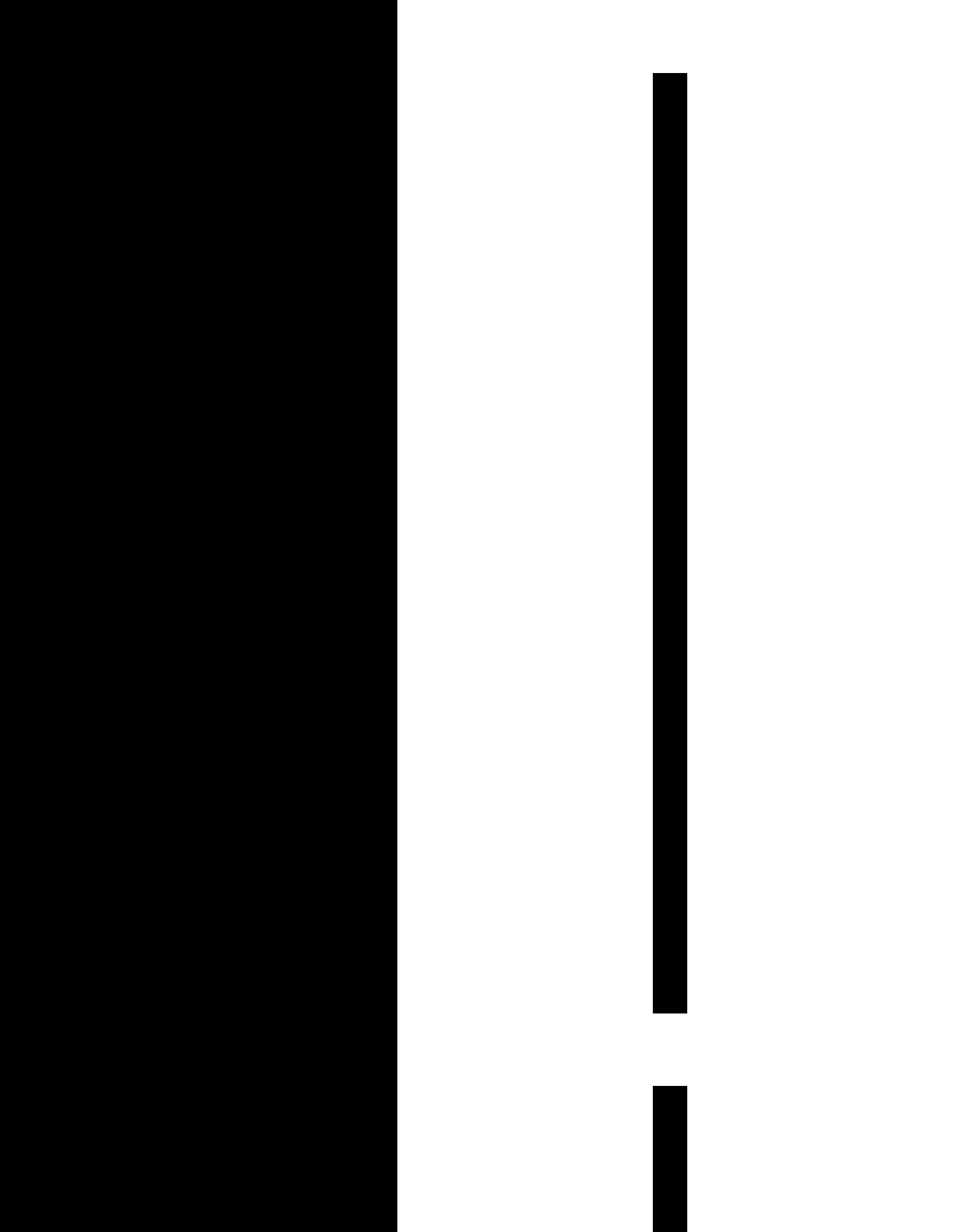


[dati]

anno 14 numero **27** / 10.19

**asimmetrie**

rivista semestrale dell'Istituto  
Nazionale di Fisica Nucleare



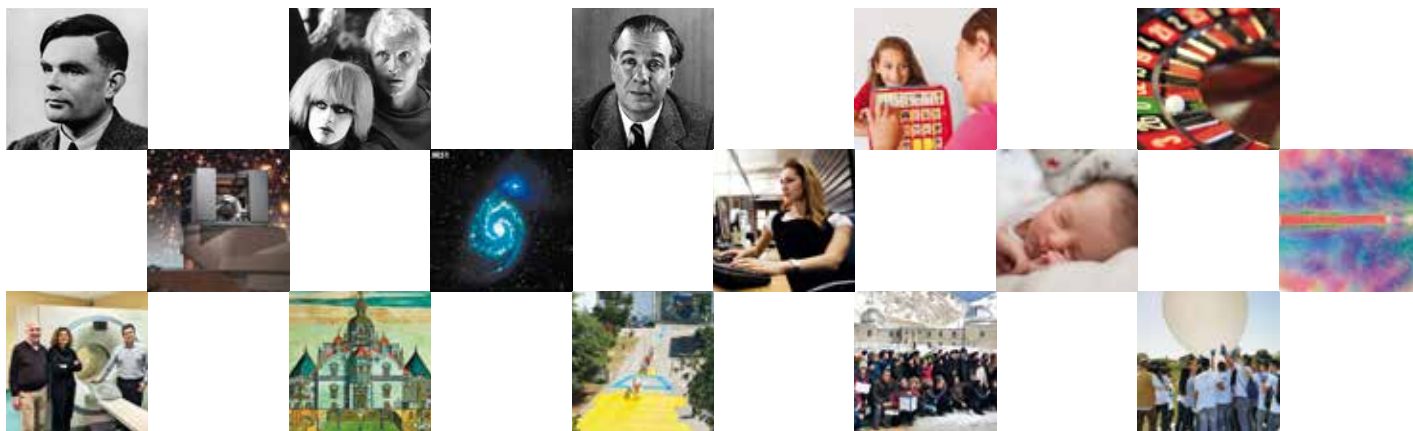
# asimmetrie

Cari lettori di Asimmetrie,

In questi anni la fisica delle particelle si è trovata ad analizzare i dati prodotti al Large Hadron Collider (Lhc) del Cern di Ginevra. Si tratta di campioni di decine di petabyte per anno, ovvero milioni di miliardi di byte di dati. Per affrontare questa sfida, i fisici hanno creato un'infrastruttura di calcolo distribuita sparsa su tutto il mondo denominata Grid, che di fatto è stata un precursore delle Cloud che adesso sono di uso corrente nella nostra società. Ora la nostra comunità è chiamata a ripensare all'organizzazione del calcolo scientifico: infatti, nei prossimi anni il mondo scientifico e quello industriale produrranno una quantità di dati senza precedenti. Nuove infrastrutture di ricerca entreranno in funzione in diversi campi, che vanno dalla fisica delle particelle, all'astronomia e all'astrofisica, come Hilumi-Lhc al Cern o il gigantesco radiotelescopio Ska in costruzione tra il Sud Africa e l'Australia, ma anche in campi come la geofisica, la biologia e le scienze mediche, umane e sociali, e avranno come caratteristica comune la capacità di produrre una quantità di dati di tanti ordini di grandezza maggiore di quella attualmente disponibile nel mondo intero. Lo stesso discorso vale per il mondo industriale, dove praticamente tutti i settori stanno subendo una "trasformazione digitale", volta all'innovazione e al sostanziale miglioramento della competitività e della sostenibilità. La sfida dei prossimi anni, sia per il mondo scientifico, che per quello industriale sarà da una parte quella di realizzare infrastrutture di calcolo, che siano in grado di trasferire, immagazzinare e processare i *big data* che verranno prodotti, anche attraverso lo sviluppo di nuove tecnologie, dall'altra quella di sviluppare nuove tecniche e nuovi algoritmi che permettano di estrarre con efficienza le informazioni rilevanti immagazzinate in questo mare di dati. Infine, la società dovrà essere in grado di formare una nuova classe di giovani che abbiano le capacità e le competenze per affrontare e vincere queste sfide. Questo numero di Asimmetrie affronterà queste affascinanti tematiche, cercando di delineare il quadro attuale e i possibili sviluppi che ci attendono nei prossimi anni.

Buona lettura.

**Antonio Zoccoli**  
*presidente Infn*



#### asimmetrie

Rivista dell'Istituto Nazionale di Fisica Nucleare

Semestrale, anno 14,  
numero 27, ottobre 2019

#### direttore editoriale

Antonio Zoccoli, *presidente Infn*

#### direttore responsabile

Catia Peduto

#### direttore comitato scientifico

Egidio Longo

#### comitato scientifico

Vincenzo Barone  
Massimo Pietroni  
Giorgio Riccobene  
Barbara Sciascia

#### redazione

Francesca Mazzotta  
Francesca Cuicchio  
(infografiche)

#### hanno collaborato

Viviana Acquaviva, Cristiano Bozza,  
Michele Caselle, Elettra Colonna,  
Eleonora Cossi, Stefano Giagu,  
Vittorio Loreto, Gaetano Maron,  
Chiara Neglia, Luciano Pandola,  
Pier Stanislao Paolucci,  
Giada Pedringa, Mario Rasetti,  
Paolo Rossi, Francesca Scianitti

#### contatti redazione

Infn Ufficio Comunicazione  
piazza dei Caprettari 70  
I-00186 Roma  
T +39 06 6868162  
F +39 06 68307944  
comunicazione@presid.infn.it  
www.infn.it

#### impaginazione

Tipolitografia Quatrini

#### stampa

Tipolitografia Quatrini

su carta di pura cellulosa  
ecologica ECF  
Fedrigoni Symbol™ Tatami  
250 - 135 g/m2

Registrazione del Tribunale di Roma  
numero 435/2005 del 8 novembre 2005.  
Rivista pubblicata da Infn.

