

Lavoro e macchine

di Claudio Sossai

Riproponiamo su *ANTIGONE* questo contributo di Claudio Sossai apparso sul numero 27-28 della rivista *I Maggio*, convinti che la limitata diffusione delle due riviste ci mette al riparo dall'accusa che una ripubblicazione a così breve distanza sia un eccesso e, nello stesso tempo, ci consente di dare conto ad un pubblico specializzato di un dibattito che ha appena preso l'abbrivio e che ha tutti i numeri per esser interessante. In un tempo in cui, per dirla col filosofo Vattimo, l'eccesso di informazione produce dibattiti sul malinteso, *repetita juvant*. Almeno ce lo auguriamo.

1. Le macchine

Schematizzando, possiamo dividere le macchine in due gruppi:

- 1) macchine termiche
- 2) macchine logiche.

La *macchina termica* è quella della vecchia manifattura, quella che sostituisce il lavoro umano come fatica fisica; per essa, la definizione fisica di lavoro ("forza x spostamento") è adeguata. In quanto alla *macchina logica*, la si può concepire come un calcolatore programmato e se ne può dare una descrizione matematica tanto semplice quanto feconda.

L'idea risale a Turing, e la forma in cui qui la presentiamo è stata data da Davis nel 1958¹.

La "macchina di Turing", rappresenta il concetto di calcolabilità stesso. La macchina di Turing è fatta da:

- 1) un nastro diviso in cellette di uguale grandezza che scorre da destra a sinistra o viceversa, *senza limiti*.
- 2) un congegno che, con questo nastro, può compiere le quattro seguenti operazioni elementari:
 - a) può scrivere 1 su una celletta del nastro
 - b) può cancellare una celletta dove compare 1 o lasciarla inalterata
 - c) può spostare il nastro di un posto verso destra
 - d) può spostare il nastro di un posto verso sinistra.

Ogni operazione richiede lo stesso tempo e viene chiamata *gradino* o *passo* nel lavoro della macchina (Cfr. Appendice 1).

Ogni operazione è controllata da un'istruzione e quindi una macchina è sostanzialmente identificabile con l'insieme di istruzioni che la definiscono. Per farla funzionare basta inserirvi un nastro con tanti 1 quanto è l'intero che le si vuole dare come input, far partire il meccanismo

e, se e quando si ferma, leggere quanti 1 ci sono scritti sul nastro: questo è il risultato del lavoro della macchina.

Non si tratta di un meccanismo molto complicato, come si può vedere; pur tuttavia si può dimostrare che tutto quello che può fare un calcolatore elettronico, può essere svolto dalla nostra "macchinetta". Qui abbiamo quindi una buona descrizione di quelli che sono i sistemi a controllo numerico e i calcolatori programmati possibili. Non solo, *ma possiamo costruire una Macchina di Turing che genera, una dopo l'altra, tutte le Macchine di Turing possibili*, nel senso che scrive, uno dopo l'altro tutti i programmi, in una lista (purtroppo) infinita.

Abbiamo così descritto tutte le macchine e possiamo allora dare una definizione che ci sarà molto utile per i prossimi sviluppi del discorso: chiamiamo *lavoro computabile*, il lavoro che può essere fatto o da una macchina termica o da una o più macchine termiche controllate da una macchina logica o, infine, da una macchina logica. La definizione così precisa di macchina ci permette di ottenere un grosso risultato: determinare ciò che le macchine -qualsiasi tipo di macchina di cui oggi possiamo disporre- non possono fare.

L'Appendice 2 fornisce un esempio formalmente semplice (concettualmente un po' meno) di un problema che nessuna macchina è in grado di risolvere.

Sui limiti delle macchine si è scritto molto dal Teorema di Gödel in poi². Banalizzando molto, si potrebbe dire che qualsiasi sistema meccanico non è in grado di controllare completamente linguaggi sufficientemente "forti" per parlare di se stessi.

La struttura delle dimostrazioni di cui si è parlato, infatti, si rifà, anche se in maniera molto più sofisticata, a quello che è un vecchio problema logico: il *paradosso*. La questione è nota; un esempio di paradosso è la frase: "io dico che mento".

1) DAVIS, MARTIN *Computability and unsolvability*, McGraw-Hill Book Company, New York.

