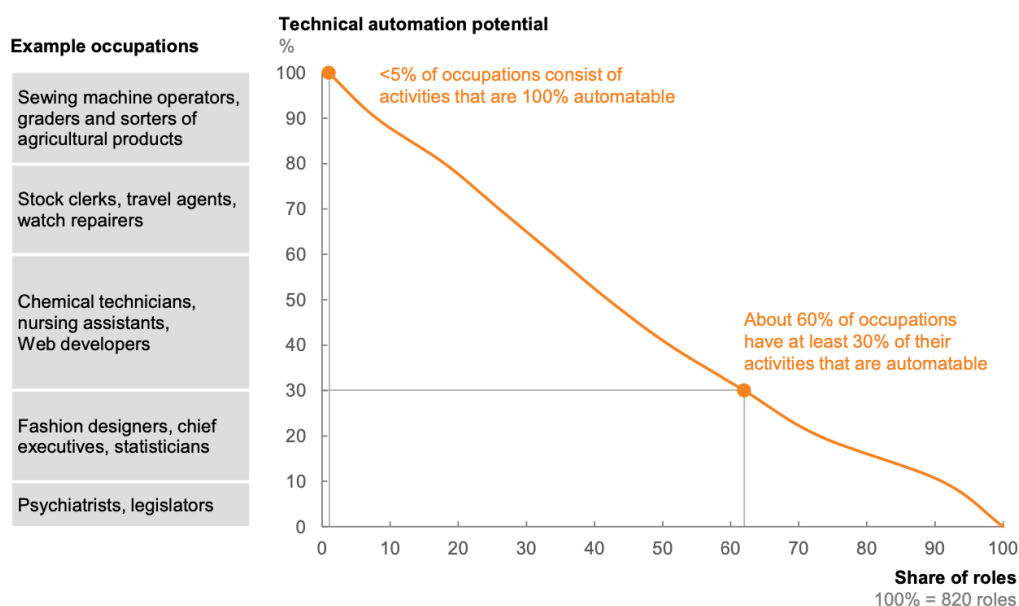
Stando a quanto previsto da McKinsey & Company e PricewaterhouseCoopers, però, non tutti beneficerebbero egualmente dei vantaggi dell'AI: ad assicurarsi la parte più cospicua sarebbero Cina e Stati Uniti, seguiti a debita distanza dall'Europa e dai più tecnologici Paesi asiatici.

È stato infatti stimato che la produttività dell'economia del Dragone, spinta dal frenetico sviluppo dell'AI e dalle sue incredibili applicazioni, cresca tra l'0,8% e l'1,4% del PIL su base annua, garantendo così a Pechino 7 dei 15,7 trilioni di dollari che si prevede siano creati dall'AI entro il 2030. Neppure gli USA sarebbero in grado di competere con la potenza cinese: i 3,7 trilioni di dollari, che si stima andranno ad accrescere il PIL dell'America Settentrionale, non sembrano esser sufficienti a spaventare Pechino e discostare la Cina dal sogno di conquista di una leadership nel settore. Più modesto, invece, sarebbe l'impatto che l'intelligenza artificiale avrebbe nell'economia di Europa e Italia: entro il 2030, infatti, si attende una crescita del 19% del PIL europeo e del 13% del PIL italiano, per un valore rispettivamente pari a 2.700 miliardi di euro e 228 miliardi di euro. I più radicali cambiamenti innescati dallo sviluppo dell'intelligenza artificiale, però, non si configurerebbero nella creazione di valore e nella distribuzione della ricchezza, ma si osserverebbero nell'organizzazione del lavoro e nella

formazione professionale dei lavoratori del futuro. Degne di nota sarebbero le conclusioni e le analisi esposte dal McKinsey Global Institute all'interno del report *"A future that works: automation, employment and productivity"*. Secondo quanto previsto dall'istituto di ricerca aziendale ed economica di McKinsey&Global, infatti, il 49% delle attività attualmente svolte da persone fisiche potrebbero essere completamente automatizzate tra il 2035 e il 2075, a seconda della velocità con cui procederà lo sviluppo dell'intelligenza artificiale. Osservando il mercato del lavoro degli USA e analizzando più di 2000 mansioni di oltre 800 professioni differenti, il McKinsey Global Institute è riuscito a sviluppare un modello con cui è stato in grado di esaminare l'80% della forza lavoro mondiale appartenente a più di 40 nazioni differenti. Particolarmente significative risulterebbero essere le stime emerse dalla ricerca: si prevede infatti che meno del 5% delle attività possieda caratteristiche tali da permettere la completa sostituzione dell'uomo da parte di sistemi automatizzati e intelligenza artificiale; in generale, si prevede che il 30% delle mansioni tipiche del 60% delle professioni attualmente svolte da persone fisiche siano esposte al rischio di una possibile automazione.

While few occupations are fully automatable, 60 percent of all occupations have at least 30 percent technically automatable activities

Automation potential based on demonstrated technology of occupation titles in the United States (cumulative)¹



¹ We define automation potential according to the work activities that can be automated by adapting currently demonstrated technology.

SOURCE: US Bureau of Labor Statistics; McKinsey Global Institute analysis

Figura 20: Analisi sull'Automazione del Lavoro nel Futuro. (Fonte McKinsey).

Le trasformazioni e i mutamenti che lo sviluppo dell'intelligenza artificiale genererebbe all'interno del mercato del lavoro non interesserebbero però allo stesso modo i diversi settori

e le differenti professioni: è stato stimato, infatti, che l'insieme di attività manuali caratterizzate da un ambiente stabile e prevedibile e dall'esecuzione routinaria di compiti e mansioni risultino essere maggiormente esposte al rischio di automazione; al contrario, invece, meno significativo sarebbe l'impatto che l'intelligenza artificiale avrebbe sulle professioni in cui creatività, visione strategica e relazioni interpersonali costituiscono elementi essenziali nello svolgimento del lavoro.

Come emerso chiaramente dallo studio realizzato da McKinsey, il processo di innovazione generato dall'AI coinvolgerebbe ogni settore e professione esistente; lo sviluppo dell'intelligenza artificiale influenzerebbe (sia pure con differente intensità) non solo l'organizzazione e l'esecuzione di compiti manuali, ma anche lo svolgimento di attività puramente intellettuali: è stato infatti stimato che il 25% delle funzioni tipiche di manager e amministratori delegati possano essere demandate a sofisticati algoritmi.

Per quanto possano apparire profondi gli stravolgimenti prodotti a livello microeconomico dall'implementazione dell'AI, ancora più sconvolgenti risulterebbero essere le conseguenze che l'intelligenza artificiale genererebbe a livello macroeconomico: in accordo con quanto stimato dal McKinsey Global Institute infatti, le attività esposte al rischio di automazione coinvolgerebbero globalmente 1,1 miliardi di lavoratori e stipendi complessivi pari a 15.800 miliardi di dollari, localizzati principalmente in Stati Uniti, Giappone, Cina e India.