Tutti i diritti sono riservati. Nessuna parte della  
rivista può essere riprodotta, rielaborata o  
diffusa senza autorizzazione scritta dell'Infn,  
proprietario della pubblicazione.

Finita di stampare nel mese di settembre 2019.  
Tiratura 20.000 copie.

#### come abbonarsi

L'abbonamento è gratuito.

Per abbonarsi compilare l'apposito  
form sul sito [www.asimmetrie.it](http://www.asimmetrie.it)

In caso di problemi contattare  
la redazione all'indirizzo  
[comunicazione@presid.infn.it](mailto:comunicazione@presid.infn.it)

#### sito internet

**Asimmetrie 27 e tutti i numeri  
precedenti della rivista sono  
disponibili anche online su  
[www.asimmetrie.it](http://www.asimmetrie.it)**

#### e-magazine

**Asimmetrie è anche disponibile  
in versione digitale, ricca di ulteriori  
contenuti multimediali, come app  
di iOS e Android sull'Apple Store  
e nel Google Play Store.**

#### crediti iconografici

Foto copertina © Istockphoto.com (fpm) //  
foto p. 4 © Istockphoto.com (gremlin); foto p.5  
©Universal Art Archive/Alamy Stock Photo;  
fig. c p. 7 ©Asimmetrie-Infn; foto d p. 8 ©A.  
Cappello-INFN; foto e p. 9 ©Archive Photos/  
Stringer // foto a p. 10 ©Grete Stern; foto  
b p. 11 ©Ingbert Gruttner, courtesy ALP  
Emilio Segrè Visual Archives, Physics Today  
Collection; foto c p. 12 © Istockphoto.com  
(Michele Jackson); fig. d p. 13 ©Asimmetrie-  
Infn; foto d p. 14 ©World Economic Forum  
/ Sikarin Thanachaiary // foto b p. 16 ©F.  
Cuicchio/Asimmetrie-Infn; foto c p. 17  
@thispersondoesnotexist.com // fig. pp. 18-19  
©Asimmetrie-Infn // foto a p. 20 © PDO/TDP/  
DRAO/Swinburne Astronomy Productions; fig.  
b p. 21 ©Asimmetrie-Infn; foto c p. 22 ©CNAF-  
INFN; fig. 1 p. 23 ©Asimmetrie-Infn // foto  
a p. 24 ©F. Cuicchio-Istockphoto(ivanastar,  
LeventKonuk); foto b p. 25 ©CERN; foto c p.  
26 © CERN; foto d p. 27 ©CPPM // foto b  
p. 29 ©Ralf Roletschek; fig. c p. 30 ©G.A.P.  
Cirrone, INFN-LNS // foto a p. 31 ©LSST;  
fig. b p. 32 ©NASA/JPL-Caltech/UCLA;  
foto c p. 33 // fig. a p. 34 ©Asimmetrie-  
Infn; fig. b p. 35 © Allen JD, Xie Y. ChenM,  
Girard L., Xiao G – PLOS ONE; fig. c p. 36 ©  
Asimmetrie-Infn // foto a p. 37 © Istockphoto.  
com (Daniela Jovanovska-Hristovska); fig. b  
p. 38 ©Andreashorn // foto a p. 40 © Photo  
Researchers Science History Images/Alamy  
Foto Stock // foto a p. 42 ©Point8; foto b  
p. 43 ©CERN // foto a p. 46 ©OCRA; foto  
b p. 47 ©V.Bocci // foto p. 48 ©C. Peduto/  
Asimmetrie-Infn.

Ci scusiamo se, per cause del tutto  
indipendenti dalla nostra volontà, avessimo  
omesso o citato erroneamente alcune fonti.

as

# 27 / 10.19

## [dati]

|                                                                      |    |                                                                                           |    |
|----------------------------------------------------------------------|----|-------------------------------------------------------------------------------------------|----|
| <b>Il dato è tratto</b><br>di Mario Rasetti                          | 4  | <b>Bio-intelligenza artificiale</b><br>di Pier Stanislaw Paolucci                         | 37 |
| <b>Creatività aumentata</b><br>di Vittorio Loreto                    | 10 | <b>[as] riflessi</b><br><i>Data mining</i> in fisica medica.<br>di Francesca Mazzotta     | 39 |
| <b>Caduti nella rete (neurale)</b><br>di Stefano Giagu               | 15 | <b>[as] radici</b><br>I dati di Tycho.<br>di Paolo Rossi                                  | 40 |
| <b>[as]</b><br>Macchine intelligenti.                                | 18 | <b>[as] traiettorie</b><br>Point 8.<br>di Catia Peduto                                    | 42 |
| <b>Questioni di core</b><br>di Gaetano Maron                         | 20 | <b>[as] selfie</b><br><i>Blazar</i> , che passione!<br>di Elettra Colonna e Chiara Neglia | 44 |
| <b>Scavare nei dati</b><br>di Cristiano Bozza                        | 24 | <b>[as] spazi</b><br>Il colore dell'invisibile.<br>di Eleonora Cossi                      | 46 |
| <b>Il casinò della fisica</b><br>di Luciano Pandola e Giada Petringa | 28 | <b>[as] illuminazioni</b><br>Immergersi nell'universo.<br>di Francesca Scianitti          | 48 |
| <b>Nello zoo galattico</b><br>di Viviana Acquaviva                   | 31 |                                                                                           |    |
| <b>Le reti della vita</b><br>di Michele Caselle                      | 34 |                                                                                           |    |

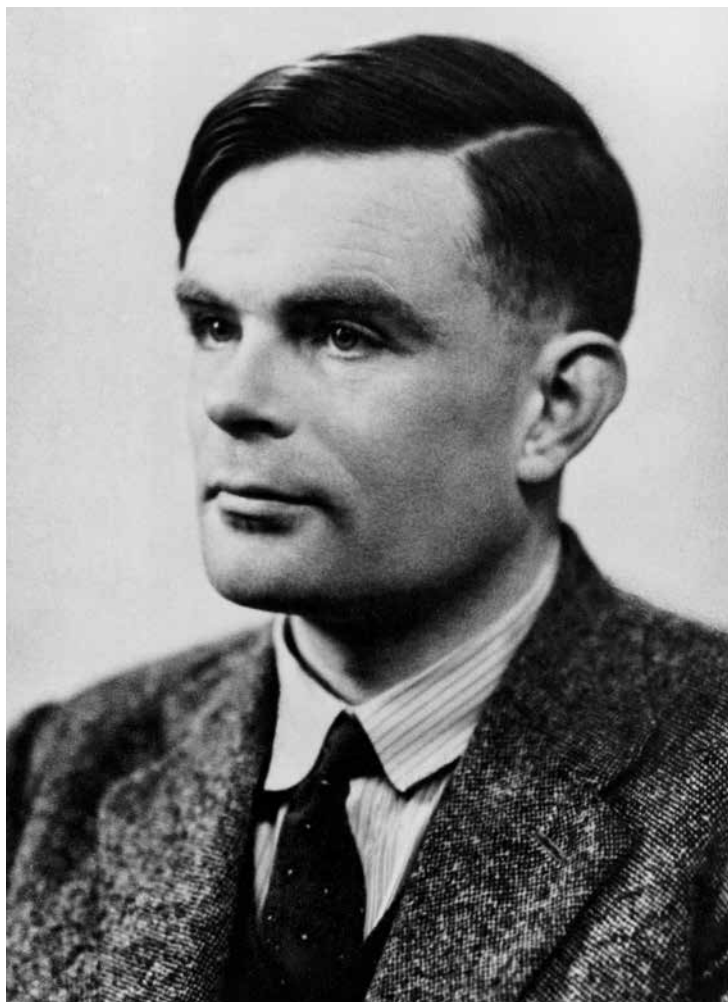


# Il dato è tratto

*Big data e intelligenza artificiale tra scienza e società*

di Mario Rasetti





a.  
Alan Turing, genio ideatore della “macchina di Turing”, il più generale schema concettuale oggi conosciuto capace di affrontare in modo universale l'esecuzione di algoritmi che calcolano funzioni “ricorsive”.

È in corso una rivoluzione, quella dei “*big data*” e dell'intelligenza artificiale, travolgente e unica: una rivoluzione culturale e industriale al tempo stesso, che condivide speranze e rischi di tutti i cambiamenti drastici che avvengono nei due ambiti. I bit faranno ben di più, in termini di spostamento degli equilibri del potere, di accesso alla conoscenza e del suo trasferimento dalle mani di pochi a comunità sempre più allargate, di quanto abbiano fatto i caratteri mobili di Gutenberg nel '400. E si produrrà un cambio di paradigma che attraverso il lavoro muterà la struttura profonda delle relazioni umane e sociali, come fece nel '700 la macchina a vapore di Watt.

È una rivoluzione che coinvolge a fondo la scienza e gli scienziati. Le scienze - specie quelle della vita e della società - s'interrogano sul fatto che per la prima volta i dati su cui operano non vengono dai loro esperimenti o dalle loro misure, ma sono proprietà di grandi società. Vacillano pilastri etici e metodologici ritenuti fuori discussione, come la ripetibilità quale criterio per la validazione dei dati. Ed è epocale che nelle neuroscienze siano emersi problemi interpretativi indecidibili secondo Turing (ma che il cervello risolve)!

Una rivoluzione che induce modi di vita sempre più

condizionati da grandi masse di dati. Quanto grandi? Numeri inimmaginabili: nel 2018 si sono generati dati per centinaia di zettabyte (un numero di Avogadro di byte!), l'equivalente di ottanta milioni di volte il contenuto della Library of Congress di Washington, in appena un anno! Il tempo di raddoppio - quello in cui si produce una quantità di dati uguale a quella generata nella storia fino a quel momento - passerà, anche a causa dei 150 miliardi di dispositivi del cosiddetto “internet delle cose” (vd. fig. c), da circa un anno di oggi a 12 ore nel 2025.