2) Si può dire che due sono i grossi risultati a cui fino ad oggi è pervenuta la logica formale:

i) la scoperta di un linguaggio simbolico all'interno del quale si possono esprimere tutti gli enunciati e i teoremi della matematica.

ii) la definizione all'interno di questo linguaggio di un metodo combinatorio di prova degli enunciati veri. In base a questi risultati i matematici cominciarono a cercare una procedura meccanica che fosse in grado di trovare automaticamente gli enunciati veri all'interno di questo linguaggio. Questo era il vecchio sogno di Leibnitz quando cercava di definire il *calculus ratiocinator*.

Gödel dimostrò nel 1930 l'impossibilità di definire un procedimento automatico per la prova degli enunciati veri dell'aritmetica. Nel fare ciò Gödel definì un'idea di procedimento di calcolo automatico che poi si vide essere equivalente alla definizione data da Turing e che noi qui abbiamo presentato.

Decidere se chi pronuncia quella frase mente o no è difficile perchè:

- a) se mente ha detto il vero, quindi non mente.
 - b) se ha detto il vero, cioè se non mente, allora mente.
- Ciò che, a conclusione di queste prime note, ci interessava mettere in evidenza è l'esistenza di una possibile attività che una macchina non è in grado di compiere. Chiameremo questo lavoro non computabile: *lavoro innovativo*.

2. Il lavoro innovativo

La definizione finora ottenuta di lavoro innovativo è una definizione in negativo: tutto ciò che non è computabile. Ciò che ci proponiamo nelle prossime righe è di descriverne alcune caratteristiche in positivo.

La prima osservazione che ci viene spontanea è che man mano che il lavoro si astrae dalla produzione di beni fisici, si avvicina sempre più alla produzione di 'conoscenza'.

Partiamo da quest'osservazione: 'conoscere qualcosa' di un fenomeno vuol dire enunciare una teoria in grado di prevedere alcuni dati relativi al fenomeno. Se la teoria non era stata precedentemente enunciata, la si chiama *innovazione*.

La "bontà" di questa teoria è misurata dalla sua capacità di diminuire l'incertezza relativa al fenomeno, prevedendone i tempi di evoluzione: "vedendo" prima alcuni dati che, successivamente, l'osservazione del fenomeno avrebbe generato³.

Detta così, la cosa è ancora troppo vaga per poterci dare delle indicazioni utili.

Per approfondire la questione bisognerà soffermarsi più dettagliatamente su una caratteristica del processo conoscitivo, cioè la sua capacità di *previsione*; ovvero: com'è possibile che una teoria descriva un'insieme di dati prima che essi vengano, per così dire, osservati.

Gli ingredienti del problema sono: cosa si intende per teoria; cosa sono i dati; cosa è il prima e il poi, cioè lo scorrere del tempo.

Una *teoria* è un insieme di enunciati di un linguaggio, derivati da un sottoinsieme proprio di enunciati (quelli che di solito si dicono assiomi), secondo le regole accettate dalla sintassi logica.

Un *dato* è un'osservazione codificata di un enunciato. Tentare di dire che cosa è lo scorrere del tempo sarebbe presuntuoso. Ci limiteremo quindi a poche osservazioni che ci saranno sufficienti a portare avanti il filo del ragionamento che ci interessa. La verifica dello scorrere del tempo è l'osservazione di un fenomeno che si ripete uguale a se stesso. Prendiamo ad esempio, il pendolo: la misura del tempo è data dal ritorno di un corpo nella stessa posizione ad intervalli regolari. Ma cosa significa "stessa posizione"? Non lo è in assoluto, visto che comunque il pendolo si muove solidale a un corpo (la Terra) che non è fermo; lo è però rispetto ad un osservatore. E' quindi l'osservatore a compiere quest'operazione di uguaglianza: rende uguale la posizione del corpo ad intervalli regolari.

E questa è un'operazione relativa al nostro modo di conoscere, o meglio: al nostro modo di usare-costruire il linguaggio (per esempio: quando diciamo: " $A=B$ "). Osserva G. Frege: "L'analisi dell'uguaglianza ci conduce a riflettere su alcuni problemi che si connettono ad essa e sono ben difficili a risolversi".

Dobbiamo vedere nell'uguaglianza un rapporto? E precisamente di che tipo? Un rapporto fra oggetti, ovvero un rapporto fra nomi (o segni) di oggetti? In un precedente lavoro mi ero pronunciato in favore di quest'ultima soluzione (che l'uguaglianza sia un rapporto fra nomi). Ecco il principale motivo che sembra a favore di essa: $a=a$ e $a=b$ sono due proposizioni di valore conoscitivo diverso, poichè $a=a$ vale a priori e deve chiamarsi, secondo Kant, analitica; mentre proposizioni della forma $a=b$ contengono spesso ampliamenti notevoli della nostra conoscenza e non sempre possono venire fondate a priori. (Per esempio che ogni mattina non sorge un nuovo sole, ma sempre il medesimo, è stata senza dubbio una delle più feconde scoperte dell'astronomia. E oggi ancora il riconoscere che in diverse osservazioni abbiamo a che fare con lo stesso piccolo pianeta o con la stessa cometa è talvolta tutt'altro che facile).

