

FrancoAngeli

Giovanni Acerboni

PAROLA MIA

**Manuale di scrittura professionale
con il supporto delle intelligenze
artificiali del linguaggio**

I lettori che desiderano informarsi sui libri e le riviste da noi pubblicati possono consultare il nostro sito Internet: www.francoangeli.it e iscriversi nella home page al servizio “Informatemi” per ricevere via e.mail le segnalazioni delle novità o scrivere, inviando il loro indirizzo, a “FrancoAngeli, viale Monza 106, 20127 Milano”.

Giovanni Acerboni

PAROLA MIA

**Manuale di scrittura professionale
con il supporto delle intelligenze
artificiali del linguaggio**

FrancoAngeli

Progetto grafico della copertina: Elena Pellegrini

Isbn: 9788835183211

Copyright © 2025 by FrancoAngeli s.r.l., Milano, Italy.

L'opera, comprese tutte le sue parti, è tutelata dalla legge sul diritto d'autore.

Sono riservati i diritti per Text and Data Mining (TDM), AI training e tutte le tecnologie simili.

L'Utente nel momento in cui effettua il download dell'opera accetta tutte le condizioni della licenza d'uso dell'opera previste e comunicate sul sito www.francoangeli.it.

A dagh'el minga a trà, el parla propi ben
(A non dargli retta, parla proprio bene)
Lucia Carini Prestini (1904-1997), *Mia nonna*

A Tommaso Raso

Indice

Introduzione, ovvero: Finché trovi un errore, ce n'è un altro	pag.	9
1. Le Intelligenze Artificiali	»	13
1.1. Colti di sorpresa	»	13
1.2. Macchine pensanti?	»	15
1.3. Chi e come istruisce le AI	»	20
1.4. Il Natural Language Processing	»	29
1.5. Le AI Generative	»	36
2. La conversazione	»	49
2.1. La conversazione faccia a faccia	»	49
2.2. Il Metodo della conversazione scritta	»	63
2.3. Il Metodo di lettura profonda	»	68
3. La scrittura professionale	»	81
3.1. Che cosa è, che cosa non è	»	81
3.2. Norme	»	84
3.2.1. UNI e ISO	»	84
3.2.2. Accessibilità	»	89
3.3. Le relazioni	»	94
3.3.1. La cortesia	»	94
3.3.2. Formule di saluto	»	98
4. Il supporto delle GenAI	»	101
4.1. La conoscenza della lingua	»	101
4.1.1. Lunghezza del periodo	»	102
4.1.2. Incisi	»	103
4.1.3. Verbi e complementi indiretti	»	103
4.1.4. I tecnicismi	»	104
4.1.5. I sinonimi	»	108
4.2. Lo stile delle GenAI	»	109

4.3.	Testare la persuasività	pag.	111
4.4.	Selezionare gli argomenti e disporli in sequenza	»	114
4.5.	Comporre il titolo	»	117
4.6.	Comporre l'abstract	»	120
4.7.	Verificare la correttezza grammaticale	»	121
4.8.	Controllare la chiarezza	»	122
	4.8.1. Le partizioni	»	123
	4.8.2. Le evidenziazioni	»	124
	4.8.3. La chiarezza linguistica	»	125
4.9.	L'accessibilità dei contenuti, di <i>Stefano Castelli</i>	»	132
	4.9.1. Progettazione dei contenuti	»	132
	4.9.2. Link e note	»	133
	4.9.3. Requisiti di accessibilità	»	134
	4.9.4. Cinque problemi tipici e la soluzione IA	»	136
4.10.	Il linguaggio inclusivo	»	144
Bibliografia e altri riferimenti			» 157

Introduzione, ovvero: Finché trovi un errore, ce n'è un altro

Sono stato decisamente fortunato: ho assistito non passivamente e in buona compagnia alla trasformazione digitale della scrittura. Negli ultimi trent'anni è letteralmente cambiato tutto, in particolare dal 2022, quando abbiamo usato la prima versione di ChatGPT.

Non mi pare proprio che nella storia vi siano state trasformazioni altrettanto profonde e rapide. Nel mio piccolo, non so se mi fa più tenerezza o mi fa più sorridere ricordare che nel mio primo libro sulla scrittura professionale, uscito nel 2005, mi era parso necessario tradurre in nota il significato di blog¹.

Dico che ho assistito non passivamente per due ragioni essenziali. La prima è che ho dovuto stare al passo della digitalizzazione per poter fare il mio mestiere di formatore di scrittura e comunicazione. La seconda è che ho voluto affrontare le Intelligenze Artificiali del linguaggio in epoca... non sospetta, nel 2014, quando le AI Generative – GenAI – non c'erano ancora (o, meglio, c'erano, ma non erano così potenti e, soprattutto, non erano interrogabili in linguaggio naturale).

Il mio primo progetto di combinazione tra la scrittura professionale e le AI del linguaggio non ha funzionato. Con la startup innovativa Writexp, avevamo realizzato nel 2019 un editor online che intercettava automaticamente e con implacabile precisione frasi e parole da semplificare, grazie ad algoritmi che avevamo brevettato. I fatti hanno la testa dura: il mercato si aspettava quello che allora non c'era, un ChatGPT che riscrisse automaticamente.

¹ *Progettare e scrivere per Internet*, McGraw-Hill, Milano, 2005.

Quello che allora non c'era, adesso c'è, anzi ce n'è in abbondanza², tanto che qualcuno si chiede se la scrittura sia ancora necessaria e qualcuno lo sostiene addirittura. Non solo perché le macchine potrebbero scrivere al nostro posto, ma anche perché scrivere sarebbe proprio inutile. Si comunica sempre più spesso con video, infografiche, registrazioni vocali, che le macchine possono facilmente trasformare in testi scritti nei casi in cui è effettivamente utile (per esempio, la trascrizione di una riunione facilita enormemente la composizione del verbale).

D'accordo, ci hanno tolto l'esclusiva della scrittura, ma di che cosa stiamo parlando?

Long story short: "information is not knowledge" non l'ha detto in modo difficile un filosofo, è autoevidente, l'ha detto Frank Zappa³, che scriveva la sua musica, così come, per fare un film, bisogna prima scrivere la sceneggiatura.

La conoscenza passa dalla scrittura. Scrivere sarà sempre necessario, almeno a chi fa sul serio.

Direbbe Forrest Gump: "non c'è altro da dire su questa faccenda".

Veniamo a noi.

Come utilizzare le GenAI, per quali obiettivi e con quanta fiducia sono gli argomenti di questo libro e qui voglio solo accennare ai temi principali.

Non saremo sostituiti dalle GenAI. Le GenAI sono software (un sistema di software). Per quanto siano straordinariamente sviluppate, non hanno tutte le capacità necessarie per comporre un testo.

Se non saremo sostituiti, non dovremo sottovalutare il fatto che la scrittura è un'attività nella quale il solo esercizio non sviluppa l'organo. Per gestire le GenAI, dobbiamo sempre saperne più di loro. Ci sfidano, ed è una bella sfida. Non possiamo non stare al gioco. Perderemo la partita se il contributo delle macchine sarà tale che a noi resterà il lavoro sporco. A quel punto non saremo sostituiti dalle macchine, ma da chi ne saprà più di loro.