Una rivoluzione, infine, che muterà i nostri valori fondamentali: lavoro, democrazia, relazioni umane, modo di fare scienza, affari, cultura, modo di perseguire salute e felicità, prospettando un futuro di progresso, crescita e qualità della vita senza precedenti, ma con i vincoli etici più severi che mai abbiamo dovuto affrontare.

L'intelligenza artificiale, che mira a decifrare il codice dell'intelligenza umana, per ora nelle sue capacità di imparare e auto-addestrarsi (il cosiddetto “*machine learning*”), ne è lo strumento e sta facendo progressi mozzafiato. Esistono già algoritmi in grado di svolgere compiti intelligenti che hanno prestazioni migliori di quelle degli uomini, e presto le capacità umane saranno superate in così tante applicazioni,



che si stima che entro un decennio il lavoro intellettuale sarà sostituito per il 60% dalla tecnologia e che metà dei lavori scompariranno.

Occorreranno enormi investimenti in cultura e formazione, incentivi alla creatività e nuove forme educative per ricreare un patrimonio di esperienza professionale e riqualificare gli espulsi dal sistema produttivo, un diverso modello sociale basato sulla redistribuzione non solo della ricchezza ma del lavoro, nuove fonti di conoscenza per generare altro sapere e creare così lavoro con nuovi lavori.

Per fortuna abbiamo lo strumento adatto: quella macchina senza pari che è il nostro cervello. Ma siamo forse troppo lenti nel costruire gli ingredienti mancanti: una visione condivisa del futuro e la capacità di convivere con la tecnologia in modo proattivo.

La scienza è cruciale in questo scenario: essa già lavora a una *roadmap* che indaga le radici della conoscenza, mirando a identificare i principi, forse le leggi, dell'intelligenza. Come Newton scoprì principi e leggi della

dinamica, quando sembrava irragionevole che ci fosse un legame tra fenomeni quali il moto dei pianeti e la turbolenza di un fluido, così oggi inizia una ricerca, altrettanto audace, dei principi generali dell'intelligenza, irragionevolmente condivisi da umani e robot. Le tecnologie digitali ci fanno credere di essere meno soli, perché sempre e sempre più connessi. Questa immensa condivisione non è che l'illusione che il nostro profilo online abbia un valore intrinseco misurato dal numero dei contatti, ma quanto più questi sono numerosi tanto più si acuisce la solitudine e si diluisce la profondità dei rapporti.

Inoltre, siamo sempre più vulnerabili alla manipolazione di opinioni, sentimenti, credenze dalla disinformazione digitale; strumento subdolo, nascosto in quella congerie di vincoli e pregiudizi sociali, cognitivi, economici cui siamo sottoposti dalla rete.

La sfida, allora, è far diventare quella straordinaria connettività, che la tecnologia

b.  
La Library of Congress di Washington, la più grande biblioteca al mondo: contiene 39,5 milioni di libri, 3,8 milioni di registrazioni, 16 milioni di fotografie, 5,6 milioni di mappe, 7,4 milioni di spartiti e 76 milioni di manoscritti.





di un sistema logico rigoroso: la sua funzione primaria è di difendere la persona dai pericoli dell'ambiente, con risposte che devono essere automatiche e ben più veloci del ragionamento razionale. Le funzioni della corteccia cerebrale, generatore della razionalità, sono in equilibrio dialettico competitivo con il ruolo primario del cervello. E, grazie a un cervello, che è la macchina più sofisticata che mai vedremo, l'intelligenza dipende da capacità ben più complesse, sottili e avanzate delle parole per comunicare e preservare le informazioni che genera: la sua vera forza è di saperle codificare nella materia.

Ciò che distingue l'intelligenza naturale da quella artificiale non è ciò che essa è, ma solo com'è fatta: l'intelligenza artificiale è l'artefatto di una particolare cultura umana e ne riflette i valori. La società ha una grande capacità collettiva di elaborare informazioni, perché la nostra comunicazione non è solo verbale. Coinvolge ben più delle

parole e implica la creazione di oggetti che trasmettono non qualcosa di fragile come un'idea, ma qualcosa di concreto come il *know-how* e l'uso della conoscenza. Gli oggetti ci "aumentano". Ci permettono di fare cose senza sapere come: la funzione "pensare" si è sviluppata conquistando il dominio prima sulla materia, poi sull'energia e infine sull'ordine fisico, cioè sull'informazione. Nel processo di evoluzione, il prossimo passo della capacità dell'*homo sapiens-sapiens* (l'uomo moderno) di generare informazione saranno proprio le reti uomo-macchina. Insieme, gli umani e le loro estensioni tecnologiche (quasi-) pensanti, continueranno a evolvere in reti sempre più complesse, asservite allo scopo di creare nell'universo nuove nicchie dove l'informazione non si riduce ma cresce, in una lotta senza fine contro l'entropia. Ma attraverso l'intelligenza artificiale gli umani diverranno più efficienti e forse più abili, non più intelligenti.

Con il continuo espandersi delle risorse

dell'informazione e della comunicazione (l'Ict, acronimo di *information and communications technology*) e il progressivo aumento della connettività, le intelligenze artificiali sapranno riprodursi, competere, evolvere. La complessità di tali processi di evoluzione rende impossibile prevedere se e quali nuove forme di intelligenza artificiale emergeranno (l'intelligenza artificiale oggi non ha ancora superato la barriera del significato, e sappiamo già che a livello semantico essa non è in grado di eguagliare le potenzialità di un cervello umano): possiamo dire solo che per la scienza ci sarà tanto da fare! Tutte le specie si estinguono. L'*homo sapiens-sapiens* non farà certamente eccezione. Non sappiamo come accadrà:

d.  
Foresta di neuroni della mostra dell'Infn "Uomo virtuale", allestita dal 4 maggio al 13 ottobre al Mastio della Cittadella di Torino.





un virus, un'invasione aliena, un asteroide, il Sole che diventa gigante rossa oppure, semplicemente, la mutazione a una super-specie evoluta, l'*homo sapiens-sapiens-sapiens* o una nuova fantascientifica forma di intelligenza artificiale pensante. Non credo a quest'ultima ipotesi. Credo invece che l'intelligenza artificiale sarà certo fonte di tante sorprese (e profitto!) per anni a venire, ma che gli umani non potranno mai essere sostituiti.

Un giorno, lontanissimo, alle macchine proveremo a chiedere di trascendere il semplice ruolo di "*problem solver*" e di diventare innovative e creative. Ma non c'è alcun segno che indichi che ci possiamo

trovare in difficoltà nel tenere l'intelligenza artificiale a freno, nel caso in cui si comporti male: siamo molto distanti dall'avere robot aggressivi, imprevedibili, machiavellici, con ambizioni di potere e la vocazione a riprodursi! Chiedersi se ipotetiche macchine pensanti possano costituire una minaccia è prudente, come per qualsiasi tecnologia. Ricordiamo, però, quante volte abbiamo creduto di essere sull'orlo di un abisso e di avere un modo, miracoloso ma controverso, per salvarci: la soluzione è sempre stata un ragionevole compromesso, non una semplicistica dialettica "sì-no".

A ben pensarci, è dell'uomo che dobbiamo preoccuparci, non delle macchine!

e.

Una scena del film Blade Runner (1982). Pris (Daryl Hannah) e Batty (Rutger Hauer, recentemente scomparso) sono due replicanti modello avanzato Nexus 6: un tentativo di *homo-sapiens-sapiens-sapiens*?

#### Biografia

**Mario Rasetti** è professore emerito di fisica teorica al Politecnico di Torino, di cui ha fondato e diretto per molti anni la Scuola di Dottorato. È presidente della Fondazione Isi, consigliere della Commissione Europea, ha vinto il premio Majorana 2011 per la fisica dei campi e la medaglia Volta. Si occupa di meccanica statistica, informazione e computazione quantistica e *big data*.

DOI: 10.23801/asimmetrie.2019.27.1

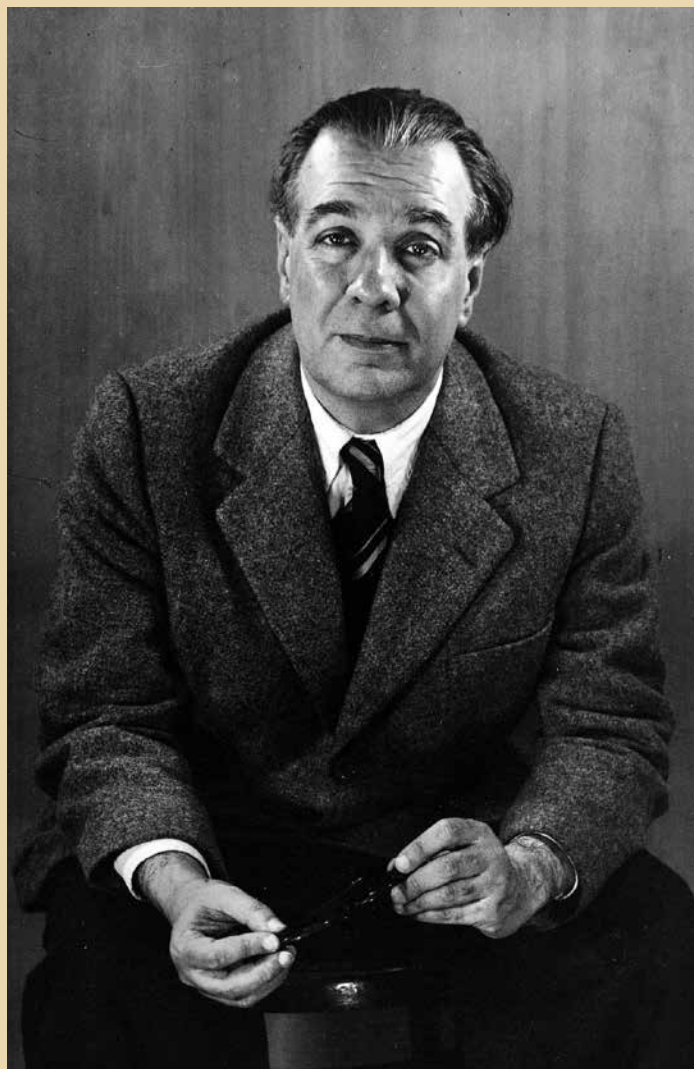


# Creatività aumentata

## Modelli teorici per l'interpretazione dei dati

di Vittorio Loreto

a.  
Una foto del 1951 del famoso  
scrittore argentino Jorge Luis  
Borges.