Orbene: se volessimo vedere nell'uguaglianza un rapporto effettivo, non fra nomi " a " e " b ", ma fra gli oggetti da essi designati, scomparirebbe ovviamente ogni diversità fra le due proposizioni " $a=a$ " e " $a=b$ ", nel caso - ben inteso - che l'oggetto " a " sia proprio uguale all'oggetto " b ". In tal caso infatti l'uguaglianza esprimerebbe un rapporto di un oggetto con sè stesso e precisamente un rapporto "sui generis" in cui ogni cosa sta con sè medesima, ma nessuna con un'altra. Ciò che si vuol dire con la proposizione " $a=b$ " sembra dunque essere questo: che i segni (o nomi) " a " e " b " significano la stessa cosa.

L'uguaglianza parlerebbe proprio di tali segni, affermerebbe un rapporto tra essi (e non fra gli oggetti).

Utilizzando questa breve osservazione sul ruolo dell'uguaglianza, possiamo descrivere un rapporto tra tempo e incertezza.

Tempo e incertezza hanno un rapporto tra loro molto più stretto di quello che possa sembrare a prima vista: il tempo misura l'incertezza e d'altra parte, l'incertezza misura il tempo.

L'incertezza, infatti, da un lato è osservabile solo nel suo svilupparsi temporale (un solo dato, una sola osservazione istantanea, non può misurare alcun grado di incertezza); dall'altro, si è appena visto come l'uguaglianza sia in grado di produrre conoscenza proprio perchè è un rapporto tra nomi di oggetti. Ora, di per sé, il sole di oggi è un oggetto diverso dal sole di ieri: basti pensare ad esempio alla variazione di massa dovuta alla dispersione di energia, eppure, l'aver dato lo stesso nome a due oggetti diversi è stata sicuramente una grossa scoperta.

3) Ad esempio è noto che il pianeta Plutone è stato "visto" prima tramite la deduzione e poi attraverso l'osservazione empirica.

In realtà un oggetto può essere uguale a se stesso solo nell'istante in cui il tempo, per astrazione, "non scorre". E' quindi questa continua verifica delle diversità che misura il "passare del tempo": l'incertezza è la misura del tempo.

Questa *struttura duale* legata al nostro modo di intendere ed usare l'uguaglianza è la lente attraverso cui osserviamo i fenomeni, ma non basta da sola a dare significato alle osservazioni stesse. Sono ben di più le cose viste, infatti, che quelle spiegate.

Il problema a questo punto diventa: se i dati generati dalla nostra struttura duale non si spiegano da soli, attraverso quali processi diventano "conoscenza"?

Un'uguale dualità la troviamo tra dati e teoria. Analogamente al tempo, le macchine sono delle "astrazioni mentali", delle invenzioni della nostra fantasia, dove i fenomeni naturali si ripetono uguali a se stessi fino a che la macchina non si rompe, imponendo alla nostra osservazione la sua origine casuale.

Il legame tra la struttura delle macchine logiche e il teorizzare è molto più stretto di quanto appaia a prima vista.

Le *regole di deduzione* oggi conosciute, il "cemento" dei nostri castelli teorici sono, infatti, meccanizzabili, nel senso che, se formalizziamo una dimostrazione, esiste una macchina che può decidere se la mia dimostrazione è sintatticamente corretta o no.

Non solo, ma esiste una macchina che genera, uno dopo l'altro, tutti i teoremi degli assiomi logici. Quindi non ci stupisce che le macchine siano, nella loro forma astratta, completamente prevedibili teoricamente; anzi, potremo addirittura dire che questa è la caratteristica che le definisce.

Viceversa, una massa di dati assolutamente prevedibile non porta nessun aumento della nostra conoscenza, e quindi, è la loro aperiodicità, non regolarità, non prevedibilità completa che modifica le nostre teorie.

Ma questa caratteristica dei dati (la loro non prevedibilità, e cioè il loro potere conoscitivo) si può rilevare solo se disponiamo di una teoria a cui li si può confrontare. Dunque: solo possedendo la ripetitività meccanica della macchina sintattico-teorica possiamo, ad esempio, dividere i dati in due gruppi: quelli che la teoria spiega, e quelli non spiegati, che rappresentano gli elementi non previsti dalle regolarità descritte dalla teoria.