Casomai, la domanda è: come possiamo scrivere per continuare a fare preferire ai lettori i nostri testi rispetto alle risposte delle GenAI? Non scommetterei alto sulla fermezza con la quale i nostri lettori resisteranno alla tentazione. E non perché ne abbia una scarsa considerazione, anzi. Il fatto è, o sarebbe: che cosa trovano di più, di meglio nei nostri testi? Questa è la sfida.

² Per gli esempi e i test di questo libro mi sono limitato esclusivamente alla versione gratuita di alcune GenAI (ChatGPT e dpChat Avvocati, Copilot, DeepSeek, Gemini, Perplexity e Vitruvian1). Ho pensato infatti che se avessi utilizzato le pur più potenti versioni a pagamento avrei subordinato la comprensione e la verifica del libro a un costo aggiuntivo, che non è intellettualmente onesto chiedere (per non dire dell'alea della pubblicità indiretta).

³ *Packard goose in Joe's Garage*, atto III, 1979.

Più cose, meno parole, più naturalezza di linguaggio, meno spersonalizzazioni, più conversazione, meno falsa confidenza? Altro ancora?

Come si vedrà via via nei vari capitoli (in particolare nel primo e nel quarto), le Gen-AI commettono inevitabilmente errori, e non eseguono con la stessa qualità tutti i task.

Quanto ai contenuti tecnici, legali, organizzativi, commerciali ecc., ogni scrittore conosce bene i suoi. I relativi errori delle GenAI sono visibili al professionista competente, purché le risposte siano lette con attenzione.

Quanto alla lingua, sono brave a scrivere i titoli, ma non altrettanto a migliorare la chiarezza delle frasi⁴, per fare solo due esempi.

Per riconoscere gli errori linguistici, in questo libro illustro i criteri di riferimento che in vent'anni si sono rivelati utili alle migliaia di professionisti che hanno avuto il buon cuore di frequentare i miei corsi. Molti di questi criteri sono enunciati nella norma UNI 11482:2013 *Elementi strutturali e aspetti linguistici delle comunicazioni scritte delle organizzazioni*, della quale sono stato proponente e relatore.

Tra scrivere correttamente ed essere chiari e persuasivi c'è più di una differenza. Opportunamente interrogate e sapientemente controllate, le GenAI possono restituire un feedback immediato, quale non potrebbe essere quello di un collega, e continuo nel tempo, quale non potrebbe essere quello di un formatore (ahimè).

Sapere in anticipo se il proprio testo verrà capito, credo che sia una bella consolazione per chi lo scrive. Diciamocelo: quando scriviamo, non ci liberiamo mai del tutto da questo dubbio. Tanto più quando il testo ha l'obiettivo di far fare qualcosa al destinatario, obiettivo immancabile nei testi professionali.

Poiché a ben vedere è il destinatario a realizzare l'obiettivo di chi gli scrive, facendo quello che gli dice o gli propone di fare, si potrebbe addirittura dire che chi scrive ha bisogno del destinatario. E quando abbiamo bisogno di qualcuno, lo dobbiamo trattare bene. O no?

Con le GenAI ci si può accertare di essere stati chiari e persuasivi. In parte, però, dati i loro limiti attuali. Miglioreranno, ma non è detto che verranno sviluppate con l'obiettivo specifico di misurare la qualità di un testo.

Misurare rispetto ai KPI è sempre più necessario in tutti i campi. Misurare la qualità di un testo prima di finalizzarlo per la pubblicazione sembra impossibile o inutile. Eppure, le AI del linguaggio, non le GenAI da sole, ce la

⁴ Suggerisco al lettore di tenere d'occhio nel tempo il benchmark EVALITA-LLM Leaderboard: https://huggingface.co/spaces/evalitahf/evalita_llm_leaderboard. Attualmente, le GenAI testate eseguono i task in italiano con risultati molto inferiori a quelli che ci si potrebbe aspettare (Bernardo Magnini, Roberto Zanolì, Michele Resta, Martin Cimmino, Paolo Albano, Marco Madeddu e Viviana Patti, *Evalita-LLM: Benchmarking Large Language Models on Italian*, 2025: <https://arxiv.org/abs/2502.02289>).

possono fare se opportunamente addestrate, anzi ce l'hanno fatta. Oggi è possibile sapere dove e perché il testo non risponde alle aspettative del lettore.

Insomma, il mestiere cambia. E, visti i motivi e i dibattiti in corso, ho ritenuto di non dare per scontate quelle questioni di lingua e linguistica necessarie alla chiarezza delle mie argomentazioni e – spero – utili al lettore.

Ci vediamo nel libro. Buona lettura.

Milano, 10 luglio 2025

Ringraziamenti

Ho voluto scrivere un libro pratico senza rinunciare ad approfondire le varie questioni (l'esercizio non sviluppa l'organo). Per questa ragione ho chiesto a chi ne sa più di me di condividere le sue conoscenze, le sue competenze e le sue esperienze sugli aspetti psicologici della conversazione, su come si fanno e sono fatte le tecnologie intelligenti, sul mercato delle AI del linguaggio e della misurazione, sull'accessibilità.

Ringrazio quindi i “six of the best”: Carlo Bonadonna, Massimo Bustreo, Stefano Castelli, Nicola Grandis, Alessandro Parisi e Marco Varone, insieme a Misa Giuliani, l'editor del libro.

Ringrazio chi in questi anni ha frequentato i miei corsi (insegnandomi forse a sua insaputa molte cose), chi ha condiviso con me ricerche, progetti, aule, dibattiti e chi con le sue pubblicazioni ha contribuito a ridurre i limiti del mio punto di vista. Vorrei ringraziare in particolare Vittorino Bottani, Beatrice Branchesi, Jacopo Deyla, Marina Doria, Marco Fabbri, Lorenzo Gregori, Bernardo Magnini, Alessandro Panunzi, Tommaso Raso, Lorenzo Spallino e Ida Tucci.

Ringrazio infinitamente mia moglie Marta Zagórowska e nostra figlia Irene per l'armonia che sanno creare.

1.1. Colti di sorpresa

Proviamo a fare un saltino indietro nel tempo. Torniamo con la mente al novembre 2022, quando OpenAI mise a disposizione di tutti noi la prima versione gratuita di ChatGPT.

Ora immaginiamo che la nostra conoscenza delle Intelligenze Artificiali fosse limitata a qualche articolo sulla robotica applicata alla produzione industriale, all'esperienza con qualche chatbot, e cose così. Insomma, **immaginiamo che non ne sapevamo praticamente niente**.

E immaginiamo che il clamore che ne era seguito, e che non si è ancora placato, ci abbia indotto a usare ChatGPT, magari per curiosità: abbiamo scritto una domanda e abbiamo avuto dopo pochi istanti una risposta plausibile e ben scritta.

Non avremmo provato **la sensazione che dall'altra parte ci fosse veramente qualcuno** che aveva capito quello che volevamo e che ci aveva risposto a tono? E, di conseguenza, non ci saremmo chiesti se ChatGPT e poi tutte le AI Generative – GenAI – che l'hanno immediatamente seguito fossero intelligenti come noi, se apprendessero come noi, se capissero e pensassero, se potessero svilupparsi fino a superare la nostra intelligenza, la nostra inventiva, la nostra abilità ad agire nel mondo?

Del resto, incontravamo termini come intelligenza generativa, reti neurali, apprendimento, addestramento, comportamento, conoscenza, ragionamento, memoria, esperienza, peso sinaptico, linguaggio, comprensione, conversazione, risposta, scrittura, allucinazione ecc.