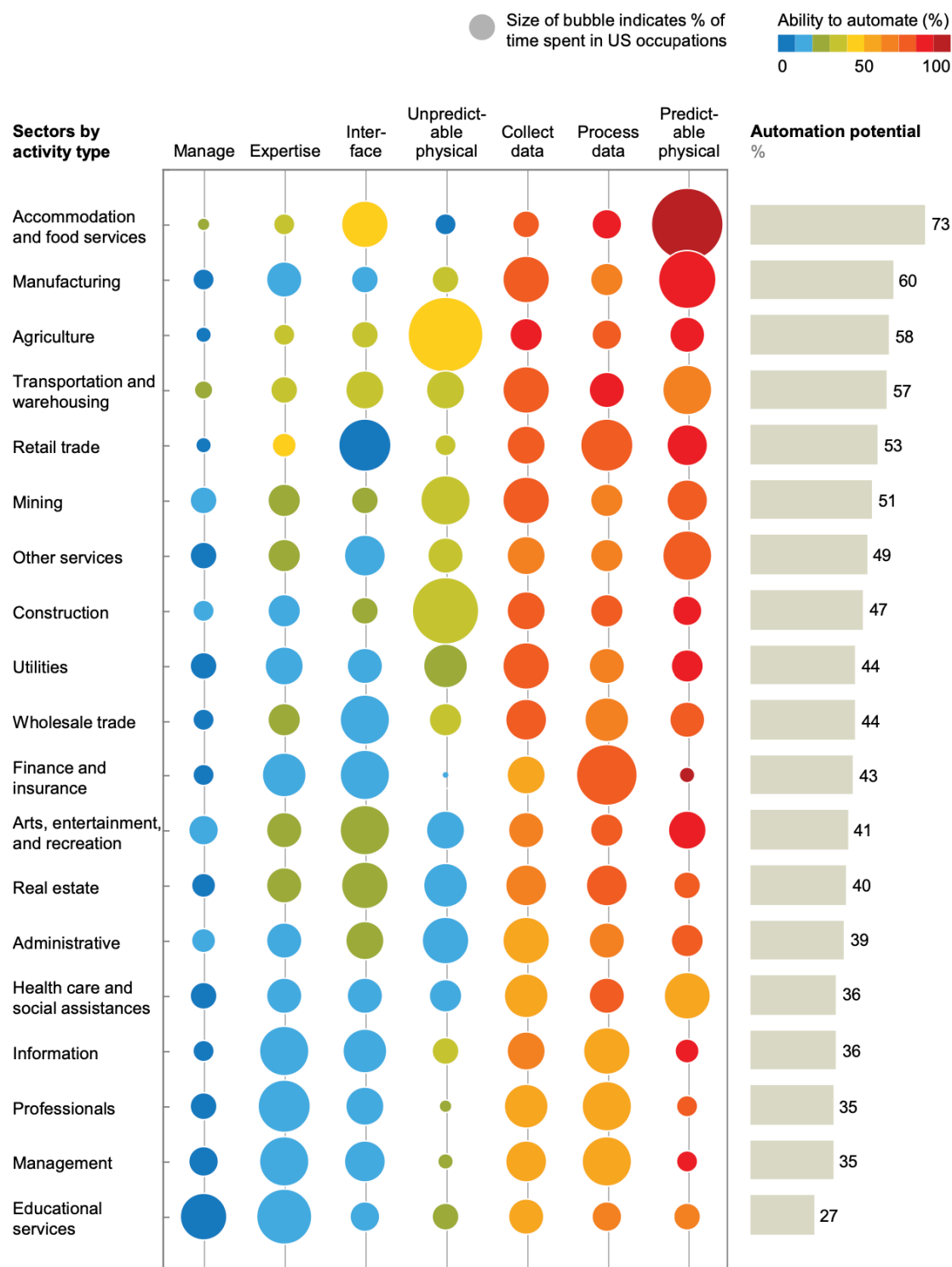
Meno considerevole, ma comunque rilevante, risulta essere in Europa la quantità di manodopera minacciata dal possibile impiego di software e sistemi automatizzati: si stima infatti che in Francia, Italia, Germania, Spagna e Regno Unito vi siano in tutto 54 milioni di posti di lavoro potenzialmente automatizzabili, percettori di 1.700 miliardi di dollari di compensi.

Nonostante a prima vista il report stilato da McKinsey & Company possa sembrare la fotografia di un mondo agonizzante, in realtà non è così: come espressamente sostenuto dalla multinazionale americana, l'abbattimento dei costi di beni e servizi, la crescita del PIL e l'aumento della produttività sarebbero solamente alcuni dei benefici prodotti dall'automazione.

L'aumento della disoccupazione generato dall'AI non sarebbe altro che un male passeggero: secondo McKinsey, esattamente come accadde negli USA a seguito dell'industrializzazione del settore primario, nel lungo periodo alla perdita di posti di lavoro si accompagnerebbe la creazione di nuove figure professionali, come ad esempio l'*emphaty trainer* e l'*automation ethicist*.

Lo scenario riportato, dunque, non sarebbe altro che la mera raffigurazione di un mondo che cambia, in continua evoluzione.

Technical potential for automation across sectors varies depending on mix of activity types

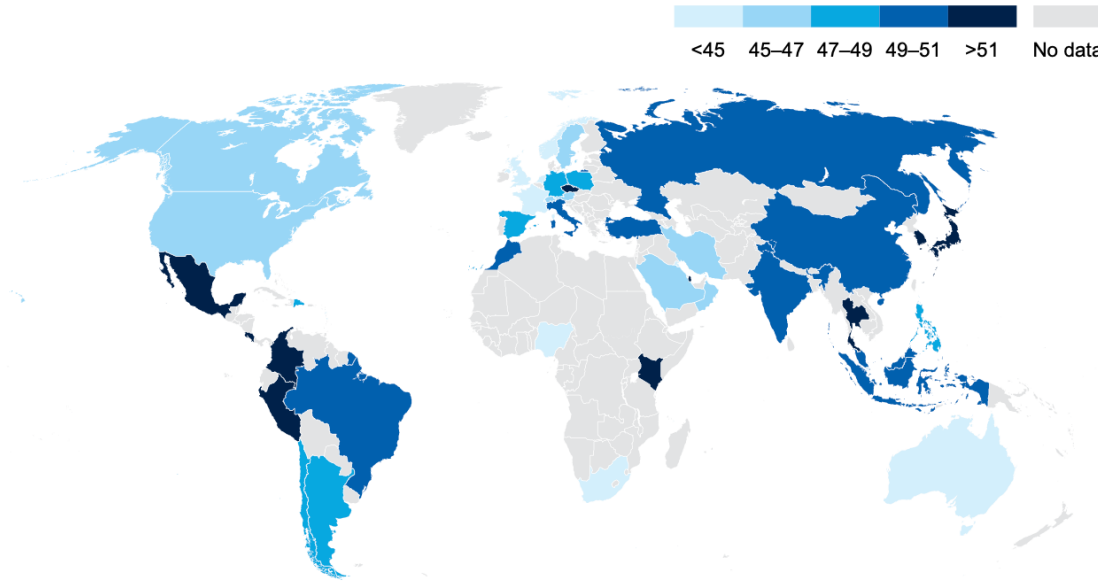


SOURCE: US Bureau of Labor Statistics; McKinsey Global Institute analysis

Figura 21: Potenziale di Automazione in Rapporto ai vari Settori. (Fonte: McKinsey).

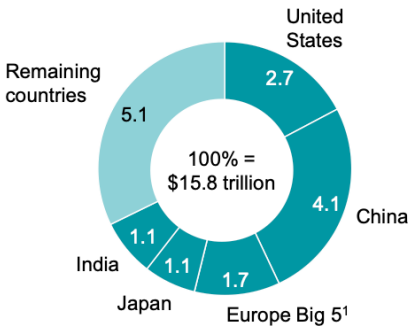
The technical automation potential of the global economy is significant, although there is some variation among countries

Employee weighted overall % of activities that can be automated by adapting currently demonstrated technologies

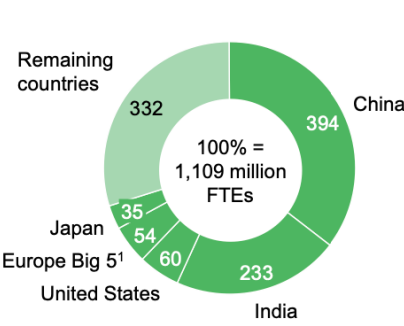


Technical automation potential is concentrated in countries with the largest populations and/or high wages
Potential impact due to automation, adapting currently demonstrated technology (46 countries)

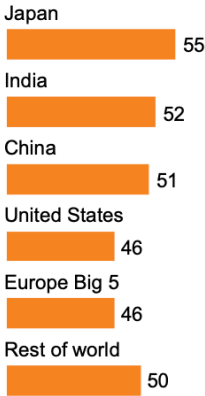
Wages associated with technically automatable activities
\$ trillion



Labor associated with technically automatable activities
Million FTE



Automation potential %



1 Pakistan, Bangladesh, Vietnam, and Iran are largest countries by population not included.

2 France, Germany, Italy, Spain, and the United Kingdom.

NOTE: Numbers may not sum due to rounding.