Quindi se non avessimo avuto la teoria non si sarebbe potuto classificare questi dati come irregolari e perciò osservare l'incertezza.

Va notato inoltre che il teorizzare ha un carattere essenzialmente locale: una teoria, proprio se "buona", se da un lato diminuisce -attraverso la sua capacità di previsione- l'incertezza relativa ad un intorno dei dati del fenomeno, dall'altro, quando ne deduciamo tutte le previsioni possibili e le confrontiamo nel tempo con i nostri dati, ci si accorge che, nello stesso modo, la stessa teoria è anche una generatrice di incertezza⁴.

Potremo dire che le teorie sono misura dei dati e viceversa.

Come già abbiamo visto nel meccanismo della previsione, i dati sono una misura del grado di "bontà" delle teorie. Comincia a diventare evidente come questi rapporti in dualità tra tempo-incertezza, teoria-dati siano elementi centrali nel processo del conoscere.

Il lavoro innovativo è giocato su queste polarità, è un muoversi in queste simmetrie. Infatti, che la teoria da sola descriva "il reale" è una finzione, tanto quanto l'esistenza "reale" delle macchine; viceversa i dati non sono leggibili, per ora, senza la chiave di lettura della "macchina teorica".

Il lavoro innovativo riassume in sé questa dualità attraverso la socializzazione e l'integrazione delle forme del conoscere: da un lato, la teorizzazione si sposta dal descrivere la realtà a inventare le realtà possibili⁵, dall'altro, questo processo ha significato solo se esiste una velocità di verifica sufficientemente alta in modo da poter scegliere in tempo reale uno fra i mondi possibili. E viceversa: la possibilità di utilizzare grosse masse di dati cioè di incertezza, dipende dall'esistenza di molti mondi possibili.

Da questo punto di vista, man mano che il lavoro innovativo diventa cosciente gioco del possibile, l'incertezza non appare più "l'altro dal conoscere", l'ignoto, ma l'*energia stessa del conoscere*, della conoscenza come fatto puramente sociale.

Per chiudere il cerchio del ragionamento resta da dire come le strutture duali che abbiamo descritto spieghino la capacità di previsione del lavoro innovativo e come questa possa essere vista come *misura* del lavoro prodotto. Allo scopo cercheremo di descrivere più dettagliatamente come la costruzione di teorie sia un metodo di lettura, di spiegazione di massa di dati.

La possibilità di leggere sistemi con alto grado di incertezza dipende dal poter costruire una grossa correlazione fra i dati e gli enunciati del linguaggio, nel senso che l'osservazione di un dato, o di un'insieme di dati, mi individua univocamente un enunciato, o un'insieme di enunciati, come molto probabile e gli altri come poco probabili.

Abbiamo detto che i dati sono, per noi, osservazioni codificate in enunciati, la costruzione di questa rete di legami possiamo vederla come la descrizione di un insieme di enunciati del tipo: "se A allora B", o negazione di enunciati di questo tipo, che si reputano altamente probabili.

4) Ad esempio la precessione del perielio dell'orbita di Mercurio, che non rientra nel quadro della meccanica newtoniana, ha portato ad ipotizzare l'esistenza di un pianeta in un'orbita più interna. Costi è stato "visto" il pianeta Vulcano che non esiste.

5) E' divertente notare come questo oscillare tra il poter essere invenzione del possibile, e dover essere descrizione del reale è un piccolo dramma che ha percorso un poco tutta la storia delle matematiche. Quante generazioni di matematici hanno sbattuto la testa per dimostrare l'oggettività del quinto postulato di Euclide, andando incontro a grossi fallimenti, quando in realtà la soluzione era che, negandolo, ci si poteva inventare il secondo campo delle geometrie non euclidee, tra cui, non ultima figlia, la Relatività einsteiniana. E lo stesso teorema di Godel nasce dal fallimento del progetto hilbertiano di dimostrare l'oggettività meccanica del conoscere.

Ora, anche formalmente, *estendere* una teoria significa proprio allargare questo insieme di legami. Infatti, aggiungere nuovi assiomi significa proprio questo: descrivere nuovi legami fra enunciati che prima erano slegati. Figurativamente potremo dire che quest'operazione rende più fine il mio "filtro teorico".

Se vediamo l'osservazione di un fenomeno come un generatore di dati, è chiaro che la comprensione o la lettura di un dato di un fenomeno che contiene un alto grado di incertezza porta in media molta più informazione di quella di un dato di un fenomeno che contiene un basso grado di incertezza.