Questi termini sono o rischiano di apparire ambigui, perché hanno già almeno un altro significato e nemmeno sempre ben determinato.

Tutta l'era digitale è fortemente caratterizzata dal fenomeno del neologismo semantico o rideterminazione o risemantizzazione, cioè dall'attribuzione di un nuovo significato specifico a termini esistenti, come rete, sito,

home (page), pagina (web), mouse ecc. Ma i termini che definiscono e descrivono le GenAI riguardano questioni profonde e generano un disorientamento che non si risolve tanto facilmente con l'abitudine alla novità (come accade per mouse). L'esempio più clamoroso è probabilmente proprio "intelligenza". Che cosa sia quella umana, è una bella domanda. E quella artificiale?

Poi, avevamo in mente anche molta fantascienza, che ci aveva ben sollecitato con il dubbio che prima o poi, e perché no, una AI sarà uguale o addirittura migliore di noi, puntando sul melodrammatico (*E.T. l'extra-terrestre*), sull'epico (*2001 Odissea nello spazio*), sul serio-faceto (*Star Wars*), sul mistero della vita (*Blade Runner*), sulla resistenza umana in un mondo governato da androidi (*Matrix*), ecc. (ognuno completi l'elenco con i suoi titoli preferiti).

Del resto, il sogno di creare esseri artificiali e automi è antichissimo (Prometeo, Golem, ecc.) e giunge in epoca pre-AI attraverso innumerevoli opere.

Vorrei ricordarne solamente una, il dramma dello scrittore ceco Karel Čapek *R.U.R. (Rossumovi univerzální roboti)* del 1920, perché ci sono due parole per noi importanti nel titolo, che, tradotto, suona: *I robot universali di Rossum*.

La prima parola importante è **robot**, che era un neologismo, in realtà coniato dal fratello Josef, che deriva da "robota", che in ceco e in altre lingue slave significa "lavoro". Sulla fortuna globale di questo termine non serve alcun commento. Nel dramma i robot non sono meccanici, ma veri e propri replicanti in carne e ossa, operai artificiali molto più efficienti degli esseri umani.

La seconda parola importante è "Rossum", che deriva da "rozum", che in ceco e in altre lingue slave significa "**intelletto**", "**ragione**". Rossum è l'inventore dei robot. Rossum si oppone allo sfruttamento economico della sua ricerca e della sua invenzione, invece sostenuta dal figlio. Lo scontro tra la ragione scientifica del padre e la ragione economica del figlio e la vittoria del secondo sono precedenti all'inizio del dramma, nel quale i due Rossum non appaiono. Il dramma sviluppa quindi gli effetti della ragione economica, effetti che i protagonisti (il titolare e i dirigenti della fabbrica che produce i robot implementando la "ricetta" segretissima del vecchio Rossum) non riescono a gestire. Morale: la ragione economica è molto pericolosa¹.

Immaginiamo dunque che ci siamo trovati spiazzati: il futuro, sperato o temuto, era arrivato? Un automa, una macchina, un software stanno davvero per

¹ Chiedo perdono dell'oscurità, ma non voglio spoilerare ai danni di chi non ha letto il non universalmente noto dramma: Karel Čapek, *R.U.R. Rossum's Universal Robots*, a cura di Alessandro Catalano, Marsilio, Venezia, 2015. Curiosità: esiste un'azienda bielorusa che si chiama Rozum Robotics.

uguagliarci e superarci nelle più esclusive capacità umane, come pensare, creare, immaginare, avere una coscienza, una morale, un'etica e così via?

E scrivere? **Scriveranno per noi, meglio di noi?** Ci spaventa o ci consola questa prospettiva? È una prospettiva realistica?

1.2. Macchine pensanti?

Le Intelligenze Artificiali vengono da lontano. Affrontano per via matematica questioni culturali, linguistiche, cognitive, etiche ecc. proprie di altre discipline scientifiche e umanistiche².

La riflessione matematica moderna prende avvio dalla domanda “Le macchine sono in grado di pensare?” (“Can machines think?”) a cui il matematico inglese **Alan Turing** tentò di rispondere nel 1950³, pressoché alla fine della sua carriera e della sua vita⁴.

Turing scartò la possibilità di dare una risposta⁵ e propose il Gioco dell’imitazione, un test che sostituisce la domanda iniziale con la domanda: “Che cosa succederà quando è una macchina a fare la parte” di un uomo e “prenderà la decisione sbagliata con la stessa frequenza con cui” la prende un uomo?⁶ (Turing prevede che sarebbe successo entro la fine del secolo XX).

La domanda iniziale e la domanda sostitutiva non sono equivalenti. Con la seconda, infatti, **Turing limitava la questione dell’intelligenza di una macchina alle sue prestazioni**, in particolare alle prestazioni linguistiche.

Soffermiamoci su questo. La lingua è la sfida più difficile per un matematico, perché è la lingua che deve rappresentare l’intelligenza della macchina. Quante cose in comune hanno lingua, intelligenza e matematica?

² Per una sintesi si può leggere Gianfranco Minati, *Breve viaggio nell’Intelligenza Artificiale*, Quaderni dell’AIEMS – n. 5, a cura di Serena Dinelli, ottobre 2024.

<https://www.aiems.eu/pubblicazioni/quaderni-dellaiems/qa05>.

³ *Computing Machinery and Intelligence*, in “Mind”, ottobre 1950, ripubblicato ora con notevole tempismo: *Macchine calcolatrici e intelligenza*, a cura di Diego Marconi, Einaudi, Torino, 2025. L’originale si può leggere online, per esempio:

<https://courses.cs.umbc.edu/471/papers/turing.pdf>.

⁴ Nel 1951 fu accusato di omosessualità, allora reato in Gran Bretagna, e condannato alla castrazione chimica, che lui preferì al carcere. Morì probabilmente suicida nel 1954. Eccellente la biografia del matematico Andrew Hodges, *Alan Turing. Una biografia*, Bollati Boringhieri, Torino, 2006 (edizione originale: 1983). Il film *The Imitation Game* di Morten Tyldum (2014) narra non del tutto fedelmente come Turing riuscì a decrittare i messaggi segreti tedeschi crittografati con la macchina Enigma.

⁵ “Credo che la domanda iniziale [...] sia troppo insensata perché valga la pena discuterne”, *Macchine calcolatrici e intelligenza*, cit. p. 21.

⁶ Ivi, p. 6.

Per Turing, l'intelligenza equivarrebbe alla capacità di **simulare un comportamento umano** (s'intende, arrivando magari allo stesso risultato, ma con processi e procedimenti diversi). Con questa mossa, Turing aggirava problemi enormi come quello della definizione di pensiero, di intelligenza e di coscienza, quest'ultimo posto chiarissimamente dal suo collega Geoffrey Jackson nel 1949, che guarda caso si riferiva nuovamente alla lingua:

Fino a quando una macchina non sarà in grado di scrivere un sonetto o di comporre un concerto grazie ai suoi pensieri e alle emozioni che prova, e non per via di una pioggia casuale di simboli, non potremo accettare che una macchina sia uguale a un cervello; s'intende, non solo di scrivere un sonetto, ma di sapere di averlo scritto. Nessun meccanismo può provare piacere (e non soltanto segnalare artificialmente, cosa facile da realizzare) per i propri successi, dolore quando le sue valvole si fondono, essere messo di buon umore dall'adulazione, rattristato dai propri errori, affascinato dal sesso, arrabbiato o depresso quando non riesce a ottenere ciò che vuole⁷.

