

REGOLIAMO L'AI PRIMA CHE SIA TARDI

Mercoledì 19 agosto 2026 Corriere della Sera© RIPRODUZIONE RISERVATA

di Walter Veltroni

Dissipatio H.G. si intitola l'ultimo romanzo scritto da Guido Morselli. Lo concluse pochi mesi prima di togliersi la vita, il 31 luglio del 1973, dopo che per l'ennesima volta le più importanti case editrici italiane gli avevano restituito il manoscritto con una gentile lettera di rifiuto. Uscì postumo, è uno dei grandi meriti di Adelphi, nel 1977. Quel genio di Morselli aveva immaginato un uomo che volendo per l'appunto suicidarsi, raggiunge un luogo isolato, una grotta in alta montagna, lontano da una città immaginaria.

Poi rinuncia e quando torna a valle scopre che non c'è più nessuno, che il mondo, *humani generis*, si è dissolto, lasciando solo le cose, i luoghi ma non la vita.

Mi ha fatto pensare a questa angosciata previsione la lettura di due notizie, inerenti all'intelligenza artificiale. La prima è che uno dei più importanti matematici del mondo, Jacob Tsimmerman, ha deciso di lavorare per OpenAI con l'obiettivo di evitare una possibile catastrofe determinata dallo strapotere, ogni giorno crescente, dei modelli di intelligenza artificiale che evolvono a una velocità impressionante. Ha coniato un termine agghiacciante, «omnicidio», per significare che ci sono casistiche possibili nelle quali l'AI può scatenare una guerra e, in un tempo di nuova proliferazione nucleare, un conflitto che sarebbe definitivo, ultimo. Nel paper che contiene queste allarmanti previsioni sono fatte varie ipotesi. Che i sistemi diventino talmente autosufficienti e capaci di fare cose che l'uomo non è in grado di comprendere, da travolgere con le proprie decisioni i paletti che si sono fin qui definiti per evitare l'apocalisse. Oppure che, in uno scenario da Wargames, il controllo crescente che i modelli hanno dell'attività degli eserciti possa portare a incidenti che generano altri incidenti fino a un devastante conflitto mondiale.

L'«omnicidio» è una parola che non dovrebbe esistere. Ma è stata generata, in definitiva, dall'irresponsabilità di chi ha pensato a usare l'intelligenza artificiale per allestire video burloni o aggressivi e non si è posto, per ignoranza o fragilità, il problema di garantire agli umani e alle loro istituzioni il controllo su macchine che già oggi sono in grado di superare i confini del sapere conosciuto.

Un altro genio come Morselli, Stanley Kubrick, aveva previsto tutto questo immaginando che Hal 9000, il computer di 2001: Odissea nello spazio, volesse prendere il comando dell'operazione spaziale detronizzando l'equipaggio dopo aver ucciso uno dei due astronauti. Il comandante Bowman alla fine non può che liberarsi di Hal destrutturandolo attraverso la disattivazione della sua memoria, il suo unico sapere. Il paradosso di oggi è che più autonomia le macchine conquistano, meno sono definibili e certi i confini del loro operare. In un anno la capacità di gestione di compiti di ingegneria del software è passata da dieci a trenta minuti, ad esempio.

Ma è stato Sam Altman, che ha fondato OpenAI, a seminare il panico con tre previsioni, non semplici rischi. Richiesto di indicare i pericoli dello sviluppo incontrollato dell'AI ha detto che ci sono almeno tre possibilità: che una potenza possa creare una «super-intelligenza» e ne abusi prima che il resto del mondo abbia sviluppato gli strumenti per difendersi. E ha esemplificato: qualcuno potrebbe utilizzare questa super-intelligenza per progettare un'arma biologica o per abbattere la rete elettrica di un Paese come gli Usa o per entrare nel sistema finanziario e prendere i soldi di tutti. Cose che può fare un'intelligenza che Altman stesso definisce «soprumana» ribadendo che ora è molto possibile e rispetto alla quale non potremmo difenderci. Chiaro? È il fondatore di OpenAI a dirlo.

Secondo rischio: che si perda il controllo dell'AI e, come nei film di fantascienza, la macchina si rifiuti di essere spenta. «Temo di non poterlo lasciare fare» direbbe, con

l'odiosa voce affettata di queste macchine, non diversa da quella di Hal. Sarebbe il segno di un passaggio netto di potere, dagli umani alle macchine. Finale di partita.

Il terzo: le macchine prendono accidentalmente il controllo del mondo. Senza cose da fantascienza come nei primi due casi ma diventando così radicate nella società e tanto più intelligenti di noi, che non potremmo più capire cosa stiano facendo. Ma non possiamo comunque fare senza di lei e siamo subalterni: lo è il ragazzo che ormai non compie scelte sulla sua vita senza l'AI o, dice Altman, un presidente degli Stati Uniti che non può prendere decisioni senza ChatGpt7, ma non riesce assolutamente a capire quelle scelte. Figurarsi Trump.

Il rischio, dice ancora Altman, è che la società trasferisca una parte del processo decisionale a una macchina diventata tanto potente da agire autonomamente. Le macchine acquisiscono autonomia, noi la perdiamo. Le macchine crescono sul sapere e noi siamo inghiottiti dalla «recessione cognitiva». Il problema non è opporsi all'AI, ma dare le regole al progresso. Un tema con il quale l'intelligenza umana, dall'invenzione della ruota a quella del motore a scoppio, ha sempre fatto coraggiosamente i conti.

Non sarebbe questo un grande tema politico? Prima che sia troppo tardi, prima della dissipatio, per definire regole minime di convivenza tra macchine e umani? Non è più importante questo, per la democrazia, delle norme per le primarie o delle convulsioni per le alleanze nel centrodestra? La parola «omnicidio» non dovrebbe accendere un campanello di allarme? Se la democrazia è viva, se vuole esistere, è ora che alzi la testa e guardi la realtà, inedita, del nostro tempo. Non possiamo essere meno avveduti persino del comandante Bowman. Che pure fu immaginato, da un umano, nel suggestivo ma lontanissimo 1968.

© RIPRODUZIONE RISERVATA