Se chiamiamo tempo stesso questo rapporto duale tra tempo lineare e incertezza, un punto in questo tempo non è semplicemente determinato dall'istante lineare (il tempo in cui accade), ma anche dall'evento casuale che l'ha generato. Cioè: nello stesso momento di tempo lineare ci sono tanti momenti diversi di *tempo stesso* quanti sono gli eventi casuali possibili. Ciò che è interessante di quest'estensione del concetto di tempo, è che ci permette di descrivere un rapporto di dipendenza non solo tra passato e futuro, ma anche tra futuro e passato. Cerchiamo di chiarire con un esempio. Ammettiamo che uno lanci una moneta ogni minuto e che, avendo una lavagna, vi scriva sopra un numero ottenuto dal precedente sommando o sottraendo 1 a seconda che sia uscito testa o croce.

Prima del primo lancio sulla lavagna poniamo ci fosse scritto 0 e mettiamo a zero l'orologio al momento del primo lancio.

A questa successione casuale, possiamo collegare un altro evento aleatorio così caratterizzato: scrivere il valore massimo del tempo dell'orologio, entro 6 minuti, in cui è comparso sulla lavagna il numero 2, se vi è comparso, altrimenti scrivere 0.

Anche se, dopo 4 lanci, sulla lavagna è comparso il valore 2, non posso conoscere il tempo del secondo evento casuale se non aspetto il succedere di altri due lanci.

Infatti, se il risultato dei lanci successivi ai primi quattro è testa-testa o croce-croce, il mio evento accade al tempo lineare "4 minuti", mentre se il risultato dei due lanci è croce-testa o testa-croce, il mio evento succede al tempo lineare "6 minuti".

Il tempo esteso in cui avviene il secondo fenomeno aleatorio dipende dal succedere di eventi casuali successivi, rispetto al tempo lineare, all'accadere dell'evento stesso. Possiamo far dipendere l'accadere di un evento A al tempo lineare t dall'accadere di uno o più eventi B al tempo lineare T con $T > t$.

Abbiamo così costruito un canale informativo che dal futuro arriva a noi, cioè: vi è una corrispondenza in probabilità che, se osserviamo l'uscita di un dato al presente, questo sia stato generato da un evento che, rispetto allo scorrere del tempo lineare, è futuro.

In questo quadro, il teorizzare è semplicemente la costruzione di una codifica in grado di leggere informazioni così osservate.

Quindi la *capacità di previsione* è una misura del grado di finezza del filtro teorico o, che è lo stesso, del *valore* del lavoro innovativo prodotto.

Questo è il risultato a cui volevamo arrivare: trovare un criterio di misura del lavoro innovativo, visto che - e lo vedremo fra poco - il tempo lineare, se può essere misura del lavoro computabile, non può esserlo per il lavoro innovativo.

3. Significato del paradosso

Con queste osservazioni sull'innovazione, torniamo ai limiti delle macchine.

Cosa vuol dire che le macchine sono in grado di dimostrare i propri limiti? La spiegazione sta nel mettere assieme due cose:

1) *l'uso che la logica corrente fa della negazione*: se ammetto come tautologia (enunciato sempre vero) "A o non -A. vuol dire implicitamente che, qualunque sia l'enunciato A, possiedo sufficiente informazione per decidere la verità di A o A verità di non-A. In fin dei conti, l'ipotesi implicita nell'assunzione della tautologia "A o non-A" (qualunque sia l'enunciato "A"), è che l'incertezza è nulla.

2) *l'uso di linguaggi sufficientemente forti per parlare di se stessi*. Infatti, per costruire le dimostrazioni di cui abbiamo parlato (limiti delle macchine e Teorema di Gödel), abbiamo sempre fatto ricorso a questo tipo di linguaggi. Ma cosa significa che il connubio fra queste due cose "produce paradossi"? Analizzando la dualità tra tempo e incertezza si è visto che, se elimino l'incertezza, il tempo non scorre più.

Se inoltre posso usare linguaggi sufficientemente forti per parlare di se stessi, allora posso parlare adesso di una cosa che sto dicendo adesso; posso così costruire una frase che "nello stesso tempo" dice una cosa e la sua negazione ed è quindi indecidibile. Dal nostro punto di vista è un'altra descrizione, in negativo, della funzione positiva dell'incertezza nel processo del conoscere.