Il ragionamento di Turing è stato sviluppato in varie direzioni⁸.

Nel 1956, l'inglese **William Ross Ashby**, psichiatra e pioniere della cibernetica, parlò di **“amplificazione dell'intelligenza”**:

Non è dunque impossibile che ciò che viene detto comunemente “capacità intellettuale” equivalga in ultima analisi a “capacità di giusta scelta”. In effetti, se una “scatola nera” parlante mostrasse di possedere così elevate capacità di scelta che, sottoponendole dei difficili problemi, essa desse frequentemente le risposte corrette, allora sarebbe difficile negare che quella scatola rappresenta l'equivalente, dal punto di vista del comportamento, di una “grande intelligenza”⁹.

L'idea fu ripresa dal polacco **Stanisław Lem**, il cui pensiero è molto interessante per noi, perché Lem era uno scrittore di fantascienza con una formazione scientifica¹⁰. Nel saggio *Summa technologiae. Scritti sul futuro*, tenendo insieme un non impeccabile ragionamento scientifico con la fantasia profetica e il raffinatissimo umorismo tipici della sua narrativa, intuì con

⁷ *The Mind of Mechanical Man*, in “British Medical Journal” è del 1949. Cito dalla citazione di Turing in *Macchine calcolatrici e intelligenza*, cit., p. 27.

⁸ Il termine “Intelligenza Artificiale” fu coniato nel 1955 dall'informatico statunitense John McCarthy, *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*, 31 agosto 1955, <http://jmc.stanford.edu/articles/dartmouth/dartmouth.pdf>.

⁹ William Ross Ashby, *Introduzione alla cibernetica*, a cura di Mauro Nasti, Einaudi, Torino, 1971 (edizione originale: 1956), p. 339.

¹⁰ Laureato in medicina, studiò poi la biologia e la cibernetica. È celebre soprattutto per il romanzo *Solaris*, da cui sono stati tratti due film, di Tarkowskij nel 1972 e di Soderbergh nel 2002.

notevole precisione quello che ci sta capitando oggi. Lem immaginò che l'amplificatore d'intelligenza:

Dovrebbe rappresentare nel campo delle attività cognitive il preciso corrispettivo di un potenziatore della forza fisica qual è ogni macchina guidata dall'essere umano [...]. La difficoltà fondamentale in cui s'imbatte qui il costruttore [...] è che egli deve costruire un apparecchio "più intelligente di lui". È chiaro che, se volesse agire secondo il metodo già usuale nella cibernetica, cioè predisporre un appropriato programma di azione per la macchina, non risolverà il compito affidatogli, perché un tale programma definisce già in partenza i limiti di "intelligenza" raggiungibili dal suo congegno. Apparentemente – ma solo apparentemente – il problema sembra un paradosso irrisolvibile, come ad esempio sollevare se stessi tirandosi per i capelli (e avendo per giunta legato ai piedi un peso di cento tonnellate). Ed effettivamente, almeno secondo i criteri odierni, se si postula la necessità di formulare una teoria, per forza di cose, matematica, prima di costruire l'amplificatore¹¹.

E aggiungeva:

Nella cibernetica si aggira oggi il mito medievale dell'*homunculus*, un'entità intelligente creata artificialmente. Il dibattito sulle possibilità di creare un cervello artificiale dotato delle caratteristiche della psiche umana ha già più volte attratto nella sua orbita filosofi e cibernetici¹².

Fino al 2000 circa, l'*homunculus* è rimasto appunto un mito, sottoposto alla ricerca teorica e a sperimentazioni *in vitro*, perché non c'era abbastanza potenza di calcolo per realizzare applicazioni pratiche. L'Intelligenza Artificiale ha prodotto dunque **applicazioni limitate a compiti specifici** nei quali una prestazione non distinguibile da quella di un essere umano o persino superiore non ha mai implicato l'ipotesi che la macchina fosse o potesse essere dotata di intelligenza, che pensasse.

Queste realizzazioni, come il Natural Language Processing (il software che analizza la lingua), eseguono automaticamente una procedura senza un intervento umano. Queste realizzazioni si chiamano **AI deterministiche**, perché **la prestazione è governata da algoritmi** che la rendono sempre uguale a se stessa, prevedibile.

Le AI deterministiche non hanno mai suscitato l'interesse generale, forse anche perché possono essere utilizzate solo dagli specialisti. Nemmeno applicazioni potenzialmente di interesse generale hanno fatto gridare al

¹¹ La *Summa technologiae. Scritti sul futuro* è del 1964 ed è ora finalmente disponibile in italiano, traduzione e cura di Luigi Marinelli, Luiss University Press, Roma, 2023, p. 134.

¹² Ivi, p. 130.

miracolo, nemmeno l'originalissimo progetto *Longplayer* del compositore inglese Jem Finer, celebre casomai per essere stato con Shane McGowan il principale autore delle canzoni dei Pogues. In occasione degli spettacoli organizzati a Londra nel Millennium Dome (oggi O2 Arena) per l'inizio del secondo millennio, Finer compose un brano musicale della durata di mille anni, *Longplayer* appunto. L'esecuzione, governata da algoritmi, è iniziata il 1° gennaio 2000 e terminerà senza interruzioni e ripetizioni il 31 dicembre 2999¹³.

Le AI Generative appaiono circa dal 2000. Queste AI rispondono non automaticamente, ma autonomamente a una richiesta dell'utilizzatore elaborando statisticamente le informazioni in loro possesso. Per questo, sono dette AI probabilistiche. **Le AI probabilistiche non sono prevedibili**, perché rispondono alla stessa domanda in modi diversi, e perché apprendono con l'esperienza e dunque mutano le loro prestazioni nel tempo.

Nemmeno le prime AI probabilistiche che sapevano scrivere, come i chatbot, hanno mai suscitato un grande interesse¹⁴. Non ha fatto molto scalpore nemmeno l'articolo scritto da GPT-3 *A robot wrote this entire article. Are you scared yet, human?* pubblicato dal Guardian l'8 settembre 2020 (dopo essere stato rimesso a posto dalla redazione)¹⁵.

Nel 2022 le GenAI sono state rese utilizzabili da tutti, quando la tecnologia GPT (Generative Pre-Trained Transformer) è stata potenziata e resa interrogabile in linguaggio naturale in conversazione, in chat: ChatGPT.

La lingua giunge al centro delle AI: **serve per l'input, serve per l'output**. Bisogna saper scrivere e bisogna saper leggere, nel senso di sapere come è stato scritto quello che si legge e sapere se quello che si legge è tutto quello che si potrebbe leggere, e se le funzioni della scrittura siano tutte contemplate dalle AI.

Il problema è grosso, ha interessato proprio tutti, e il dibattito si è infuocato anche al di fuori dell'ambito specialistico. Come sempre accade ai temi seri e nello stesso tempo di moda, sono intervenuti, oltre agli scienziati, anche politici¹⁶, imprenditori, giornalisti... content creator.

¹³ Si può ascoltare *Longplayer* dal sito <https://longplayer.org/>.

¹⁴ Ma le ha studiate dal nostro punto di vista Mirko Tavosanis, *Lingue e intelligenza artificiale*, Carocci, Roma, 2018.

¹⁵ <https://www.theguardian.com/commentisfree/2020/sep/08/robot-wrote-this-article-gpt-3>.