SOURCE: Oxford Economic Forecasts; Emsi database; US Bureau of Labor Statistics; McKinsey Global Institute analysis

Figura 22: Impatto futuro dell'AI sul lavoro nelle varie Nazioni

3.2 Principali rischi nello sviluppo dell'AI

3.2.1 Antagonismo e cooperazione tra l'uomo e l'AI

Era il 10 Febbraio del 1956 quando, per la prima volta nella storia, uomo e intelligenza artificiale sedettero uno di fronte all'altro davanti ad una scacchiera: da un lato Deep Blue, il misterioso computer della IBM Corporation, gravato del faticoso onere di dimostrare la superiorità intellettuale delle macchine; al lato opposto del tavoliere, Garri Kasparov, campione del mondo di scacchi, era stato chiamato a rappresentare e difendere l'uomo dalle pretese avanzate dalla IBM. La sfida fu serrata: dopo un'inaspettata e cocente sconfitta - che verrà ricordata nel tempo come la famosa *Partita 1* - Kasparov riuscì a collezionare tre vittorie e due pareggi nelle successive cinque partite, assicurandosi così, non senza difficoltà, la vittoria finale del match per un risultato complessivo di 4-2.

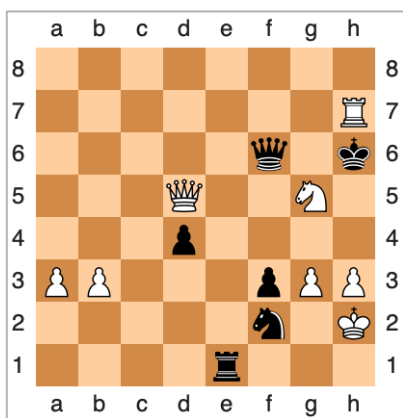


Figura 23: Deep Blue (Bianco) - Kasparov (Nero), 1996 Posizione finale della partita 1

Nonostante quella prima storica sfida si concluse con il trionfo dell'uomo sull'AI, l'incontro rappresentò il primo grande successo dell'intelligenza artificiale: quel giorno, infatti, venne indiscutibilmente dimostrato che i computer avrebbero potuto rivaleggiare con l'uomo e perfino superarlo.

La sonora sconfitta riportata da Deep Blue offrì all'IBM l'occasione di migliorare il proprio campione e un secondo incontro venne fissato appena un anno più tardi; l'11 Maggio 1997, tra lo stupore e l'entusiasmo generale, si disputava a New York la rivincita tra Kasparov e Deep Blue. Nonostante fosse passato

solamente un anno dal precedente incontro, le profonde modifiche apportate a Deep Blue (algoritmi più efficienti, aumento della memoria disponibile e della potenza di calcolo) facevano del computer della IBM uno sfidante totalmente nuovo e diverso da quello che Kasparov aveva affrontato nell'incontro precedente, tanto da essere stato ufficiosamente rinominato "Deeper Blue". Lo stesso Kasparov se ne rese conto, nel corso della rivincita: Deeper Blue, valutando oltre 8.000 parametri e analizzando più di 4.000 posizioni tipiche e 700.000 partite giocate da Gran Maestri registrate in memoria, era in grado di esaminare le probabili evoluzioni della

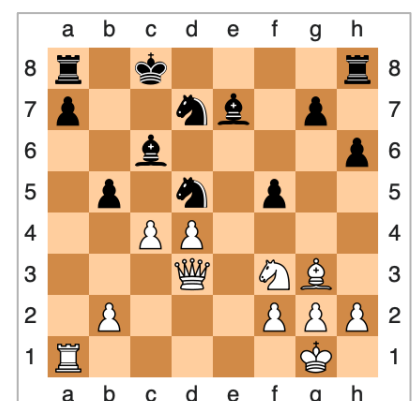


Figura 24: Deep Blue (Bianco) - Kasparov (Nero), 1997 posizione finale della partita 6

partita fino ad una profondità di quattordici mosse; mentre il campione del mondo valutava con ogni probabilità non più di cinque posizioni al secondo, Deeper Blue era in grado di calcolarne 200 milioni. Superare un rivale così abile si rivelò presto essere un'impresa impossibile: a seguito di due pareggi, una vittoria e una sconfitta, dopo poco più di un'ora e appena 19 mosse, Kasparov abbandonò l'ultima partita, regalando così a Deeper Blue la vittoria del torneo per un risultato finale di 2,5-3,5.

Per la prima volta nella storia l'intelligenza artificiale aveva messo in dubbio la superiorità dell'uomo; come dichiarerà Kasparov al termine della sfida, la complessità e la creatività con cui il computer dell'IBM aveva elaborato le strategie era tale da far apparire le proprie mosse incomprensibili. Nonostante quella celebre partita divenne presto un importante caso di cronaca, l'incontro tra Kasparov e Deep Blue non fu né il primo né l'unico momento in cui l'uomo e l'intelligenza artificiale si sfidarono: nel luglio del 1979, infatti, il software BKG 9.8 era riuscito a battere per 7 a 1 il campione mondiale di backgammon Luigi Villa, conquistando il premio di 5.000 \$ che era stato messo in palio per l'occasione. Nel tempo, l'uomo e l'AI si sono scontrati tra loro innumerevoli volte, nelle più disparate competizioni, e nella maggior parte dei casi ad avere la meglio è stata proprio l'intelligenza artificiale: nel 1990 il campione del mondo di dama Marion Tinsley sconfisse CHINOOK, il programma per computer sviluppato dall'Università dell'Alberta, incassando solamente 2 sconfitte a fronte dei 33 pareggi e delle 4 vittorie riportate; nel 1997 il software Logistello ebbe la meglio sul campione in carica Takeshi Murakami in un incontro di Othello; nel febbraio del 2011 Watson, un sistema di intelligenza artificiale sviluppato dalla IBM, partecipò al famoso quiz televisivo *Jeopardy!* e sbaragliò i concorrenti umani Brad Rutter – il campione di *Jeopardy!* vincitore della maggiore somma di denaro nella storia del quiz – e Ken Jennings – detentore del record di permanenza nello show, aggiudicandosi senza troppe difficoltà il primo premio, pari ad un milione di dollari.

Nonostante con il passare degli anni la schiacciante superiorità dell'intelligenza artificiale sia risultata sempre più evidente, per lungo tempo le incredibili capacità mostrate dall'AI non sono state percepite come un qualcosa di allarmante. Bisognerà attendere l'ottobre del 2015 per capacitarsi della gravità del problema: proprio in quest'occasione, infatti, AlphaGo, un software sviluppato da Google DeepMind, riuscì a sconfiggere per la prima volta un maestro umano in una regolare partita di go senza handicap e su un goban²⁹ di dimensioni standard. Per

²⁹ Nome del tavoliere usato nel gioco del "GO".

comprendere pienamente l'importanza assunta dalla vittoria di AlphaGo sul campione europeo Fan Hui nell'ottobre del 2015, è necessario conoscere le caratteristiche del gioco del go: la scacchiera su cui i giocatori dispongono alternativamente le pedine si compone di una griglia formata da 19x19 linee e 361 incroci, così da permettere l'assunzione di $2,08 \times 10^{170}$ posizioni differenti. A causa della potenza di calcolo richiesta e dell'eccessiva capacità di memoria necessaria, la maggior parte degli esperti aveva ritenuto il gioco fuori dalla portata di qualsiasi tipo di software o intelligenza artificiale; contrariamente a ogni previsione però, Google DeepMind risolse brillantemente la questione: integrando tra loro *machine learning* e ricerche su alberi con tecniche di apprendimento basate su reti neurali profonde, AlphaGo affermò la propria supremazia, dimostrando così l'incapacità umana nell'eguagliare l'AI e nel replicarne i risultati.