4. Il tempo delle macchine e il tempo dell'innovazione.

Abbiamo visto come il processo innovativo sia sostanzialmente irriducibile al processo meccanizzabile. Quindi il processo innovativo non è descrivibile con un insieme, per quanto complesso, di operazioni che si ripetono uguali nel tempo. E, pertanto non è misurabile con il tempo delle macchine.

Eppure, quando misuriamo il valore del lavoro in "tempo di lavoro" facciamo proprio quest'operazione di riduzione: ridurre tutto il lavoro a lavoro computabile. Questo però crea delle contraddizioni, soprattutto oggi che si dispiega il processo di sostituzione del lavoro computabile fatto dall'uomo, con le macchine.

Per vedere meglio queste contraddizioni, cominciamo dall'osservare più da vicino il processo di sostituzione. Tutte le macchine sono riducibili a quella che abbiamo chiamato "macchina logica": le macchine, cioè, anche se molto potenti, sono riducibili a poche operazioni

fondamentali: le quattro operazioni della Macchina di Turing. Il processo di sostituzione è in pratica un processo di omogeneizzazione, di codificazione in un linguaggio omogeneo (sempre le operazioni della Macchina di Turing) del sistema produttivo.

Sappiamo che, con l'aumentare del grado di omogeneità di un sistema, diminuisce la prevedibilità del sistema stesso e aumenta anche l'incertezza contenuta nel sistema (cfr. Appendice 3). A questo processo contribuisce, e di questo processo fa uso, il lavoro innovativo, con le seguenti modalità:

- 1) il lavoro computabile fatto dalle macchine libera il tempo disponibile per il lavoro innovativo;
- 2) l'omogeneizzazione del sistema è il modo per aumentare la velocità di circolazione dell'incertezza che si è visto essere la base e l'energia del lavoro innovativo;
- 3) buona parte del lavoro innovativo oggi è ancora usata per progettare e realizzare il completamento del processo di sostituzione.

Alcune brevi note conclusive.

Anticipiamo alcune osservazioni che verranno sviluppate in un lavoro successivo, anche per non lasciare al lettore che ci ha seguiti fino a qui l'impressione di assoluta inutilità dell'apparato teorico sviluppato.

Lo svilupparsi del lavoro innovativo e del processo di sostituzione richiede un sistema tanto strutturalmente instabile quanto difficilmente prevedibile. Mentre un sistema basato sul lavoro computabile fatto dall'uomo è sostanzialmente stabile per lunghi periodi, a meno di variazioni critiche, prevedibile, quindi, controllabile esogeneamente al lavoro stesso.

Inoltre, i metodi di valutazione del valore prodotto dal lavoro, basandosi ancora sulla misura del lavoro in tempo lineare di lavoro, diventano sempre più inadeguati, e il peso politico del lavoro computabile, anche per il processo di sostituzione in atto, diventa sempre più irrilevante; da qui la crisi di rappresentatività delle forme politiche e del vecchio sindacato del lavoro computabile, e l'enorme difficoltà a rappresentare il lavoro innovativo che non coincide di per sé con il terziario più o meno avanzato. Nello stesso tempo però il processo di valorizzazione del lavoro innovativo si fa strada anche se in maniera contraddittoria e inconsapevole.

Abbiamo detto che la misura del lavoro innovativo prodotto è la sua capacità di previsione, ed in effetti la capacità di previsione sembra sia stata, anche se in maniera rudimentale, un sistema di spostamento, e quindi di valorizzazione di grosse masse di ricchezza; basta pensare all'estendersi del gioco finanziario e delle micro imprese.

Abbiamo osservato lo svilupparsi di un'economia dei valori attesi, come verifica della capacità del sistema di raccogliere la sua forza di previsione o, che è lo stesso, di verificare/valorizzare il lavoro innovativo prodotto. Il problema è che siccome questa capacità di previsione è

limitata, sia dal punto di vista della sua qualità che nella sua estensione, essendo ristretta ad un piccolo numero di operatori con grossa potenza economica, in realtà si è cercato di forzare le cose a senso unico, cioè nel senso della crescita senza fluttuazioni.

Ora un sistema instabile, e con processi dipendenti dalla circolazione dell'informazione, come quello che attualmente si configura, può costruire un'immagine di sé spostata dal feed-back in maniera proporzionale alla concentrazione/controllo della circolazione dell'informazione. (Si sconta così l'inadeguatezza del canale rispetto all'incertezza esistente). Ed è precisamente quello che è successo, tonfi periodici compresi.

Appendice 1 Macchina di Turing.