*“Noi, in un’occhiata, percepiamo: tre bicchieri su una tavola. Funes: tutti i tralci, i grappoli e gli acini d’una pergola. Sapeva le forme delle nubi australi dell’alba del 30 aprile 1882, e poteva confrontarle, nel ricordo, con la copertina marmorizzata di un libro che aveva visto una sola volta, o con le spume che sollevò un remo, nel Rio Negro, la vigilia della battaglia di Quebracho”. (“Funes, o della memoria” di Jorge Luis Borges, in “Finzioni”)*

Iniziamo con questa citazione di Borges per far comprendere, con un esempio forse paradossale, come la disponibilità di enormi moli di dati non sempre rappresenti un vantaggio. Ireneo Funes, condannato a una prodigiosa capacità di registrare e ricordare anche il più piccolo dettaglio del suo universo, sarebbe stato incapace di riconoscere un suo familiare la mattina dopo perché nella notte, mentre dormiva, la disposizione dei capelli era cambiata. Quando la mole di dati da “processare”, di cui, per fortuna o purtroppo, tutti noi disponiamo, è smisurata, possiamo essere incapaci di creare categorie o cogliere regolarità. Questa incapacità, associata alla mancanza di strumenti per dare un senso alla vastità di dati, può seriamente menomare la nostra possibilità di comprendere e spiegare i fenomeni che ci circondano. E in ultima analisi la nostra capacità di anticipare il futuro e fare previsioni.

Non era di questo avviso, Chris Anderson, che nel 2008, allora direttore della rivista Wired, scrisse un pezzo intitolato provocatoriamente *“The end of theory”* (vd. link sul web). Anderson sosteneva, al contrario, che l’enorme abbondanza di dati, unita alle straordinarie capacità computazionali di cui disponiamo, potesse dar luogo ad algoritmi di *machine learning* capaci addirittura di surclassare gli esseri umani e la loro capacità di costruire teorie. L’affermazione chiave (argomentata principalmente attraverso le scoperte prodotte pochi anni prima dal sequenziamento del genoma umano, vd. p. 34, ndr) era: “Le correlazioni sostituiscono la causalità, e la scienza può avanzare persino senza modelli coerenti, teorie



b.  
Henri Poincaré (1854 - 1912)  
e Mark Kac (1914 - 1984,  
nella foto) provarono due  
risultati importantissimi per la  
predicibilità dei sistemi dinamici  
deterministici. Poincaré dimostrò  
il celebre teorema di ricorrenza,  
che stabilisce che un sistema  
dinamico deterministico ritornerà  
dopo un tempo opportuno vicino  
allo stato iniziale. Esistono dunque  
in teoria degli stati analoghi di  
cui si potrebbe trovare traccia nel  
passato. La domanda è quanto  
tempo impiegherà il sistema a  
tornare vicino allo stato iniziale. Il  
lemma di Kac mostrò come tale  
tempo di ricorrenza può essere  
arbitrariamente grande per sistemi  
anche di bassa complessità.

unificate e addirittura qualsiasi tipo di spiegazione meccanicistica”.

Ma è veramente così?

Il criterio fondamentale per comprendere la bontà di una teoria o di una costruzione matematica dovrebbe essere la capacità di effettuare previsioni e verificarne l’attendibilità in opportuni contesti. La predizione del futuro rappresenta da sempre una sfida per le nostre capacità cognitive. E anticipare il futuro significa indovinare l’esito più probabile di una data situazione, a partire da un certo numero di possibilità. Una delle strategie più ovvie per effettuare previsioni del futuro è sempre stata quella di appoggiarsi alla nostra esperienza del passato e, in particolare, alla frequenza con cui si sono verificati nel passato determinati eventi. La straordinaria abbondanza di “serie storiche” (cioè l’insieme di variabili casuali ordinate rispetto al tempo) dei fenomeni più disparati, da quelli fisici a quelli riguardanti gli individui e la società, sembrerebbero far

ben sperare. Uno dei criteri storicamente considerato più promettente è il cosiddetto “metodo degli analoghi”. Se in una serie temporale passata si scoprissero delle situazioni analoghe a quelle che viviamo oggi, potremmo dedurne, anche in assenza di un contesto teorico ben definito, che il futuro attuale potrebbe non essere molto diverso dal futuro di quel distante passato. Sappiamo ora che così non è e che il metodo degli analoghi è per la maggior parte dei casi fallimentare. Già James Clark Maxwell, verso la metà del XIX secolo, mise in guardia la comunità scientifica dai rischi di quella che considerava una “dottrina metafisica”, ossia che dagli stessi antecedenti derivino conseguenze simili. Da un punto di vista matematico furono i contributi di Henri Poincaré e Mark Kac a farci comprendere la fallacia dell’approccio: la chiave è che anche in un sistema deterministico il tempo necessario perché il sistema evolva in uno stato simile a quanto già

osservato in passato è tanto maggiore quanto meno probabile è questo stato. Benché in linea di principio si possa catturare la complessità di un dato fenomeno a partire dalle serie temporali delle sue osservazioni, dal punto di vista pratico tale operazione è resa impossibile non appena la complessità del fenomeno è minimamente grande. È il caso delle previsioni del tempo, oggi basate sulla simulazione di un modello matematico che ci permette, con l’aiuto determinante di dati sulle serie storiche, di proiettare il sistema nel futuro generando gli scenari più attendibili. Da tutto ciò evinciamo che un compromesso intelligente tra modellizzazione teorica e *machine learning* sembra essere la migliore strategia di previsione. Chiediamoci ora se il metodo degli analoghi, e più in generale la strategia di guardare al futuro con gli occhi del passato, non sia problematica anche per quelle che oggi chiamiamo intelligenze artificiali guidate da algoritmi di *machine learning*. Esemplare, a tal proposito, è

il caso delle auto senza guidatore della Volvo che, come riportato dal Guardian nell'estate del 2017 (vd. link sul web), si bloccarono totalmente davanti ai canguri australiani. Essendo macchine svedesi, "addestrate" su esempi di alci e renne, il primo incontro con un canguro si rivelò per loro un'esperienza incomprensibile. Questo esempio ci permette di mettere a fuoco una limitazione fondamentale degli attuali "algoritmi di inferenza", ossia di quel complesso di tecniche, sviluppate nella teoria della probabilità, che permettono di dedurre regolarità (in inglese "*patterns*") a partire da moli di dati, e stimare la probabilità di eventi futuri basandoci sull'osservazione del passato. L'enigma su cui generazioni di scienziati si sono misurati può essere enunciato nel modo seguente: come possiamo stimare la probabilità di un evento che non si è mai verificato in precedenza? Non ne abbiamo la minima idea! Questo è un problema molto generale, che riguarda anche le cosiddette macchine "intelligenti". Il problema generale con cui umani e macchine si confrontano è quello del

nuovo, di comprendere cioè cosa sarà possibile e, ancor di più, quale dei futuri possibili ha maggiori probabilità di realizzarsi.