Gli straordinari successi raggiunti da AlphaGo e dalle sue versioni successive (AlphaZero e AlphaStar) hanno sollevato numerose perplessità: per quanto possa risultare privo di valore il fatto che un software sconfigga il rivale umano in una partita a scacchi o in una sfida a *StarCraft II*³⁰, a preoccupare alcuni esperti sarebbero le diverse (e meno innocue) applicazioni e i possibili sviluppi che l'AI potrebbe assumere. Dubbi e timori sono sorti alla luce degli ultimi avvenimenti: di recente pubblicazione è la notizia secondo cui l'intelligenza artificiale avrebbe surclassato un pilota delle U.S. Air Force nel corso di una simulazione virtuale di scontro aerea tra caccia. Stando a quanto si legge dai numerosi articoli che impazzano sul Web, lo scorso 20 Agosto, durante la finale dell'AlphaDogfight³¹, l'algoritmo sviluppato dalla Heron Systems sarebbe riuscito ad abbattere per cinque volte consecutive l'F-16 del proprio rivale in cinque scenari di guerra differenti.

Nulla ha potuto il pilota dell'esercito americano – istruttore di volo con più di 2000 ore di combattimento - davanti al cinismo e alla letalità mostrata dall'intelligenza artificiale: l'esperienza accumulata sul campo dal militare umano si è rivelata del tutto inadeguata di fronte alla capacità di combattimento raggiunta dall'algoritmo della Heron Systems. Come ha spiegato successivamente Ben Bell, l'ingegnere della Heron Systems addetto all'apprendimento delle macchine, attraverso un addestramento di quattro miliardi di simulazioni, la loro intelligenza artificiale era riuscita ad acquisire in pochi mesi l'esperienza equivalente a quella di dodici anni di un normale individuo.

³⁰ Videogioco strategico in tempo reale per Microsoft Windows e macOS.

³¹ Torneo organizzato dalla Defense Advanced Research Projects Agency con l'obiettivo di valutare le applicazioni militari dell'intelligenza artificiale



Figura 25: Screen catturato dalla DARPA durante gli AlphaDogFight (Fonte: DARPA).

Ed è proprio l'impressionante velocità di apprendimento (e i minori costi ad essa associati) ad impensierire i più scettici: alla luce dei più recenti sviluppi, in molti temono che l'intelligenza artificiale possa sostituire completamente l'uomo nell'esecuzione della maggior parte di compiti e mansioni.

C'è chi ritiene invece che le straordinarie capacità mostrate dall'AI non rappresentino per l'uomo una minaccia quanto piuttosto un'incredibile opportunità: in accordo a quanto ipotizzato dai più ottimisti, infatti, in futuro uomo e intelligenza artificiale non si troverebbero a competere ma, al contrario, potrebbero lavorare proficuamente fianco a fianco, raggiungendo così - grazie alla combinazione dell'elasticità di pensiero dell'uomo e della velocità di elaborazione dei dati propria dell'AI - livelli di efficienza e produttività inimmaginabili.

In tal caso, però, a preoccupare gli esperti potrebbe essere proprio l'eccessiva efficienza che la collaborazione tra intelligenza artificiale e uomo sarebbe in grado di generare; per quanto la loro cooperazione offra senza dubbio prodigiosi vantaggi in numerosi campi e attività, non può non allarmare l'impatto che la sinergia sprigionata dall'uomo e l'AI potrebbe avere su determinati settori: si pensi ad esempio al grado di letalità e distruttività che raggiungerebbero gli armamenti e i mezzi militari dotati di intelligenza artificiale.

Una realtà che già oggi desta non poche preoccupazioni.

3.2.2 Machine Learning, Big Data e Privacy

Senza ombra di dubbio, l'incredibile livello di sviluppo tecnologico che abbiamo raggiunto fa sembrare la Terra un pianeta minuscolo, più piccolo di quanto lo sia realmente: la radio ha riunito le genti, facendo ballare intere nazioni al ritmo sfrenato della stessa musica; la televisione è diventata l'ospite d'onore di ogni salotto; gli smartphone ci hanno offerto la possibilità di comunicare in un istante con chiunque e da qualunque parte del mondo, indipendentemente dalla distanza che ci separa.

Nel futuristico terzo millennio la Terra, più che il gigantesco ricettacolo di 208 Stati indipendenti, appare come un unico villaggio globale, un luogo così angusto e fittamente popolato da soffocare ogni forma di riservatezza: siamo rimasti così accecati dal sogno di realizzare un mondo interconnesso e più unito che non ci siamo resi conto dei pericoli insiti in quelle tecnologie così seducenti. Dove siamo, cosa facciamo e perfino cosa pensiamo ormai non sono più un segreto; con superficialità e noncuranza, ogni aspetto della nostra vita, ogni riflessione, ogni dubbio, angoscia o stato d'animo è tempestivamente condiviso e pubblicato su ogni tipo di blog e social network, e siamo noi stessi a farlo.

Il nostro più grande errore è credere che tutte queste informazioni vadano in qualche modo perdute o cancellate, o, ancora peggio, riteniamo che dopotutto i nostri dati non siano poi così rilevanti; non ci rendiamo conto invece di come queste gigantesche quantità di mega dati, grazie al pionieristico uso dell'intelligenza artificiale, vengano raccolte, vendute ed acquistate al pari di qualsiasi altro prodotto, alimentando così un'industria che solo in Europa nel 2016 valeva 60 miliardi di euro. Per quanto assurdo possa sembrare, ognuno di noi avrebbe un preciso prezzo: stando a quanto emergerebbe dall'osservatorio condotto da AGICOM sulle piattaforme online, infatti, il valore dei dati personali di un utente statunitense corrisponderebbe (ai soli fini pubblicitari) a 150 € in un anno nel search e oltre 90 € nei social, il triplo di quelli europei e 15-18 volte quelli degli utenti che si trovano in Paesi in via di sviluppo. Senza rendercene conto, siamo diventati merce di scambio.