Come si è già detto, una Macchina di Turing è fatta di un *nastro* e di un *congegno* in grado di compiere 4 operazioni che indicheremo con i seguenti simboli:

1 per l'operazione di *stampare* se il nastro è vuoto, oppure di *lasciare* 1 se sul nastro c'è già 1;

0 per *cancellare* quello che c'è sul nastro o lasciarlo invariato se c'è già zero;

D per *passare a destra* di una casella (celletta);

S per *passare a sinistra* di una casella.

L'esecuzione di una delle operazioni elementari sarà chiamata *gradino* nel lavoro della macchina.

Alla conclusione di ogni gradino, la macchina, distinta dal nastro, memorizza la propria configurazione interna.

Queste configurazioni sono gli *stati interni* e saranno indicati con i simboli (i), dove $i=0,1,2,3...$

1 e 0 denoteranno anche le condizioni possibili cui può trovarsi il nastro.

Un'istruzione della macchina è dunque formata (descritta) da quattro simboli nella seguente maniera:

il *primo* indica uno stato interno;

il *secondo* una possibile condizione della cella del nastro osservata dalla macchina;

il *terzo* un'operazione;

il *quarto*, infine, il nuovo stato interno in cui la macchina si trova dopo l'esecuzione dell'operazione indicata.

Ad esempio, è un'istruzione:

(0)10(1)

Un qualsiasi insieme finito di queste quadruple è una Macchina di Turing purché soddisfi alla seguente condizione: due quadruple distinte devono differire per il primo e il secondo simbolo. Questo significa richiedere che non esistano due istruzioni che impongano alla macchina due operazioni diverse nello stesso tempo.

Tuttavia non richiediamo che nel programma della macchina per ogni stato (i) ci siano entrambe le combinazioni possibili: (i)1 e (i)0, quindi in certe occasioni può succedere che la macchina non compia nessuna operazione, in questo caso diciamo che la macchina si ferma.

Esempio. La macchina determinata dalle istruzioni: (0)10(2), (2)0D(0) cancella tutti gli 1 consecutivi al primo che trova sul nastro e poi si ferma. Infatti, se ad esempio sul nastro ci sono

cinque 1 consecutivi e tutti 0, e poniamo la macchina sul primo 1 a sinistra nello stato (0) e facciamo partire la macchina, succede questo: la macchina esegue la prima istruzione (0)10(2) che dice "se nello stato (0) trovi 1, cancellalo e passa allo stato (2)". Quindi, dopo aver eseguito quest'istruzione la macchina "legge" l'istruzione (2)0D(0) che dice "se nello stato (2) trovi 0 passa a destra e torna allo stato (0)". Cancellando 1 con la prima istruzione, la macchina, nello stato (2) legge ora 0 e dunque "passa" a destra dove trova 1; poichè l'istruzione la rimanda allo stato (0) essa eseguirà ancora l'istruzione (0)10(2), cioè la prima istruzione, eseguendone l'operazione: cancellare 1. E così avanti fino ad aver cancellato tutti gli uno consecutivi; a quel punto, nello stato (0), la macchina non troverà più 1 ma 0 e quindi si fermerà.

Appendice 2

Limiti delle macchine.

Si è detto che esiste una macchina che genera, una dopo l'altra, tutte le macchine possibili. Prendiamo questa enumerazione come standard: si può quindi parlare della prima macchina, della seconda, ..., della n-esima macchina, ... ecc.

E' facile costruire una macchina che non si ferma mai: le istruzioni che la caratterizzano: (0)11(0), (0)00(0); la macchina comunque "gira su se stessa".

Proviamo ad alleggerire il linguaggio.

Indichiamo con A_n la macchina che sta all'n-esimo posto nella lista; se a questa macchina diamo y come input e, dopo un numero finito di passi, si ferma, diremo che $A_n(y)$ converge e indicheremo con $A_n(y)$ anche il suo output. Altrimenti diremo che la macchina diverge.

Con quanto finora esposto, siamo in grado di costruire un problema che nessuna macchina è in grado di risolvere.

Il problema è: se esiste una macchina in grado di decidere se, presa una qualsiasi macchina e un qualsiasi input, questa macchina con quell'input termini il calcolo in un numero finito di passi (cioè converge) o no (cioè diverge). Espresso formalmente il discorso è:

"Esiste una macchina che, data una qualsiasi coppia di numeri interi (x, y) , prende la x-esima macchina A_x , le fa calcolare $A_x(y)$; se $A_x(y)$ si ferma dopo un numero finito di passi, la nostra macchina scrive 1; altrimenti scrive 0".