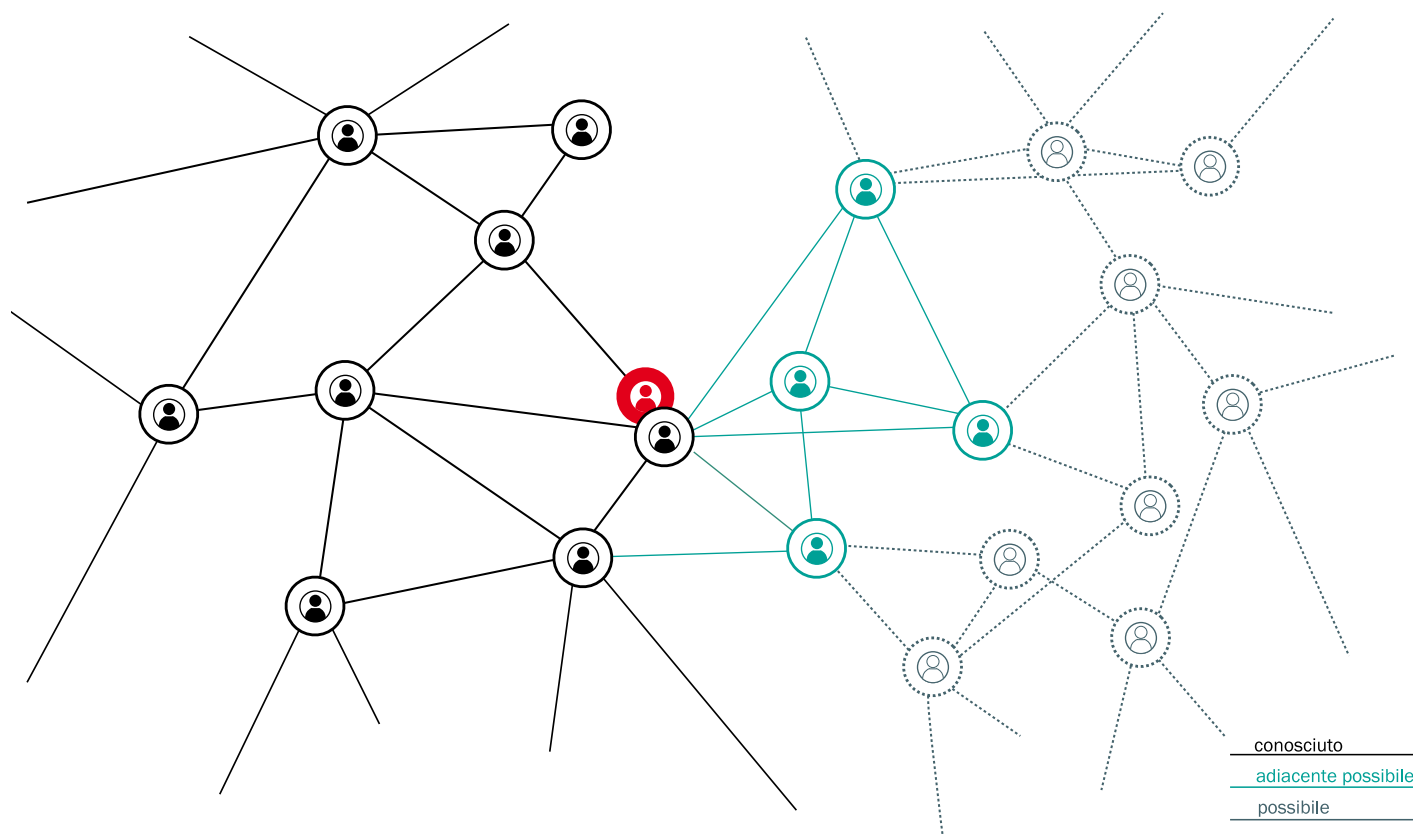
Queste considerazioni sono state splendidamente riassunte dal famoso biologo e scienziato dei sistemi complessi Stuart Kauffman. Kauffman ha coniato il termine "adiacente possibile", per indicare tutte quelle cose (che siano molecole, idee o strumenti tecnologici) che sono a un passo da ciò che esiste realmente e che quindi possono derivare da modifiche incrementali e ricombinazioni di materiale esistente (vd. fig. d, ndr). Nelle parole di Steven Johnson (giornalista e scrittore che si occupa di innovazione tecnologica): "La strana e bella verità sull'adiacente possibile è che i suoi confini crescono man mano che si esplorano". Questo è il punto chiave dell'adiacente possibile: l'espansione (o la ristrutturazione) dello spazio delle possibilità, condizionata all'esplorazione dello spazio stesso e all'esperienza del nuovo. La nozione di adiacente possibile è affascinante. Racconta una storia su

uno spazio, lo spazio delle possibilità, molto difficile da tracciare. Eppure questo è lo spazio che esploriamo e modifichiamo costantemente nella nostra vita quotidiana mentre prendiamo decisioni e agiamo. Sarebbe desiderabile sapere molto di più su questo spazio, sulla sua struttura, sulla sua ristrutturazione, sulle sue dinamiche, sul modo in cui lo esploriamo, individualmente o collettivamente. Recentemente, in collaborazione con Vito Servedio, Steven Strogatz e Francesca Tria, abbiamo introdotto una struttura matematica per indagare le dinamiche del nuovo attraverso l'adiacente possibile. Il modello teorico fornisce una struttura matematica all'idea dell'adiacente possibile e sviluppa l'idea che i processi di innovazione non siano lineari e che le condizioni per il verificarsi di un determinato evento potrebbero realizzarsi solo dopo che è successo qualcos'altro. YouTube, per esempio, non avrebbe potuto svilupparsi prima dell'avvento della banda larga.



c.  
Nel 2017 la Volvo testò delle auto senza guidatore, che si bloccarono totalmente davanti ai canguri australiani perché, essendo macchine svedesi, erano "addestrate" su esempi di alci e renne. Un esempio perfetto della fallacia del metodo degli analoghi nell'intelligenza artificiale.





L'espansione dell'adiacente possibile, condizionata al verificarsi di una novità, è l'ingrediente cruciale in questo schema teorico, per derivare previsioni testabili e quantitative. Per esempio, le previsioni sul tasso di innovazione, cioè della velocità con cui osserviamo nuove cose, insieme alle correlazioni tra eventi di innovazione, l'emergere di ondate di novità, tendenze, ecc. Le previsioni basate sulla nozione di adiacente possibile sono state validate in moltissimi sistemi sociali e tecnologici: dalla scrittura di testi allo sviluppo collaborativo delle pagine di Wikipedia, dalla fruizione della musica alla scrittura di open software, dalla creazione di hashtag su Twitter alla scrittura di brevetti. Si può dire che, al di là di una sterile contrapposizione tra intelligenza umana e intelligenza artificiale, per affrontare le sfide scientifiche e sociali del futuro sarebbe interessante potersi muovere verso una cosiddetta "creatività aumentata". Questo implica comprendere in che modo la straordinaria capacità umana di creare, di concepire modelli e intuire soluzioni possa essere supportata dalle nuove tecnologie nella ricerca di percorsi rari che possano

portare a innovazioni utili e dirompenti. Forse ci è data un'opportunità unica per innescare un circolo virtuoso, in cui le macchine possono evolversi per imparare dalla natura ancora inafferrabile della creatività umana e gli esseri umani possono essere stimolati e sfidati nelle loro attività creative dall'onnipresenza di una nuova generazione di macchine.

Con un minimo di riflessione, possiamo identificare due meccanismi principali attraverso i quali le macchine ci stanno già aiutando a essere più creativi. Da un lato c'è quello che potremmo definire un supporto "passivo". L'esempio forse più eclatante è la sfida che ha opposto AlphaGo, un programma per computer che riproduce il gioco da tavolo Go, a Lee Sedol, uno dei migliori giocatori di Go di sempre. Sedol ha perso 4 partite su 5 ed è riuscito a vincere una sola partita grazie a una mossa straordinariamente sorprendente, divenuta famosa tra gli esperti come "mossa 78". AlphaGo ha semplicemente giudicato quella mossa troppo improbabile (uno su diecimila) da giocare e quella mossa è stata la pietra angolare dell'unica vittoria di Lee Sedol.

d. Schema semplificato dell'adiacente possibile. In questo esempio, il pallino rosso rappresenta un individuo che si muove in un certo spazio (ad esempio uno spazio di amicizie). I cerchietti neri rappresentano le persone già conosciute mentre quelli verdi rappresentano l'adiacente possibile dell'individuo, ossia l'insieme delle persone che potrebbe incontrare in un prossimo futuro. Il punto cruciale è che l'insieme dei cerchietti verdi diventa disponibile solo nel momento in cui si effettua un nuovo incontro. Ad esempio, si fa una nuova conoscenza e solo allora gli amici del nuovo amico diventano potenziali nuovi amici. L'insieme dei cerchietti grigi, infine, rappresenta lo spazio delle persone che potrebbero essergli presentate solo dopo aver incontrato alcuni degli amici dell'adiacente possibile.



e.  
Lee Sedol al World Economic Forum  
in Tianjin (China) nel 2016.

In questo caso specifico, la macchina programmata per vincere in modo efficiente i giochi di Go ha creato la condizione per Lee Sedol di andare oltre i suoi limiti e trovare una soluzione originale per il cambio di passo. Un altro modo, forse ancora più interessante, in cui le macchine possono supportare e aumentare la nostra creatività, è attraverso un ruolo più attivo. Invece di creare ostacoli da superare, le macchine potrebbero suggerire nuovi percorsi da esplorare nello spazio dell'adiacente possibile. Percorsi forse molto rari, che potrebbero stimolare ulteriormente l'immaginazione umana circa i loro potenziali esiti. Gli esempi sono davvero molti. È ad esempio il caso dei progetti di strumenti di co-creazione per musica o testi. O dei progetti per la creazione di piattaforme per la sostenibilità, il cui obiettivo è quello di comprendere il presente, attraverso i dati, e proiettarlo nel futuro attraverso una previsione realistica di come un cambiamento nelle condizioni attuali influenzerà e modificherà lo scenario futuro (i cosiddetti "*whatif games*"). Anche in questo caso l'idea è di supportare la creatività umana nella ricerca di nuove soluzioni efficaci e sostenibili con l'aiuto di tecniche di apprendimento automatico o di agenti

personalizzati artificiali. Il percorso verso gli scenari sopra menzionati è lungi dall'essere lineare e facile. Se l'intelligenza artificiale può trattare sorprendentemente bene "problemi difficili", ha difficoltà a capire i contesti e le strategie da adottare in ambienti mutevoli. Ancora più difficile è la nuda applicazione delle tecniche di intelligenza artificiale in situazioni in cui lo spazio delle possibilità non è solo ampio e complesso, ma anche aperto. Ciò implica un cambio di paradigma: inventare nuove soluzioni invece di cercare soluzioni. E questa è la situazione

più comune quando si svolge ricerca scientifica, quando si affrontano problemi sociali complessi, quando si apprende o quando ci si esprime artisticamente. Secondo qualcuno, l'autoreferenzialità - intesa come l'abilità di un sistema cognitivo di osservare sé stesso da un livello più alto di astrazione - è una chiave per capire veramente l'intelligenza. Sulla stessa linea, l'interazione tra intelligenza umana e artificiale potrebbe essere fondamentale anche per capire meglio i percorsi cognitivi alla base della nostra intelligenza e sfruttarli per affrontare efficacemente le sfide della società.