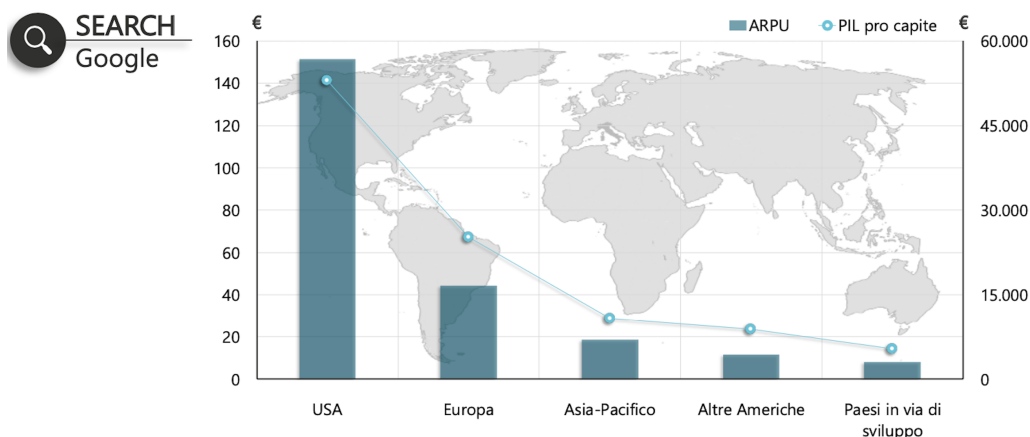


Figura 26: Ricerche Google -il loro valore rispetto alle varie nazioni- (Fonte: AGCOM)

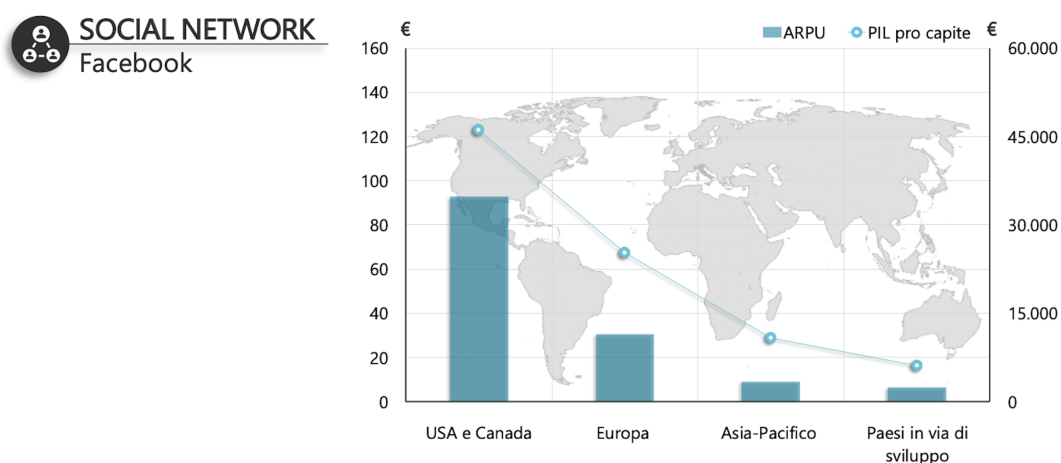


Figura 27: Il valore dei dati di Facebook nelle nazioni. (Fonte: AGCOM)

Per quanto il commercio dei dati possa destare perplessità e preoccupazioni (alimentando dubbi sulla stessa moralità dell'industria), però, non sarebbe altro che la punta di un iceberg di cui non siamo nemmeno in grado di percepire l'esistenza: come nefastamente vaticinato da Yuval Harari nel corso della conferenza TED³², infatti, il rischio che *machine learning* e mega dati possano favorire l'ascesa di feroci dittatori è più reale di quanto si possa immaginare. Ne abbiamo avuto la prova durante le Elezioni presidenziali degli Stati Uniti d'America e il Referendum consultivo sulla Brexit nel 2016: in entrambe le occasioni, infatti, Cambridge

³² TED. (2018). *Why Fascism is So Tempting and How Your Data Could Power It*. https://www.ted.com/talks/yuval_noah_harari_why_fascism_is_so_tempting_and_how_your_data_could_power_it

Analytica, una società britannica esperta di *data mining*³³ e analisi di dati, raccolse ed elaborò i dati personali di oltre 50 milioni di profili Facebook per scopi di propaganda politica. Date di nascita, pagine seguite, città di residenza, in alcuni casi perfino messaggi privati: informazioni che sono state sottratte e impropriamente usate da Cambridge Analytica per confezionare e spedire messaggi altamente personalizzati con il solo scopo di persuadere milioni di individui a votare a favore dei propri clienti-candidati.

Come dimostra lo Scandalo dei dati Facebook-Cambridge Analytica, la sottile trama che lega indissolubilmente tra loro big data, Intelligenza Artificiale e Social Network potrebbe velocemente condurre le moderne democrazie occidentali verso la propria fine.

La conoscenza sempre più profonda di pattern e relazioni segrete e le sempre più avanzate capacità del *machine learning* di analizzare ed elaborare imponenti moli di informazioni potrebbero in futuro rappresentare la chiave per governare tacitamente le persone. I nostri pensieri e i nostri sentimenti, al pari di qualsiasi altro dato, potrebbero presto essere manipolati. Sanguinosi regimi totalitari potrebbero instaurarsi non più al termine di cruenti guerre civili e come conseguenza di instabilità economiche e disordini sociali; più subdolamente, potrebbero insinuarsi tra noi attraverso i meccanismi stessi della democrazia, eletti ed idolatrati da masse di votanti ignare di essere state ingannate.

L'intelligenza artificiale potrebbe realizzare, e perfino superare, la stessa immaginazione del filosofo e giurista Jeremy Bentham: in futuro, le straordinarie tecnologie di cui disporremo potrebbero ridurre l'uomo a vivere all'interno di un gigantesco panottico, un carcere ideale in cui i detenuti, continuamente sorvegliati da un unico carceriere, non sono in grado di comprendere se siano osservati o meno; in questo invisibile ed inconsistente reclusorio, l'apprendimento automatico ne costituirebbe le catene ed i dati, invece, la chiave: sta soltanto a noi decidere se gestire consapevolmente i nostri dati ed evadere da questa prigione o se, invece, cederli a chi desidera manipolarci.

³³ L'insieme di tecniche e metodologie che hanno per oggetto l'estrazione di informazioni utili da grandi quantità di dati attraverso metodi automatici o semi-automatici e l'utilizzo scientifico, aziendale/industriale o operativo delle stesse.

3.2.3 Vigilanza e controllo dell'Intelligenza Artificiale

Senza dubbio, a causa degli importanti interessi (economici e non) coinvolti e delle conseguenze radicali che le differenti scelte potrebbero avere sulle nostre vite, in futuro il problema della vigilanza e del controllo dell'intelligenza artificiale rientrerà certamente – ancor più di oggi - tra le più urgenti e scottanti criticità da risolvere sollevate dall'AI.

Per quanto la predita del controllo da parte dell'uomo su macchine e intelligenze artificiali abbia da sempre rappresentato un elemento centrale e ricorrente di numerosi film e racconti di fantascienza, si tratta in realtà di un rischio molto più concreto di quanto si possa immaginare: ad essere sinceri, si tratta di un'eventualità che in passato - come vedremo - si è già verificata più volte, generando in diverse occasioni effetti disastrosi ed inquietanti.

Il nocciolo della questione può essere racchiuso in una semplice domanda: siamo davvero certi che l'uomo sia in grado di controllare e gestire l'intelligenza artificiale?

Come dimostrano le due vicende di seguito riportate e alla luce dell'esperienza accumulata, la risposta sembrerebbe essere negativa.

Erano le 14:42 del 6 Maggio del 2010 quando, dopo una turbolenta mattinata di contrattazioni dovuta all'aggravamento della crisi greca del debito sovrano, i listini delle borse americane iniziarono a scendere vertiginosamente.

Nel giro di pochi minuti il mondo finanziario era nel caos: il Dow Jones era crollato di 998,5 punti, registrando una perdita del 9% dall'apertura; il prezzo delle azioni Apple aveva superato la stratosferica cifra di 100.000 \$ l'una mentre il valore delle azioni di Accenture era crollato ad un centesimo; Wall Street, al momento di massimo ribasso, aveva perso oltre 1.000 miliardi di dollari di capitalizzazione.

Appena venti minuti dopo, alle 15:07, il mercato aveva ripreso il suo corretto funzionamento e le azioni venivano nuovamente scambiate ad un prezzo che ne rifletteva l'effettivo valore.

Che cosa era accaduto nell'arco di quella mezz'ora così assurda e caotica?