Ammettiamo che tale macchina esista e chiamiamola $E(x, y)$. Dunque: $E(x, y)$ scrive: 0 se $A_x(y)$ diverge; e scrive 1 altrimenti (cioè, converge).

Per calcolare $E(x, y)$, quindi, si procede in questo modo: si prende la x-esima macchina della nostra lista, A_x , si fa partire la macchina A_x con y come input; se questa si ferma dopo un numero finito di passi $E(x, y)=1$, altrimenti $E(x, y)=0$.

Se E è una macchina, possiamo definirla un'altra così: $B(x)$ scrive: 1 se $E(x, x)=0$ e diverge se $E(x, x)=1$ (nota: la scrittura $E(x, x)$ può essere interpretata così: "la macchina parla di se stessa").

Se E è una macchina anche B lo è; infatti, per calcolare $B(x)$ basta vedere cosa fa $E(x, x)$.

Se è 0, basta far stampare un 1 a B ; se invece $E(x, x)=1$, basta mandare B nel gruppo di istruzioni della macchina che diverge sempre (abbiamo già visto quale tipo di istruzione fa girare la macchina all'infinito).

Allora: se B è una macchina starà sulla nostra lista.

Sia quindi y_0 il suo posto nella lista. Vogliamo *vere* cosa fa B quando le assegniamo y_0 come input.

$B(y_0)$ dà output se e solo se $E(y_0, y_0)=0$, così infatti lo abbiamo definito: ma $E(y_0, y_0)=0$ se e solo se $A_{y_0}(y_0)$ non dà output, ma $A_{y_0}(y_0)=B(y_0)$. Ad analoga contraddizione si giunge se si assume che $B(y_0)$ diverge.

Aver assunto che E sia una macchina ci ha perciò portato a una contraddizione e, dunque, a "rigor di logica", possiamo dire che E non è una macchina, o più correttamente, che le operazioni ora descritte non sono meccanizzabili.

Naturalmente saremmo pervenuti alla contraddizione anche partendo dall'ipotesi che $B(y_0)$ diverge.

Una dimostrazione così semplice dell'esistenza di problemi irrisolvibili meccanicamente è uno dei grossi risultati della matematica degli ultimi 50 anni: semplice e anche molto elegante.

N.B.: I limiti delle macchine così individuati non hanno niente a che vedere con limiti di natura ingegneristica (ad esempio: potenza della memoria e velocità di calcolo). Infatti nel caso qui considerato si è addirittura supposto che la macchina abbia memoria infinita (infinite cellette sul nastro) e non si è posto alcun limite temporale al calcolo. Si tratta di limiti intrinseci alle nostre attuali forme conoscitive in generale.

Appendice 3

Lo schema utilizzato in Teoria dell'informazione per descrivere la situazione a cui si applicano i metodi della teoria è questo: **SORGENTE CODIFICATORE CANALE DECODIFICATORE DESTINATARIO**

Un esempio può essere il telegrafo:

La "sorgente" è un messaggio che si vuole trasmettere, il "codificatore" è lo strumento che trasforma il messaggio in punti e linee - e questi in impulsi elettrici, il "canale" è il filo, il "decodificatore" fa l'operazione inversa al codificatore, cioè trasforma gli impulsi in punti e linee - e questi in lettere. Se prendiamo una sorgente che emette segnali scelti in un alfabeto finito $a_1, a_2, a_3, \dots, a_n$, con probabilità $p(a_1), p(a_2), \dots, p(a_n)$, si definisce la quantità di informazione contenuta in una lettera a_i dell'alfabeto come:

$$-\log p(a_i)$$

Il senso della definizione è questo: $-\log p(a)$ è una funzione che cresce man mano che $p(a)$ ("probabilità di a ") diventa piccolo; il che vuol dire che l'uscita di una lettera molto frequente (cioè: $p(a)$ grande) porta poca informazione, mentre l'uscita di una lettera con probabilità più bassa ci dà più informazione. La quantità di informazione media della sorgente è l'entropia della sorgente.

Chiariamo ora, con un esempio, la relazione tra entropia e omogeneità del sistema.

Un tale vuole comunicare ad un altro il risultato del lancio di una moneta. Di quanta informazione ha bisogno? Se la moneta è equilibrata (sistema omogeneo, cioè probabilità ugualmente ripartita tra testa e croce), la quantità di informazione è massima. Quanto più la moneta è squilibrata in favore di una faccia (sistema disomogeneo), tanto meno informazione è necessaria.

L'idea è che una struttura più disomogenea è più prevedibile e quindi abbisogna di meno informazione per essere descritta.