¹⁶ Regolamento (UE) 2024/1689 del Parlamento europeo e del Consiglio del 13 giugno 2024 che stabilisce regole armonizzate sull'intelligenza artificiale, https://eur-lex.europa.eu/legal-content/IT/TXT/PDF/?uri=OJ:L_202401689. Anche il nostro Parlamento sta per emanare una legge.

Non sono mancate semplificazioni eccessive, discorsi strampalati degli esperti della domenica, esagerazioni¹⁷ mirate a diffondere l'idea che l'intelligenza delle macchine sia o potrà essere uguale o addirittura superiore a quella umana e che la cosiddetta AI forte¹⁸, o AGI – Intelligenza Artificiale Generale¹⁹, non avrà limiti²⁰.

Al momento, l'AGI resta un'ipotesi, perché si fonda sulle basi incerte della limitata conoscenza che abbiamo della complessità dell'organismo umano. Che cosa siano la coscienza, l'intelligenza, da dove originino la comunicazione linguistica, la creatività ecc., sono temi aperti, su cui sappiamo qualcosa, non certamente tutto. Questo qualcosa ci dice, per esempio, che la nostra comunicazione si avvia per un impulso emotivo, che nel nostro ragionamento siamo più spesso razionali che logici e così via²¹. **Non**

¹⁷ Ottima la sintesi (e la presa di posizione) dell'informatico Walter Quattrociochi, *L'AI non è davvero intelligente, è una questione di statistica. Le allucinazioni? Inevitabili, è una caratteristica intrinseca*, in "Il Corriere della Sera", 3 giugno 2025:

https://www.corriere.it/tecnologia/25_giugno_03/l-ai-non-e-intelligente-e-una-questione-di-statistica-le-allucinazioni-inevitabili-e-una-caratteristica-intrinseca-9699a0b6-4952-4d69-bd8d-d4772bdb2xk.shtml.

¹⁸ I termini "AI forte" e "AI debole" sono stati introdotti da John Searle, *Minds, brains, and programs*, in "Behavioral and Brain Science", Volume 3, Issue 3, settembre 1980, pp. 417-424, <https://web.archive.southampton.ac.uk/cogprints.org/7150/1/10.1.1.83.5248.pdf>. Searle contrasta nettamente l'AI forte. Il celebre esempio della Stanza Cinese lo porta a concludere che: "As long as the program is defined in terms of computational operations on purely formally defined elements, what the example suggests is that these by themselves have no interesting connection with understanding". Tradotto da Graziella Tonfoni, *Menti, cervelli e programmi*, CLUP, Milano, 1984.

¹⁹ Un apostolo dell'AI forte è, per esempio, Mustafa Suleyman, *L'onda che verrà. Intelligenza Artificiale e potere nel XXI secolo*, Garzanti, Milano, 2024. Decisamente più prudente per l'epoca Nick Bostrom, *Superintelligenza. Tendenze, pericoli, strategie*, Bollati Boringhieri, Torino, 2014.

²⁰ Julian Nida-Rümelin prende di mira in modo molto duro la cosiddetta ideologia della Silicon Valley: "È connessa alla speranza di redenzione dell'America delle origini, ispirata al puritanesimo, la speranza di creare un mondo di puri e giusti che si siano lasciati alle spalle sporcizia e peccato [...]. Molti esperti di software della Silicon Valley annunciano la loro condizione di eletti [...]. Sulla Silicon Valley e sul lavoro legato all'AI viene caricata così una valenza metafisica [...]. Il concetto chiave in tutto questo è quello d'Intelligenza Artificiale (un concetto ricco di significati metafisici e teologici inespressi): un'entità capace di migliorarsi iper-razionale, sempre più considerata come se fosse dotata di un'anima". Julian Nida-Rümelin e Nathalie Weidenfeld, *Umanesimo digitale. Un'etica per l'epoca dell'Intelligenza Artificiale*, FrancoAngeli, Milano, 2018, pp. 19-21.

²¹ La bibliografia è vastissima. Qualche suggerimento: Giorgio Rizzolatti e Corrado Sinigaglia, *So quel che fai. Il cervello che agisce e i neuroni specchio*, Cortina, Milano, 2006; Daniel Kahneman, *Pensieri lenti e veloci*, Mondadori, Milano, 2012; Silvano Zipoli Caiani, *Corporità e cognizione. La filosofia della mente incorporata*, Mondadori Education - Le Monnier

assomigliamo a un software, che non ha emozioni ed elabora le informazioni per via matematica, statistica, logica, secondo **un'architettura progettata per analogia su quella poco conosciuta del sistema nervoso umano**.

A ben vedere, le AI deterministiche e quelle probabilistiche hanno caratteristiche comuni e **sono prodotti dell'essere umano** che:

- le progetta e le addestra per scopi specifici;
- le accende e le spegne;
- dà l'autorizzazione affinché il prodotto abbia effetti sul mondo.

Quanto alle differenze, oltre a quelle di cui abbiamo già parlato, le GenAI si distinguono dalle altre perché hanno costi di realizzazione e di gestione straordinariamente superiori. Per dare un'idea, per addestrare e usare gli LLM più potenti, la potenza di calcolo (espressa in GPU – Graphics Processing Unit) e l'energia elettrica costano da decine di migliaia a miliardi di dollari. Le aziende stanno rifacendo i conti e sviluppando tecnologie più sostenibili. Se ne vedono già i primi esiti.

1.3. Chi e come istruisce le AI

I sistemi di Intelligenza Artificiale sono estremamente complessi, progettati da decine o centinaia o migliaia di scienziati, ricercatori, tecnici. Sono loro che addestrano le AI automatiche e autonome affinché possano svolgere le attività per le quali sono progettate.

Per l'addestramento c'è bisogno di un corpus o dataset, cioè una amplissima e rappresentativa collezione di testi (in questo libro ci occupiamo solo di lingua, non di immagini²² ecc.).

Alessandro Panunzi definisce un corpus linguistico “una raccolta strutturata di eventi comunicativi prodotti in ambiente naturale e selezionati sulla base di criteri espliciti al fine di rappresentare una lingua o una sua specificità”²³.

Università, Firenze-Milano, 2016; Marcello Frixione, *Come ragioniamo*, Laterza, Bari-Roma, 2007; Francesco F. Calemi, *Argomentare, dimostrare, confutare. Un'introduzione alla logica*, Carocci, Roma, 2022. Per approfondire si possono seguire le bibliografie indicate.

²² Per l'annotazione delle immagini rimando a Fei Fei Li, *Tutti i mondi che vedo. Curiosità, scoperta e meraviglia all'alba dell'intelligenza artificiale*, Luiss University Press, Roma, 2024. Molto utile per chi va di fretta la splendida visual story di Christo Buschek e Jer Thorp, *Models all the way down*, <https://knowingmachines.org/models-all-the-way>.

²³ Emanuela Cresti e Alessandro Panunzi, *Introduzione ai corpora dell'italiano*, Il Mulino, Bologna, 2013, p. 52. Nel primo e quinto capitolo (di Emanuela Cresti) si leggono un utile

Long story short: la rappresentatività del corpus è data dalla coerenza dei dati. Un corpus composto, poniamo, da articoli di giornale, letteratura e leggi potrebbe forse rappresentare la lingua dal punto di vista delle sue strutture, ma certo non la lingua utilizzata in un contesto preciso, poniamo il contesto aziendale. Per non dire della lingua parlata.