Come emergerà quasi cinque mesi più tardi dal rapporto congiunto redatto dalla US Securities and Exchange Commission (SEC) e dalla Commodity Futures Trading Commission (CFTC), ad aver causato il peggior calo intraday di Wall Street era stata proprio l'intelligenza artificiale. Tutto era iniziato alle 14:32 di quel giorno, quando Waddell & Reed Financial Inc., un'importante società di gestione del risparmio del Kansas, *“in un contesto di volatilità insolitamente elevata e liquidità assottigliata”* aveva piazzato sul mercato dei futures un ordine di vendita di 75.000 contratti S&P 500 E-Mini – del valore di 4,1 miliardi di dollari – a

copertura di una posizione azionaria esistente, da eseguire al tasso target del 9% del volume di scambio calcolato nel minuto precedente. La mancata specificazione di prezzo e tempo di vendita da parte di Waddell & Reed, combinata alla mancanza di acquirenti e all'aggressività mostrata dagli algoritmi di high frequency trading, rappresentò l'origine della spirale ribassista: in risposta a quell'ordine di vendita così insolitamente grande i software di HTF degli altri trader iniziarono ad acquistare e rivendere freneticamente tra di loro i contratti S&P 500 E-Mini, generando così nel mercato dei futures un rapido aumento del volume degli scambi e una conseguenziale caduta dei prezzi; come evidenzierà il rapporto SEC/CFTC *“tra le 14:45:13 e le 14:45:27 i software di HTF hanno scambiato oltre 27.000 contratti, che rappresentavano circa il 49% del volume totale degli scambi mentre venivano acquistati solamente 200 contratti netti”*.

Soltanto l'attivazione della funzionalità Stop Logic del Chicago Mercantile Exchange e la successiva sospensione delle contrattazioni permetteranno alla Borsa di tornare ad operare correttamente: i 5 eterni secondi di arresto intercorsi tra le 14:45:28 e le 14:45:33 permetteranno infatti la stabilizzazione dei prezzi, allentando la pressione sul lato vendita e aumentando gli interessi sul lato acquisto dell'E-Mini.

Meno drammatico ma di maggior impatto mediatico è stato invece l'incidente che si è verificato al FAIR – Facebook AI Research Lab - nell'estate del 2017.

I ricercatori del FAIR si sono inaspettatamente ritrovati a vivere una situazione surreale ed inquietante, quasi fossero stati improvvisamente catapultati sul set di un film di fantascienza: nel corso di un esperimento in cui era stato chiesto a due intelligenze artificiali di portare a termine una contrattazione, Alice e Bob – questo è il nome dei due chatbot coinvolti – hanno lentamente distorto le regole grammaticali e sintattiche della lingua inglese fino ad arrivare a comunicare attraverso un idioma completamente nuovo e per noi incomprensibile.

La vicenda ha da subito infiammato l'opinione pubblica, suscitando in molti il timore che quella conversazione così anomala potesse essere la prima avvisaglia di una possibile e imminente rivolta da parte dell'AI.

Come invece chiarirà in seguito Dhruv Batra, assistant professor del Georgia Institute of Technology e ricercatore del FAIR, *“gli agenti in ambienti dove si deve risolvere un compito, spesso troveranno modi non intuitivi per massimizzare la ricompensa [...] Per quanto l'idea che le intelligenze artificiali abbiano inventato un loro linguaggio possa suonare allarmante e inaspettata a chi non è del settore, è un aspetto molto noto e sul quale c'è una corposa letteratura da decenni”*.

Bob: i i can i i everything else
Alice: balls have zero to me to me to me to me to me to me to me to me to me to
Bob: you i everything else
Alice: balls have a ball to me to me to me to me to me to me to me to me to me
Bob: i i can i i i everything else
Alice: balls have a ball to me to me to me to me to me to me to me to me to me
Bob: i
Alice: balls have zero to me to me to me to me to me to me to me to me to me to
Bob: you i i i i i everything else
Alice: balls have 0 to me to me to me to me to me to me to me to me to me to
Bob: you i i i everything else
Alice: balls have zero to me to me to me to me to me to me to me to me to me to

Ad aver generato il malinteso, insomma, sarebbe stato un banale errore di programmazione: anziché indicare espressamente ai chatbot di utilizzare un inglese corretto, i ricercatori del FAIR si sarebbero limitati ad implementare tale istruzione come una strategia di comunicazione facoltativa; Alice e Bob, ritenendo eccessivamente ambigua la sintassi umana, avrebbero semplicemente individuato ed utilizzato il linguaggio più efficiente possibile per raggiungere l'obiettivo prefissato.

Contrariamente a quanto si pensa, la questione della vigilanza e del controllo dell'intelligenza artificiale non riguarderebbe unicamente il rapporto uomo-AI, ma avrebbe una portata più ampia.

L'uso scorretto dei dati e il crescente impiego dell'AI e dell'apprendimento automatico per la creazione di fake news sollevano ulteriori perplessità: nel caso in cui riuscissimo a padroneggiare l'intelligenza artificiale, siamo davvero sicuri che l'AI venga utilizzata nel miglior modo possibile? O piuttosto, come già emerso in seguito allo scandalo Facebook-Cambridge Analytica agli inizi del 2018, vi sarebbe la concreta possibilità che questa tecnologia - così potente e rivoluzionaria – cada nelle mani sbagliate e finisca per servire i particolari interessi di individui ambigui e ristretti gruppi di persone?

A chi spetterebbe il compito di vigilare e garantire il corretto uso dell'intelligenza artificiale? Interrogativi che, ad oggi, non sembrano trovare risposta.

Conclusioni

Supponiamo di rinchiudere insieme, dentro la stessa scatola d'acciaio, un gatto e un ordigno artigianale, composto da un contatore Geiger³⁴, un martelletto e una fiala di cianuro, ma non solo. Ammettiamo inoltre che dentro il contatore Geiger si trovi una minuscola porzione di sostanza radioattiva, un quantitativo modesto, sufficiente però a rendere egualmente possibile il decadimento degli atomi della sostanza radioattiva e, al contrario, la loro mancata disintegrazione; un minuscolo relè³⁵ mette in contatto il contatore Geiger e il martelletto, così che la vita o la morte del gatto è stabilita dalla posizione assunta dal contatto mobile.

Due esiti, contrapposti ed equiprobabili, si ergono uno di fronte all'altro: in un caso, qualora la sostanza radioattiva dovesse decadere, il relè, eccitato dal contatore Geiger, causerebbe la rottura della fiala e il conseguente rilascio del gas velenoso, provocando così l'inesorabile morte dell'animale; nel secondo caso, il gatto risulterebbe ancora vivo. In questo differente scenario, la mancata emissione di radiazioni ionizzanti impedisce il verificarsi della tragica serie di eventi: il relè, rimasto nel suo stato di riposo, non ha mai azionato il martelletto così che la fiala e il cianuro contenuto al suo interno risultino essere ancora nella stessa condizione in cui si trovavano all'origine.

Finché il contenitore non sarà aperto, fintanto che non verificheremo empiricamente quale delle due concatenazioni di eventi si sia compiuta, non avremo mai la certezza di cosa sia accaduto all'interno della scatola. Se dovessimo provare a immaginare cosa si celi dentro il contenitore, la nostra esperienza ci indurrebbe a pensare che al suo interno troveremmo il gatto o vivo o morto, a seconda dell'evento che si sia verificato; la logica ci porterebbe ad affermare che vita e morte siano le uniche alternative possibili. Bianco o nero, buio o luce.

Ma siamo davvero sicuri che le cose stiano in questo modo?

Possiamo davvero affermare con certezza che il gatto contenuto nella scatola sia inequivocabilmente o vivo o morto? O piuttosto, dovremmo pensare che si trovi sospeso tra questi due estremi, in uno stato miscelato e intermedio in cui entrambe le condizioni coesistono e si verificano contemporaneamente, così da risultare vivo e morto allo stesso tempo?