#### Biografia

**Vittorio Loreto** è professore di Fisica dei Sistemi Complessi presso la Sapienza Università di Roma e membro di Facoltà del Complexity Science Hub di Vienna. Attualmente dirige il Sony Computer Science Lab a Parigi, dove guida il gruppo di "Innovation, Creativity and Artificial Intelligence". La sua attività scientifica è focalizzata sulla fisica statistica dei sistemi complessi e le sue applicazioni interdisciplinari. Ha coordinato diversi progetti a livello europeo.

#### Link sul web

<https://www.wired.com/2008/06/pb-theory/>

<https://www.theguardian.com/technology/2017/jul/01/volvo-admits-its-self-driving-cars-are-confused-by-kangaroos>

<http://whatif.cslparis.com/>

[https://www.ted.com/talks/vittorio\\_loreto\\_need\\_a\\_new\\_idea\\_start\\_at\\_the\\_edge\\_of\\_what\\_is\\_known/](https://www.ted.com/talks/vittorio_loreto_need_a_new_idea_start_at_the_edge_of_what_is_known/)

DOI: 10.23801/asimmetrie.2019.26.2

# Caduti nella rete (neurale)

## Gli algoritmi dell'intelligenza artificiale

di Stefano Giagu



a.  
Prototipo di una scheda elettronica del sistema di trigger per muoni che verrà installato nell'esperimento Atlas. Nel centro della scheda è visibile il processore di tipo Fpga.

Il concetto di intelligenza artificiale esiste dalla metà degli anni '50 e lo studio e sviluppo di algoritmi di *machine learning*, che ne costituiscono il motore computazionale, è in corso da più di venti anni, ma è solo recentemente che le applicazioni hanno cominciato a diffondersi esponenzialmente in tutti i settori della società, incluso l'ambito della ricerca di base in fisica delle particelle. L'idea di "rete neurale artificiale" nasce nel 1943 come un modello di tipo connessionista, proposto dal neurofisiologo Warren McCulloch e dal logico matematico Walter Pitt, per descrivere il funzionamento dei neuroni nel cervello umano. Tale idea si è evoluta negli anni, grazie soprattutto all'algoritmo "*perceptron*" di Frank Rosenblatt, psicologo americano pioniere nella ricerca sui sistemi cognitivi alla Cornell University, in grado di imparare ad associare informazioni in input con risposte, in modo apparentemente simile all'apprendimento negli umani, e successivamente nei moderni algoritmi di apprendimento basati su "*backpropagation*", che costituiscono la spina dorsale delle moderne architetture di reti neurali artificiali in uso oggi. Le più recenti evoluzioni degli algoritmi di *machine learning* basate sul *deep learning* (in pratica, un qualunque algoritmo basato su una o più reti neurali) stanno trovando un terreno molto fertile nelle grandi collaborazioni al Large Hadron Collider (Lhc) del Cern, caratterizzate da apparati costituiti da milioni di elementi, i cui segnali producono enormi quantità

di dati complessi che devono essere analizzati e semplificati per poterne distillare l'informazione utile. Tradizionalmente il problema viene affrontato tramite una serie di passaggi, che prendono il nome di "analisi dati", che operano sugli eventi per ridurne la complessità. Questo approccio, che ha permesso scoperte fondamentali, segna però il passo di fronte all'aumento della complessità e della quantità dei dati prodotti dagli esperimenti, e in tali condizioni il *deep learning*, grazie alla capacità di apprendere e condensare le informazioni rilevanti, può fornire l'ingrediente necessario a colmare il divario tra l'approccio tradizionale e quello ottimale. Applicazioni di intelligenza artificiale in fisica delle particelle non si limitano all'analisi dati. La flessibilità e la velocità delle reti neurali, che possono essere eseguite in tempo reale su dispositivi di calcolo veloci come le Field Programmable Gate Array (Fpga) (vd. fig. a), possono essere utilizzate per rimpiazzare i sistemi hardware di selezione degli eventi prodotti nelle collisioni di Lhc, più costosi e complessi da progettare e mantenere (vd. in *Asimmetrie* n. 17 p. 33, ndr). Il *deep learning* viene utilizzato anche nel campo della ricostruzione "offline" degli eventi, per individuare le traiettorie delle particelle elementari prodotte nelle collisioni, per identificarne la tipologia e, più recentemente, per simularne il comportamento all'interno dei rivelatori di Lhc, un passo cruciale per confrontarsi con le previsioni dei modelli teorici.



Tra i diversi algoritmi di *machine learning*, i Boosted Decision Tree (Bdt) (vd. p. 19, ndr) hanno tradizionalmente avuto un ruolo di primo piano nell'analisi dati degli esperimenti per le intrinseche caratteristiche di semplicità e trasparenza. I Bdt combinano insieme un numero molto grande di algoritmi elementari chiamati "alberi di decisioni binarie". Questi emulano la sequenza di domande/decisioni utilizzata per esempio nel gioco "Indovina Chi?" per cercare di individuare un personaggio segreto sulla base di una serie di domande sull'aspetto del personaggio, alle quali si possono ottenere solo risposte binarie del tipo "vero-falso". Ad ogni domanda successiva il campione di possibilità (i personaggi nel gioco) si restringe sempre più fino idealmente ad arrivare alla risposta corretta. La stessa procedura è utilizzata dai fisici sperimentali quando cercano di identificare eventi appartenenti a un processo fisico interessante, per esempio eventi contenenti un bosone di Higgs che decade in due fotoni, tra i tantissimi eventi prodotti a Lhc. Gli eventi interessanti vengono selezionati tramite una sequenza di richieste su diverse grandezze fisiche (l'equivalente delle domande nel gioco), scelte in modo da massimizzare la purezza del segnale che si vuole selezionare. Quali grandezze fisiche utilizzare e quale tipo di richieste fare costituisce l'addestramento dell'algoritmo, che viene effettuato sulla base di campioni di eventi di esempio, spesso simulati in modo da riprodurre il segnale anche prima di averlo trovato. Un Bdt combina diversi alberi di decisione binaria per produrre un *ensemble* di classificatori più efficace e potente. Ancora più promettenti sono gli algoritmi basati sulle "reti neurali profonde", in particolare le cosiddette "reti convoluzionali" o Cnn (Convolutional Neural Network), sviluppate nel campo della

*computer vision* per il riconoscimento di immagini fotografiche. Le Cnn sono specializzate nell'identificazione di caratteristiche geometriche ricorrenti presenti nelle immagini con cui vengono addestrate. Caratteristiche che possono essere combinate in strutture complesse, per esempio volti di persone (vd. in Asimmetrie n. 20 approfondimento p. 33, ndr) o specifici oggetti, animali, ecc., con lo scopo di riconoscerle successivamente in altre immagini. Molti dei segnali ricostruiti dai rivelatori degli esperimenti di fisica delle particelle si prestano a essere rappresentati come immagini fotografiche, in cui ad ogni pixel dell'immagine corrisponde per esempio l'energia rilasciata da una particella in una delle zone sensibili del rivelatore. Una rete Cnn può quindi essere addestrata a riconoscere caratteristiche specifiche nelle "immagini" dei rivelatori, e quindi per esempio a identificare in tali immagini la presenza di un decadimento del bosone di Higgs, al pari di quanto fa una Cnn in uno smartphone per riconoscere il volto del proprietario.

Una variante interessante è costituita dalle "reti neurali ricorrenti" o Rnn (Recurrent Neural Network), algoritmi sviluppati per l'elaborazione del linguaggio naturale. Esempio tipico è il funzionamento dei traduttori online, dove la traduzione di una frase viene raffinata man mano che si aggiungono parole (vd. p. 19, ndr). Analogamente, le Rnn analizzano ogni elemento di una sequenza di dati di input mantenendo "memoria" di ciò che è stato calcolato precedentemente, e sono molto efficaci nell'elaborare le lunghe sequenze di dati correlati che compaiono in molti problemi in fisica delle particelle. Tipico esempio è quello dell'identificazione di quark beauty prodotti nelle interazioni di alta energia di Lhc. I quark beauty si presentano sotto forma di



b.

Il gioco "Indovina Chi?" ci fornisce un esempio di "albero di decisioni binarie": per cercare di indovinare il personaggio segreto dell'avversario, il giocatore fa una serie di domande sull'aspetto del personaggio, scartando man a mano i personaggi non pertinenti (p.es. i personaggi con gli occhiali o con i capelli biondi, ecc.).



un flusso collimato di particelle, chiamato “getto adronico”, che vive per un certo tempo prima di decadere in altre particelle. Identificare questi getti adronici è molto impegnativo dal punto di vista computazionale per cui negli algoritmi tradizionali si introducono semplificazioni che ne limitano le prestazioni. L’uso di Rnn ha permesso di migliorare sostanzialmente la precisione con cui i quark beauty vengono identificati negli esperimenti di Lhc (in particolare in Atlas e Cms), estendendone la sensibilità di scoperta per effetti di nuova fisica.

Infine, un’altra classe di reti neurali chiamate “reti generative avversarie” o Generative Adversarial Network (Gan) (vd. fig. c), possono essere sfruttate per simulare in modo preciso e veloce il comportamento atteso dei rivelatori in presenza di diversi segnali fisici interessanti. Una Gan è fatta da due reti neurali profonde, il generatore e il discriminatore. La prima genera immagini, mentre la seconda riconosce un’immagine reale da una generata. Le due reti vengono addestrate “affrontandosi” l’una contro l’altra. In pratica il discriminatore cerca di identificare le differenze tra le immagini di un campione vero (per esempio immagini ricostruite dai segnali di un rivelatore colpito da un dato tipo di particelle) e quelle generate dal generatore, mentre quest’ultimo cerca di ingannare il discriminatore nell’accettare le sue immagini, e così facendo impara a generare immagini realistiche. I risultati ottenuti con questa tecnica sono molto incoraggianti e, se le promesse iniziali verranno confermate, l’utilizzo delle Gan permetterà in futuro di ridurre sostanzialmente il costo computazionale necessario per simulare l’enorme mole di eventi richiesta per la fase di funzionamento ad altissima intensità dell’acceleratore Lhc.

c.  
Immagini di persone non esistenti, generate da una rete neurale di tipo Gan.