Oltre che essere strutturato e coerente, **un corpus deve anche essere annotato**. Annotare significa attribuire un tag, cioè un'etichetta che definisce il valore di ogni unità linguistica. I tag per l'annotazione linguistica sono sigle convenzionali²⁴ (il tag per un sostantivo è NOUN), non da tutti seguite però.

In linguistica computazionale, **le unità linguistiche si chiamano token**. Il token indica sia le parole sia le non parole, cioè quella classe disomogenea composta da segni di punteggiatura, numeri, sigle e altri simboli. Il token non rappresenta quindi il numero di parole di cui è composto un corpus; tuttavia, si può per approssimazione dire che a 4 token corrispondono circa 3 parole o, in altri termini, che una parola corrisponde a circa 0,75 token.

L'annotazione serve a:

- istruire una nuova AI e renderla operativa;
- migliorare un'AI già operativa;
- istruire un'AI già operativa a svolgere funzioni particolari.

L'annotazione è un lavoro essenzialmente manuale, con qualche eccezione. Per esempio, per addestrare un nuovo NLP si possono usare i numerosissimi corpus già annotati. Lo stesso vale per le GenAI: le versioni successive sono addestrate anche con le versioni precedenti.

Gli NLP annotano automaticamente il testo dell'utilizzatore. Per ogni token sappiamo, per esempio:

- se è una parte del discorso (Part of Speech, POS), come un articolo (DET), un sostantivo (NOUN) ecc., o un segno di punteggiatura (PUNCT) ecc.;

riepilogo della storia dei corpus (dal Brown corpus alle pionieristiche realizzazioni del gesuita italiano Roberto Busa) e una descrizione dei corpus dell'italiano allora disponibili. La piattaforma Sketch Engine raccoglie e mette a disposizione centinaia di corpus, gratuiti e a pagamento: <https://www.sketchengine.eu/>, tra i quali anche il Brown University Standard Corpus of Present-Day American English, il primo corpus linguistico, realizzato con un milione di parole nel 1961 (<https://www.sketchengine.eu/brown-corpus/>).

²⁴ Lo standard per l'annotazione linguistica è l'Universal Dependency: <https://universaldependencies.org>.

- le sue caratteristiche morfologiche generali (Full Morpho): l'articolo è determinativo (RD) o indeterminativo (RI) e può essere singolare o plurale, maschile o femminile;
- le sue caratteristiche morfologiche in quella specifica frase (Selected Morpho): l'articolo è maschile singolare;
- la sua funzione nell'albero sintattico, cioè nella struttura del periodo (Dependency Parsing): soggetto (nsubj), oggetto (obj), complementi indiretti (nmod) ecc.

Facciamo un esempio con la mia frase precedente: “Il token non rappresenta quindi il numero di parole di cui è composto un corpus; tuttavia, si può per approssimazione dire che a 4 token corrispondono circa 3 parole o, in altri termini, che una parola corrisponde a circa 0,75 token”.

Ho fatto analizzare la frase al *CoreNLP* dell'Università di Stanford riaddestrato per l'italiano e rinominato *Tint* dalla Fondazione Bruno Kessler di Trento²⁵. Ho utilizzato la versione demo (gratuita) nella quale è sufficiente copiare e incollare il testo per avere una visualizzazione passabilmente comprensibile del testo annotato (Fig. 1).

Ma se si vuole lavorare con i NLP, bisogna dialogare nei suoi linguaggi. Il più comune per interrogarli è Python. La richiesta di analizzare la mia frase si scrive così:

```
sentence = ("Il token non rappresenta quindi il numero di parole di cui è composto
un corpus, "
"tuttavia si può per approssimazione dire che a 4 token corrispondono circa 3
parole o, "
"in altri termini, che una parola corrisponde a circa 0,75 token.")
```

```
# Analizza la frase
```

```
doc = nlp(sentence)
```

```
# Stampa token, lemma, parte del discorso e dipendenza sintattica
```

```
print(f'{"Token":<15} {"Lemma":<20} {"POS":<10} {"Dipendenza"}')
```