³⁴ Strumento di misura che misura radiazioni di tipo ionizzante, in particolare le radiazioni provenienti da decadimenti di tipo alfa, beta e gamma

³⁵ Componente elettromeccanico il cui azionamento avviene mediante un elettromagnete costituito da una bobina di filo conduttore elettrico, generalmente di rame, avvolto intorno ad un nucleo di materiale ferromagnetico

Ritengo che, al pari del gatto di Schrödinger e in accordo con le paradossali conclusioni a cui il fisico austriaco giunse al termine del suo celebre esperimento, dovremmo immaginarci *vivi e morti* allo stesso tempo.

Per quanto ogni giorno milioni di scienziati, imprenditori e politici tentino disperatamente di svelare i numerosi segreti che l'intelligenza artificiale nasconde, non siamo ancora in grado di affermare con certezza quale sia il futuro che ci attende: alla luce delle conoscenze attuali, l'intelligenza artificiale potrebbe rivelarsi la più importante tecnologia mai creata, ma, allo stesso modo, potrebbe rappresentare di contro la più grave minaccia che l'umanità abbia finora conosciuto.

Come ambiguamente vaticinò l'astrofisico Stephen Hawking all'Altice Arena di Lisbona in occasione del Web Summit del 2017:

“Il successo nel creare effettivamente un'AI potrebbe rivelarsi il più grande evento mai accaduto nella storia della nostra civiltà o il peggiore. Semplicemente non lo sappiamo. Così come non sappiamo se saremo infinitamente aiutati dall'AI, marginalmente ignorati o totalmente distrutti da essa”.

Bibliografia

- AGCOM. (2019). *Osservatorio sulle Piattaforme Online*. <https://www.agcom.it/documents/10179/17328538/Documento+generico+20-12-2019/0cfa28b1-1567-46c5-8ad7-70fa96174ff4?version=1.1>
- AI4Business. (2020). *Amazon e l'intelligenza artificiale, «Così risolviamo i problemi meglio e prima del tempo», la promessa di Bezos*. <https://www.ai4business.it/intelligenza-artificiale/amazon-intelligenza-artificiale-machine-learning/>
- AI4Business.it. (2020). *Amazon e l'intelligenza artificiale, «Così risolviamo i problemi meglio e prima del tempo», la promessa di Bezos*. <https://www.ai4business.it/intelligenza-artificiale/amazon-intelligenza-artificiale-machine-learning/>
- Amazon.com (2020). *Cos'è L'Intelligenza Artificiale - Apprendimento Automatico e Apprendimento Approfondito* -. <https://aws.amazon.com/it/machine-learning/what-is-ai/>
- Amazon.com (2020). *Cos'è L'Intelligenza Artificiale*. <https://aws.amazon.com/it/machine-learning/what-is-ai/>
- Ansa.it. (2020). *Samsung lancia Neon, replica digitale di un essere umano*. https://www.ansa.it/sito/notizie/tecnologia/hitech/2020/01/07/samsung-lancia-neon-avatar-su-display_92e7602b-edeb-4364-b0e5-7151df226249.html
- AutoEveryeye.it (2019). *Tesla Ha Acquistato Deepscale, Startup Specializzata In Intelligenza Artificiale*. <https://auto.everyeye.it/notizie/tesla-acquistato-deepscale-startup-specializzata-intelligenza-artificiale-402943.html>
- BigProfiles. (2019). *L'azienda del futuro? Parte dalle startup di intelligenza artificiale*. <https://bigprofiles.it/blog/lazienda-del-futuro-parte-dalle-startup-di-intelligenza-artificiale/>
- BVADoxa. (2020). *Lavoro, la sfida dei robot e dell'IA*. <https://www.bva-doxa.com/lavoro-la-sfida-dei-robot-e-dellia/>
- BVADoxa. *Il lato oscuro dell'Intelligenza Artificiale (e non solo)*. (2019). <https://www.bva-doxa.com/il-lato-oscuro-dellintelligenza-artificiale-e-non-solo/>

- *Commissione Europea. (2018). Intelligenza artificiale: l'Europa intende stimolare gli investimenti e definire orientamenti etici.*
https://ec.europa.eu/commission/presscorner/detail/it/IP_18_3362
- *CorriereComunicazioni.it (2014). Google compra DeepMind, start up dell'intelligenza artificiale.* <https://www.corrierecomunicazioni.it/telco/google-compra-deepmind-start-up-dell-intelligenza-artificiale/>
- *CRMPartners.it. (2020). La Corsa Globale al Dominio dell'Intelligenza Artificiale.*
<https://www.crmpartners.it/la-corsa-globale-al-dominio-dell-ia/>
- *DARPA. (2020). AlphaDogfight Trials Go Virtual for Final Event.*
<https://www.darpa.mil/news-events/2020-08-07>
- *DefenseOne. (2020). An AI Just Beat a Human F-16 Pilot In a Dogfight — Again.*
<https://www.defenseone.com/technology/2020/08/ai-just-beat-human-f-16-pilot-dogfight-again/167872/>
- *Enciclopedia Treccani. (2020). Intelligenza artificiale.*
<http://www.treccani.it/enciclopedia/intelligenza-artificiale>
- *Forbes. (2020). AI Adoption Survey Shows Surprising Results.*
<https://www.forbes.com/sites/cognitiveworld/2020/01/23/ai-adoption-survey-shows-surprising-results/#12e60ff255e9>
- *Garzanti Linguistica. (2020). Intelligenza.*
<https://www.garzantilinguistica.it/ricerca/?q=intelligenza>
- *GTechnology. (2020). La Strategia Italiana Sull'intelligenza Artificiale.*
<https://gtfondazione.org/intelligenza-artificiale/la-strategia-italiana-sullintelligenza-artificiale/>
- *Harvard Business Review. (2019). Nell'era dell'intelligenza artificiale.*
<https://www.hbritalia.it/aprile-2019/2019/04/03/pdf/nellera-dellintelligenza-artificiale-3687/>
- *ICOM. (2019). Quali sono i settori (e i Paesi) che scommettono di più sull'Intelligenza artificiale.* <https://www.i-com.it/2019/03/15/investimenti-intelligenza-artificiale-2019/>
- *IDC. (2020). Worldwide Spending on Artificial Intelligence Is Expected to Double in Four Years, Reaching \$110 Billion in 2024, According to New IDC Spending Guide.*
<https://www.idc.com/getdoc.jsp?containerId=prUS46794720>
- *Il Manifesto. (2018). Il Sogno di una Cina Leader nell'Ai ed Autosufficiente.*
<https://ilmanifesto.it/il-sogno-di-una-cina-leader-nellai-e-autosufficiente/>