#### Biografia

**Stefano Giagu** insegna alla Sapienza Università di Roma e svolge attività di ricerca in fisica delle particelle elementari e nello sviluppo di metodi di intelligenza artificiale e *machine learning* in fisica di base e applicata. Ha partecipato agli esperimenti L3 al Cern e Cdf al Fermilab ed è membro della collaborazione Atlas, al Large Hadron Collider del Cern di Ginevra, della quale dal 2019 è responsabile nazionale.

#### Link sul web

[https://it.wikipedia.org/wiki/Apprendimento\\_automatico](https://it.wikipedia.org/wiki/Apprendimento_automatico)  
[https://en.wikipedia.org/wiki/Artificial\\_neural\\_network](https://en.wikipedia.org/wiki/Artificial_neural_network)  
[https://en.wikipedia.org/wiki/Deep\\_learning](https://en.wikipedia.org/wiki/Deep_learning)  
<https://cerncourier.com/the-rise-of-deep-learning/>  
<https://www.datascience.com/blog/convolutional-neural-network>

DOI: 10.23801/asimmetrie.2019.27.3



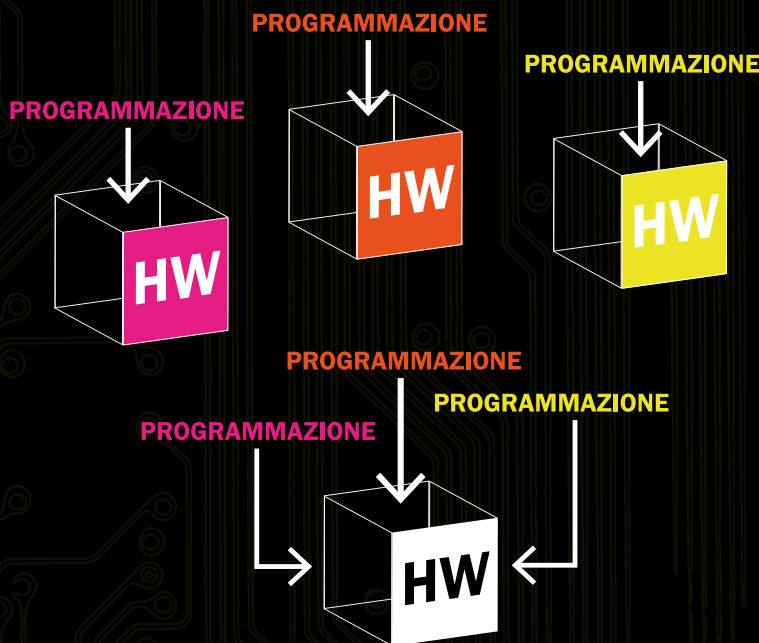
# [as] Macchine intelligenti.

## FPGA/ASIC

Gli FPGA (Field Programmable Gate Array) sono costituiti da molti blocchi logici elementari, la cui funzione è definita in fase di programmazione. Sono riprogrammabili a differenza degli ASIC (Application Specific Integrated Circuit), molto più convenienti se prodotti in gran numero.

## CPU/GPU

Le CPU (Central Processing Unit) e le GPU (Graphics Processing Unit, di fatto una CPU ottimizzata per la gestione della grafica digitale) possono eseguire molti tipi di programmi che non modificano l'hardware su cui girano.



## PROGRAMMAZIONE CLASSICA

La programmazione "classica" definisce - attraverso uno dei tanti linguaggi disponibili - una sequenza di operazioni da eseguire in modo da ottenere il risultato voluto (l'output) a partire da una serie di dati in ingresso (l'input).



INPUT



PROGRAMMA



OUTPUT

## PROGRAMMAZIONE RETI NEURALI E DATA MINING

Dall'ideazione nel 1958 del perceptron, un modello matematico ipersemplificato del funzionamento dei neuroni, si è sviluppata la programmazione delle reti neurali (BDT, CNN, RNN, GAN,...) dove la serie di passi tra l'input e l'output è sempre meno sotto il controllo diretto del programmatore e i programmi migliorano man mano che analizzano dati (*machine learning*).



INPUT



PROGRAMMA



OUTPUT

L'evoluzione in corso in questi anni prevede l'analisi della più grande quantità possibile di dati alla ricerca di informazioni o relazioni interessanti che aiutino a risolvere problemi. Il programmatore ha un ruolo marginale nella definizione sia dell'input che dell'output.



INPUT



PROGRAMMA



OUTPUT

## BDT

### BOOSTED DECISION TREE

Un "albero delle decisioni" usa alcune caratteristiche dei dati in ingresso per separarli in modo ricorsivo in base alle caratteristiche stesse possedute da ciascun dato. A ciascun "nodo" i dati sono divisi in base al valore di una delle caratteristiche. Ciascun nodo terminale o "foglia" è una possibile classificazione del dato in ingresso, spesso accompagnata da un valore di probabilità. L'uso di più alberi diversi sullo stesso campione di dati amplifica l'abilità nel separare i dati nelle categorie volute.

## CNN

### CONVOLUTIONAL NEURAL NETWORK

Le CNN sono usate principalmente per classificare immagini, raggrupparle per somiglianza o identificare oggetti all'interno delle immagini stesse. Usando livelli successivi (quelli intermedi sono un esempio di *deep learning*) che identificano e assegnano diversa importanza a varie caratteristiche, le CNN distinguono i diversi "oggetti" rappresentati, arrivando a riconoscere volti, distinguere segnali stradali, identificare formazioni tumorali o qualsiasi altro aspetto che sia rappresentabile graficamente.

## RNN

### RECURRENT NEURAL NETWORK

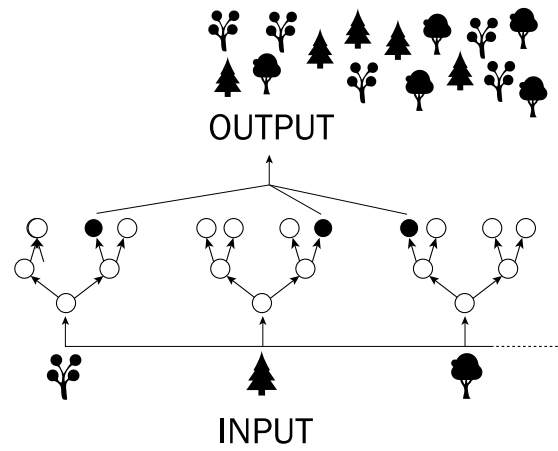
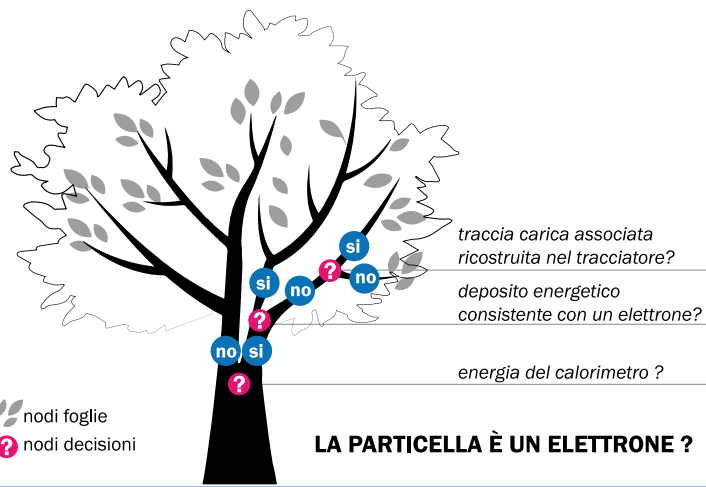
A differenza di altre reti, le RNN usano la struttura temporale dei dati, modificano il risultato ottenuto mano a mano che vengono analizzati i dati in ingresso. Le principali applicazioni sono nei traduttori automatici e nel riconoscimento vocale, nel controllo dei robot e nella predizioni di comportamenti futuri a partire da dati registrati come le previsioni del meteo o i comportamenti dei mercati finanziari.