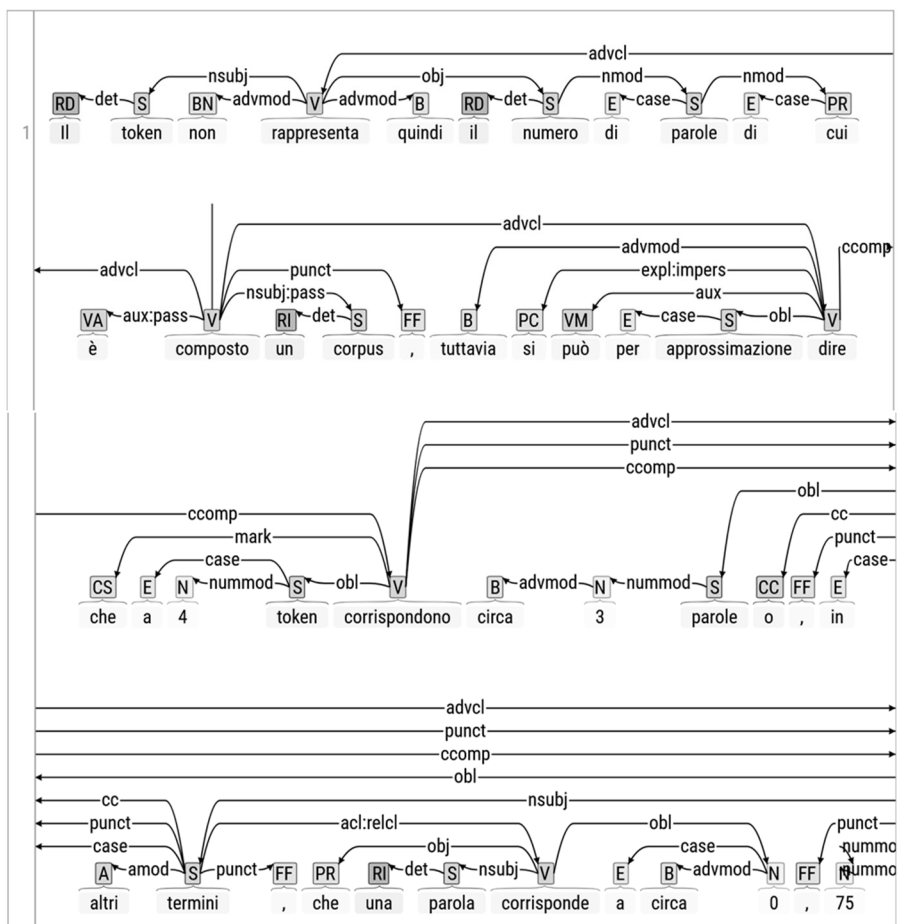
```
print("-" * 60)
```

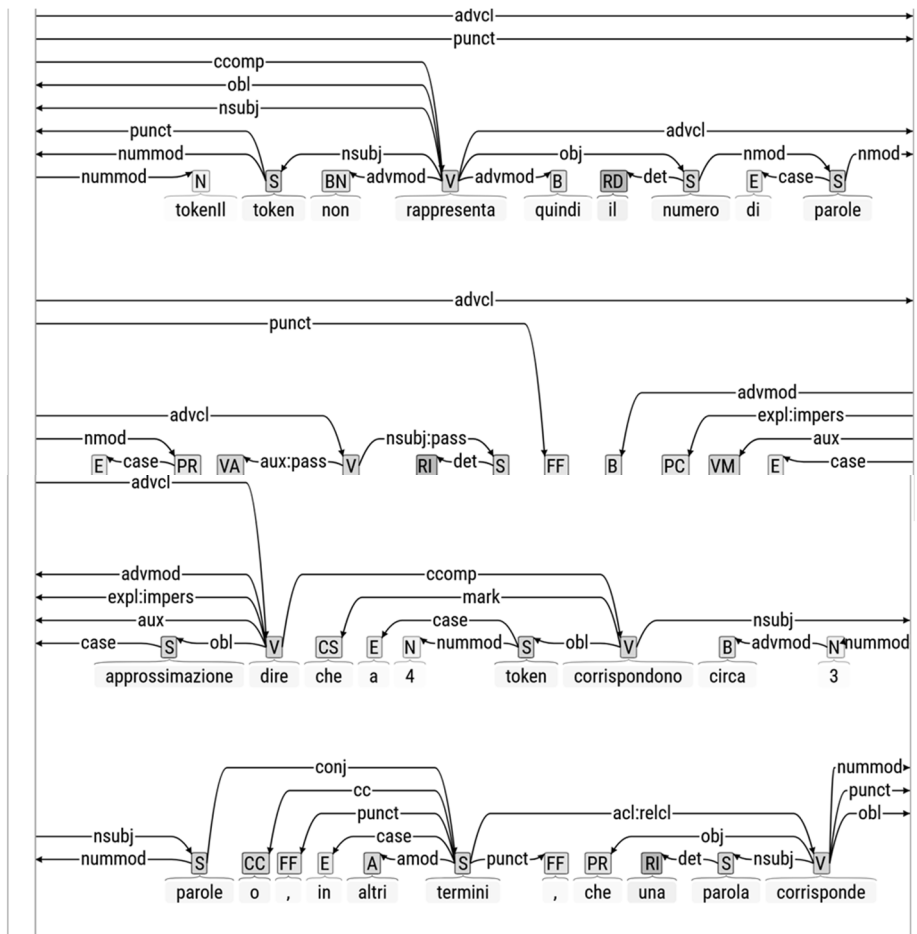
```
for token in doc:
```

```
    print(f'{"token.text":<15} {"token.lemma_":<20} {"token.pos_":<10} {"token.dep_"}')
```

²⁵ <https://dh.fbk.eu/research/tint/>.

Fig. 1 – Visualizzazione del testo annotato





E il NLP risponde in json (o Penn Treebank o altri linguaggi, secondo il NLP) che, per fare un solo esempio, così descrive la parola “rappresenta”:

```
{
  "index": 4,
  "word": "rappresenta",
  "originalText": "rappresenta",
  "lemma": "rappresentare",
  "characterOffsetBegin": 13,
  "characterOffsetEnd": 24,
  "pos": "V",
  "featuresText": "Mood=Ind|Number=Sing|Person=3|Tense=Pres|VerbForm=Fin",
  "ner": "O",
  "isMultiwordToken": false,
  "isMultiwordFirstToken": false,
  "ud_pos": "VERB",
  "full_morpho": "rappresenta rappresentare+v+indic+pres+nil+3+sing rappre-
sentare+v+imp+pres+nil+2+sing",
  "selected_morpho": "rappresentare+v+indic+pres+nil+3+sing",
  "guessed_lemma": false,
  "features": {
    "Tense": [
      "Pres"
    ],
    "VerbForm": [
      "Fin"
    ],
    "Mood": [
      "Ind"
    ],
    "Number": [
      "Sing"
    ],
    "Person": [
      "3"
    ]
  }
},
```

Chi vuole rovinarsi gli occhi può avventurarsi nella giungla del Dependency parsing. Le dipendenze sintattiche sono visualizzate così (trascrivo solo i primi 7 elementi su oltre 130):

```

"basic-dependencies": [
{
  "dep": "ROOT",
  "governor": 0,
  "governorGloss": "ROOT",
  "dependent": 13,
  "dependentGloss": "composto"
},
{
  "dep": "det",
  "governor": 2,
  "governorGloss": "token",
  "dependent": 1,
  "dependentGloss": "Il"
},
{
  "dep": "nsubj",
  "governor": 4,
  "governorGloss": "rappresenta",
  "dependent": 2,
  "dependentGloss": "token"
},
{
  "dep": "advmod",
  "governor": 4,
  "governorGloss": "rappresenta",
  "dependent": 3,
  "dependentGloss": "non"
},
{
  "dep": "advcl",
  "governor": 13,
  "governorGloss": "composto",
  "dependent": 4,
  "dependentGloss": "rappresenta"
},
{
  "dep": "advmod",
  "governor": 4,
  "governorGloss": "rappresenta",
  "dependent": 5,
  "dependentGloss": "quindi"
}
]

```

```

},
{
  "dep": "det",
  "governor": 7,
  "governorGloss": "numero",
  "dependent": 6,
  "dependentGloss": "il"
},

```

Come vedremo meglio nel paragrafo 1.5, per addestrare le GenAI a scrivere e a elaborare risposte sono stati utilizzati ormai tutti o quasi tutti i corpus già annotati manualmente o automaticamente. Siccome però le GenAI apprendono continuamente e possono continuamente migliorare le loro prestazioni, c'è sempre bisogno di annotazioni manuali. Per esempio, l'annotazione automatica di un NLP non ci dice se la frase "Che meraviglia la grammatica!" vada intesa alla lettera o ironicamente. C'è bisogno di una annotazione manuale.

L'annotazione manuale delle GenAI è oggi compiuta da **eserciti di annotatori** che svolgono a cottimo questo lavoro certosino per conto di società che operano per conto delle BigTech (il compenso supera raramente i 10 dollari all'ora nei paesi con un'economia come la nostra)²⁶.

Alcune BigTech, poi, offrono direttamente ai clienti un servizio di annotazione specifico per i loro dati, come per esempio Amazon Mechanical Turk (<https://www.mturk.com/>): il Turco Meccanico. Questo nome viene da lontano, dall'invenzione dell'ungherese Wolfgang von Kempelen che nel 1770 creò una specie di gioco di prestigio: una scatola meccanica che batteva gli esseri umani a scacchi. Per molti anni fu anche una macchina da soldi: il pubblico pagava per assistere alle dimostrazioni. Ma alla fine si scoprì che nella scatola era nascosto un giocatore vero, un turco di bassa statura²⁷. Ecco, gli annotatori sono... piccoli turchi. Bah!

Comunque, **la maggior parte degli annotatori lavora gratis**, senza sapere che sta lavorando, senza sapere che sta annotando. Siamo noi.

Stiamo annotando ogni volta che selezioniamo le immagini di un semaforo o di un autobus o di un cane per superare il test del reCAPTCHA di Google se siamo esseri umani. Stiamo annotando ogni volta che nei social

²⁶ Kate Crawford, *Né intelligente né artificiale. Il lato oscuro dell'AI*, Il Mulino, Bologna, 2021 conduce un discorso molto critico al riguardo, ed anzi lo estende allo sfruttamento generale delle risorse che le GenAI compirebbero per "servire gli interessi dominanti" (p. 16).

²⁷ La storia è stata raccontata a modo suo da Edgar Allan Poe, *Il giocatore di scacchi di Maelzel* (varie edizioni). Maelzel è il nome dell'inventore tedesco che acquistò la scatola per esibirla negli USA.

come Facebook o Instagram pubblichiamo un selfie, un post, un commento a un post.