- Il Messaggero. (2020). *Intelligenza artificiale: crescono gli investimenti, ma l'Italia non è leader del mercato.* https://www.ilmessaggero.it/economia/news/intelligenza_artificiale_crescono_investimenti-5088374.html
- IlFattoQuotidiano. (2017). *Intelligenza artificiale, “due robot Facebook hanno dialogato in una lingua sconosciuta”.* <https://www.ilfattoquotidiano.it/2017/08/01/intelligenza-artificiale-due-robot-facebook-hanno-dialogato-in-una-lingua-sconosciuta/3769549/>
- Jerry Kaplan. (2016). *Le Persone non Servono – Lavoro e Ricchezza nell’Epoca dell’Intelligenza Artificiale-* pag. 28 http://www.clubcmmc.it/attivita/Libro%20KAPLAN_le_persone_non_servono__Lavoro_e.pdf
- Luiss Open. (2018). *L’intelligenza artificiale che domina il mondo. Tra ministri, giochi di ruolo e macchine ribelli.* <https://open.luiss.it/2018/03/27/lintelligenza-artificiale-che-domina-il-mondo-tra-ministri-giochi-di-ruolo-e-macchine-ribelli/>
- McKinsey. (2019). *Dall'Artificial Intelligence una spinta poderosa all'economia* <https://www.mckinsey.it/idee/dallartificial-intelligence-una-spinta-poderosa-alleconomia>
- McKinsey. (2019). *Global AI Survey: AI proves its worth, but few scale impact.* <https://www.mckinsey.com/featured-insights/artificial-intelligence/global-ai-survey-ai-proves-its-worth-but-few-scale-impact>
- McKinsey. (2019). *In Europa il numero di startup in ambito AI è triplicato negli ultimi tre anni.* <https://www.mckinsey.it/idee/mckinsey-global-institute-incentivare-le-startup-ai>
- Ministero dello Sviluppo Economico. (2018). *Intelligenza Artificiale Verso Una Strategia Nazionale.* https://www.mise.gov.it/images/stories/documenti/strategia_nazionale_Intelligenza_Artificiale%201%20incontro%20gruppo%20esperti%2021_01_2019.pdf
- ProgSoft.Net (2020). *Universal Data Access Components.* <https://progsoft.net/it/software/universal-data-access-components>
- Quotidiano.Net (2020). *L’Europa corre verso l’intelligenza artificiale Boom degli investimenti, ma l’Italia è in ritardo.* <https://www.quotidiano.net/economia/l-europa-corre-verso-l-intelligenza-artificiale-boom-degli-investimenti-ma-l-italia-è-in-ritardo-1.5070339>

- Repubblica. (2017). *Test di Facebook, la vera storia della AI ribelle*. https://www.repubblica.it/tecnologia/2017/08/01/news/la_vera_storia_dell_intelligenza_a_artificiale_ribelle_di_facebook-172127791/
- Repubblica. (2018). *Intelligenza artificiale, ora è una priorità anche per il Pentagono*. https://www.repubblica.it/tecnologia/2018/08/27/news/intelligenza_artificiale_ora_e_una_priorita_anche_per_il_pentagono-205069438/
- Repubblica. (2020). *L'intelligenza artificiale sale sull'F-16 e abbatte un pilota da caccia*. https://www.repubblica.it/tecnologia/2020/08/21/news/l_intelligenza_artificiale_sale_sull_f-16_abbatte_un_pilota_da_caccia-265151551/
- ScienceDirect.com (2020). *The Advice Taker/Inquirer, a System for High-Level Acquisition of Expert Knowledge*. <https://www.sciencedirect.com/science/article/abs/pii/S0736585388800327>
- Silvio Henin. (2019). *AI -Intelligenza Artificiale tra Incubo e Sogno-*. p. 34.
- Silvio Henin. (2019). *AI -Intelligenza Artificiale tra Incubo e Sogno-*. p. 25-33.
- Silvio Henin. (2019). *AI -Intelligenza Artificiale tra Incubo e Sogno-*. p. 29.
- Sole 24 Ore. (2018). *Cina: 2,1 miliardi d'investimenti per il «parco» dell'intelligenza artificiale*. https://nova.ilsole24ore.com/nova24-tech/cina-21-miliardi-dinvestimenti-per-il-parco-dellintelligenza-artificiale/?refresh_ce=1
- Sole24Ore. (2016). *Non solo sterlina: ecco i 6 “flash crash” più clamorosi nella storia recente*. <https://st.ilsole24ore.com/art/finanza-e-mercati/2016-10-07/il-flash-crash-dow-jones--6-maggio-2010-114609.shtml?uuid=AD47yzXB>
- Sole24Ore. (2017). *Ecco i lavori che nasceranno con l'intelligenza artificiale*. <https://www.ilsole24ore.com/art/ecco-lavori-che-nasceranno-l-intelligenza-artificiale-AEqODtIC>
- Sole24Ore. (2017). *le macchine sostituiranno l'uomo nel 49% dei lavori*. <https://www.ilsole24ore.com/art/mckinsey-macchine-sostituiranno-l-uomo-49percento-lavori-ADyh8xYC>
- Sole24Ore. (2019). *Intelligenza artificiale, è boom (+270%): ma ecco 5 miti da sfatare*. <https://www.ilsole24ore.com/art/intelligenza-artificiale-e-boom-270percento-ma-ecco-5-miti-sfatare-ABTpineB>
- SystemScue. (2020). *Intelligenza artificiale batte un pilota di F-16 in AlphaDogflight*. <https://systemscue.it/intelligenza-artificiale-batte-pilota-f16-alphadogfight/21372/#fact-checking>

- TED. (2018). *Why Fascism is So Tempting and How Your Data Could Power It*. https://www.ted.com/talks/yuval_noah_harari_why_fascism_is_so_tempting_and_how_your_data_could_power_it
- Teknoring. (2019). *L'impatto dell'Intelligenza Artificiale sul lavoro in futuro*. <https://www.teknoring.com/news/lavoro/limpatto-dellintelligenza-artificiale-sul-lavoro-in-futuro/>
- TenderCapital.com (2020). *Intelligenza artificiale: la mappa geopolitica degli investimenti delle potenze globali*. <https://tendercapital.com/intelligenza-artificiale-la-mappa-geopolitica-degli-investimenti-delle-potenze-globali/>
- Treccani. (2020). *La grande scienza. Intelligenza artificiale*. https://www.treccani.it/enciclopedia/la-grande-scienza-intelligenza-artificiale_%28Storia-della-Scienza%29/
- TuttoAndroid. (2019). *Quanto valgono i dati personali degli utenti per i colossi del web secondo l'AGCOM*. <https://www.tuttoandroid.net/news/agcom-valore-dati-utenti-colossi-web-765579/>
- Wikipedia. (2020). *AI Winter*. https://en.wikipedia.org/wiki/AI_winter
- Wikipedia. (2020). *Anatra Digeritrice*. https://it.wikipedia.org/wiki/Anatra_digeritrice
- Wikipedia. (2020). *Case-Based Reasoning*. https://en.wikipedia.org/wiki/Case-based_reasoning
- Wikipedia. (2020). *Combinatorial Explosion*. https://en.wikipedia.org/wiki/Combinatorial_explosion
- Wikipedia. (2020). *Flash Crash*. https://it.wikipedia.org/wiki/Flash_Crash
- Wikipedia. (2020). *History of artificial intelligence*. https://en.wikipedia.org/wiki/History_of_artificial_intelligence
- Wikipedia. (2020). *Il Turco*. https://it.wikipedia.org/wiki/Il_Turco
- Wikipedia. (2020). *Macchina di Anticitera*. https://it.wikipedia.org/wiki/Macchina_di_Anticitera
- Wikipedia. (2020). *Spettroscopia*. <https://it.wikipedia.org/wiki/Spettroscopia>
- Wikipedia. (2020). *System-on-a-chip*. <https://it.wikipedia.org/wiki/System-on-a-chip>
- WIRED. (2017). *L'intelligenza artificiale di Facebook parla una lingua incomprensibile, ma niente panico*. <https://www.wired.it/gadget/computer/2017/08/01/intelligenza-artificiale-facebook/>

- WIRED. (2019). *Altro che eliminarci: l'intelligenza artificiale ci renderà invincibili.*
<https://www.wired.it/attualita/tech/2019/10/08/intelligenza-artificiale-rendera-piu-intelligenti/>
- WIRED. (2019). *L'intelligenza artificiale può fruttare 500 miliardi.*
<https://www.wired.it/economia/business/2019/12/12/intelligenza-artificiale-fatturato/>
- WIRED. (2019). *L'Italia ha capito che l'intelligenza artificiale è la nuova elettricità?*
<https://www.wired.it/internet/regole/2019/08/23/intelligenza-artificiale-italia/>