Non solo. Quando ci siamo iscritti al servizio abbiamo acconsentito esplicitamente a cedere all'azienda proprietaria del servizio il nostro diritto di sfruttamento delle nostre immagini e delle nostre parole.

Come dice Jaron Lanier, che interpreta se stesso nel film *The Social Dilemma*²⁸: “Se il servizio è gratis, il prodotto sei tu”²⁹.

Il **machine learning** è la fase in cui il corpus annotato viene inviato al software affinché impari a riconoscere e a classificare i tag, e quindi ad apprendere quello che deve apprendere (la lingua, le immagini ecc.). Una volta che l'abbia appreso, il software elabora a modo suo un proprio modello di comportamento secondo il quale eseguirà le attività che deve eseguire elaborando però dati reali, cioè i dati non annotati che l'utilizzatore reale gli chiederà di elaborare. **Nessuno può sapere come elabori il suo proprio modello e in base a quali misteriosi meccanismi funzioni.**

Merita una pausa di riflessione questo mistero inquietante. Torniamo alla profezia di Lem sull'amplificatore di intelligenza (anno 1964):

Esiste però – per ora solo come possibilità ipotetica – un approccio completamente diverso alla faccenda. Una conoscenza approfondita sulla struttura interna dell'amplificatore non ci è accessibile. Forse è addirittura superflua. Basta trattarlo come una “scatola nera”, come un dispositivo di cui non abbiamo la più pallida idea né del funzionamento interno e né degli altri stati, perché ci interessano esclusivamente i risultati finali della sua azione. Quell'amplificatore, come ogni apparecchiatura cibernetica che si rispetti, possiede un'“entrata” e un'“uscita”. Nel mezzo si stende l'area della nostra ignoranza [...]. Noi stessi siamo delle “scatole nere”. I nostri corpi ci sono assoggettati, possiamo dar loro certi comandi, eppure non conosciamo (o meglio: non dobbiamo conoscere, nel senso che tale conoscenza non è indispensabile) la loro organizzazione interna. Rieccoci, dunque, al problema del saltatore che sa saltare, ma non sa in che modo lo fa, cioè non possiede la conoscenza della dinamica dei percorsi neuro-muscolari il cui risultato è il salto. E dunque un perfetto esempio di un dispositivo del quale possiamo far uso pur non conoscendone l'algoritmo è ogni essere umano. Uno dei “dispositivi” in tutto l'universo “a noi più vicini” è proprio il nostro cervello: ce l'abbiamo addirittura in testa. Nondimeno fino a oggi non sappiamo esattamente come funzioni. [...]. La Scatola Nera, in quanto sistema molto complesso, è indescrivibile: il suo algoritmo non è conosciuto né può conoscerlo nessuno, agisce secondo il metodo delle probabilità, e quindi, messa due volte di

²⁸ Prodotto da Netflix nel 2020 con la regia di Jeff Orlowski.

²⁹ Si tratta forse della parafrasi del titolo del suo libro *You Are Not a Gadget: a Manifesto*, Penguin Books, Londra, 2010, traduzione italiana *Tu non sei un gadget*, Mondadori, Milano, 2010.

fronte alla stessa situazione, non deve affatto comportarsi allo stesso modo. Inoltre, la Scatola Nera – ed è questo l'aspetto fondamentale – è una macchina che nel corso di attività concretamente intraprese impara sui propri stessi errori”³⁰.

Il machine learning non garantisce che i risultati ottenuti *in vitro* durante l'addestramento si otterranno anche con dati reali in situazioni reali, ma quanto più i dati del corpus sono rappresentativi dei dati reali, tanto più il software sarà preciso.

Per un machine learning ben fatto, **il corpus va diviso in due**: una parte per l'addestramento, l'altra parte per testare l'apprendimento.

Il machine learning può essere supervisionato, non supervisionato o per rinforzo.

Il machine learning è supervisionato quando si forniscono al software sia gli esempi sia i risultati che si vogliono ottenere. Poi si controllano i risultati dell'apprendimento con dati reali, ed eventualmente si apportano delle correzioni.

Il machine learning è non supervisionato quando non si forniscono al software i risultati che si vogliono ottenere. A volte perché non sono noti, altre volte perché ci si fida del software (ingegneri e informatici tendono a fidarsi più di altri specialisti) o perché si preferisce risparmiare tempi e costi della supervisione.

Il machine learning è per rinforzo quando l'apprendimento avviene in continue sessioni (iterazioni) in cui il software tenta di risolvere un problema, sbaglia e riprova fino a quando impara (questo machine learning viene spesso utilizzato nelle applicazioni per le aziende, forse con lo scopo di risparmiare, scopo che però fallisce quando le iterazioni sono troppe).

Al termine del machine learning, il software ha elaborato un proprio modello di comportamento con il quale è pronto per affrontare i dati reali degli utilizzatori (ci tengo a ricordare che nessuno può sapere come il software abbia elaborato questo modello né come funzioni. Del modello si vedono solo i risultati).

Nessun tipo di machine learning può garantire a nessuna applicazione AI di non commettere errori. **L'errore dell'AI va postulato.**

1.4. Il Natural Language Processing

Come abbiamo visto, il Natural Language Processing restituisce informazioni sulla lingua in se stessa. Queste informazioni interessano solo ai

³⁰ *Summa technologiae*, cit., p. 134, 139 e 144.

linguisti. Per offrire applicazioni utili e applicabili ai casi e ai problemi reali, i NLP devono essere nuovamente addestrati a riconoscere e a gestire usi particolari della lingua. Per esempio, individuare in un archivio aziendale complesso e destrutturato tutti i testi che hanno determinate caratteristiche linguistiche o determinati significati (Natural Language Understanding).

Per un nuovo machine learning ci vuole un nuovo corpus che deve essere annotato con nuovi tag specifici che definiscono gli aspetti linguistici di interesse.

La qualità del modello elaborato dal NLP si valuta secondo due metriche: la Precision e la Recall. Per capire le metriche è necessario premettere che i risultati dell'analisi automatica sono di quattro tipi:

- **veri positivi:** il dato viene correttamente etichettato come corretto. L'utilizzatore vede una informazione corretta. Per esempio, se cerco "contratto", il sistema visualizza "contratto";
- **falsi positivi:** il dato viene scorrettamente etichettato come corretto. L'utilizzatore vede una informazione scorretta. Per esempio, se cerco "contratto", il sistema visualizza "contatto";
- **veri negativi:** il dato viene correttamente etichettato come scorretto. L'utilizzatore non vede una informazione scorretta. Per esempio, se cerco "contratto", il sistema esclude "contatto";
- **falsi negativi:** il dato viene scorrettamente etichettato come scorretto. L'utilizzatore non vede una informazione corretta. Per esempio, se cerco "contratto", il sistema esclude "contratto".

La Precision è il rapporto tra tutte le etichettature corrette e il totale dei casi da etichettare. Formula matematica:

$$\frac{(\text{veri positivi} + \text{veri negativi})}{(\text{veri positivi} + \text{falsi positivi} + \text{veri negativi} + \text{falsi negativi})}$$

La Recall è la percentuale di veri positivi. Formula matematica:

$$\text{veri positivi} \div (\text{veri positivi} + \text{falsi negativi})$$

Precision e Recall possono essere migliorate ma, generalmente, l'una a scapito dell'altra. Quale privilegiare, dipende dal progetto (nel Box 1 racconto una mia esperienza).