## GAN

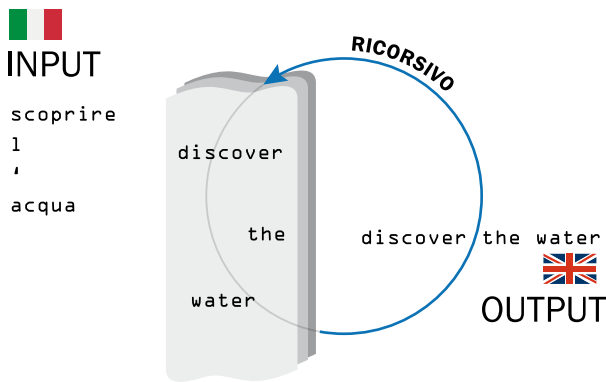
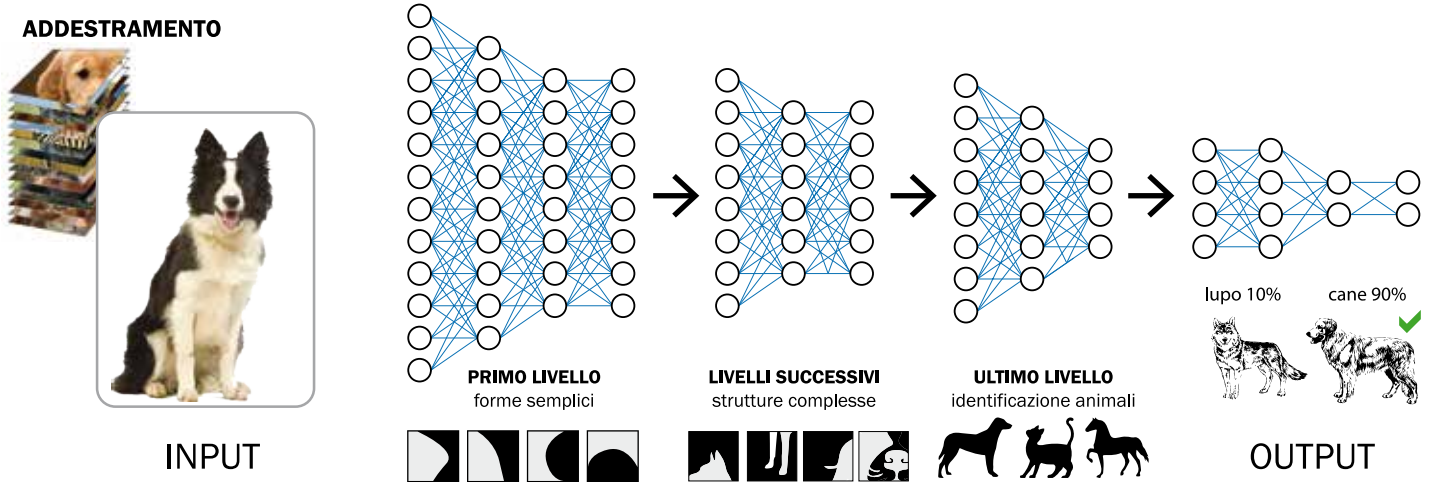
### GENERATIVE ADVERSARIAL NETWORK

La caratteristica principale delle GAN è l'uso simultaneo e ricorsivo di due reti neurali. La prima, il generatore, analizza forme note e le trasforma nell'output voluto, per esempio un volto. La seconda, il discriminatore, confronta l'output generato con campioni di dati, cercando di distinguerli. Il risultato di questo confronto viene usato dal generatore per migliorare la qualità dell'output generato.

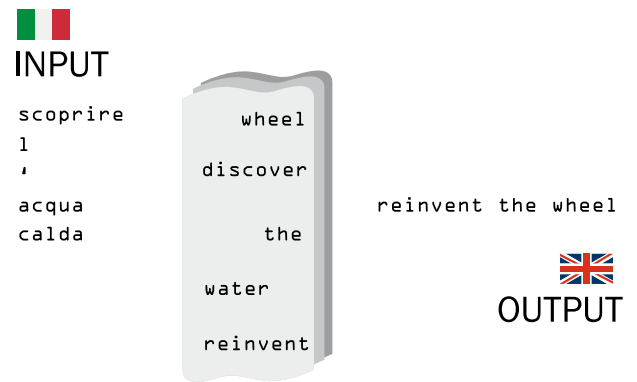




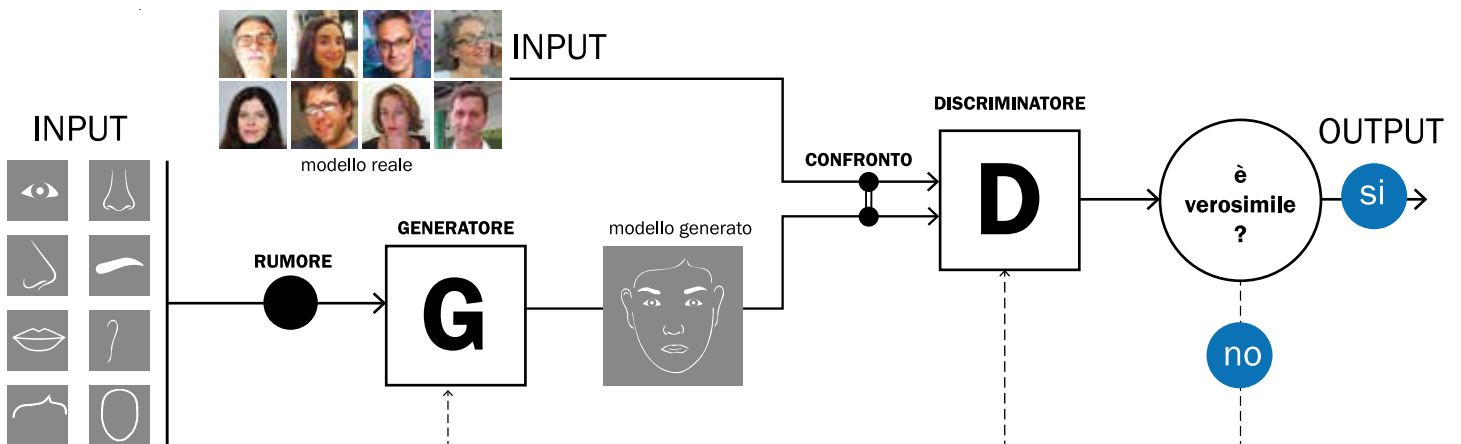
## ADDESTRAMENTO



## PROCESSI DI TRADUZIONE NASCOSTI



## PROCESSI DI TRADUZIONE NASCOSTI

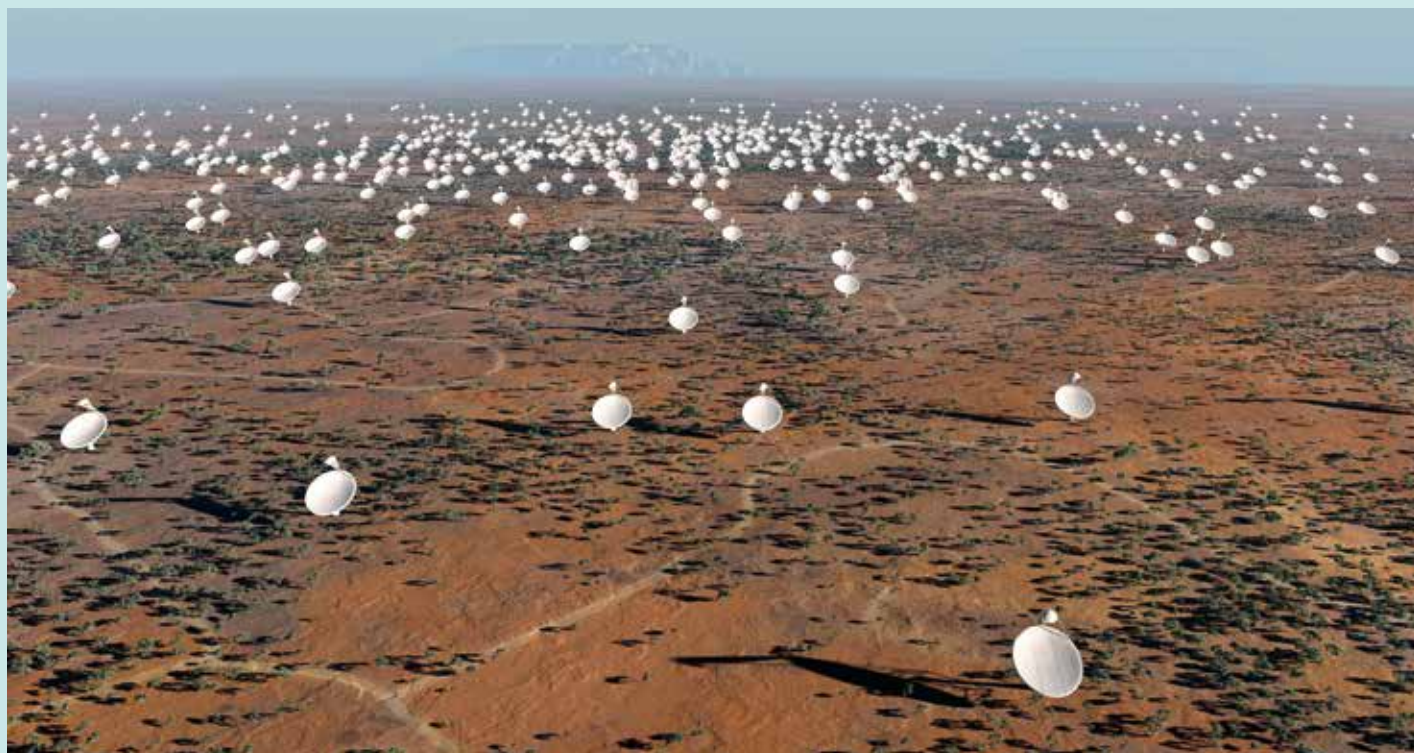


# Questioni di *core*

## Processori per i *big data*

di Gaetano Maron

a.  
Visualizzazione artistica del nucleo centrale, del diametro di circa 5 km, del radiotelescopio Square Kilometre Array (Ska). A regime il progetto prevede l'installazione di 130.000 antenne distribuite nei due siti in Sud Africa e Australia, che permetterà di avere una zona complessiva di raccolta delle onde radio provenienti dallo spazio di circa 1 km<sup>2</sup>.



L'analisi dati nella fisica sperimentale, in particolare in quella delle alte energie, è oggi dominata, per dimensioni, potenza di calcolo necessaria e complessità degli algoritmi, dagli esperimenti del Large Hadron Collider (Lhc) del Cern. Anche la complessità e le dimensioni degli esperimenti di fisica astroparticellare e di astrofisica sta rapidamente evolvendo e la loro analisi dati presenta ormai caratteristiche simili a quelle di Lhc. Nel prossimo decennio, due giganti della *Big Science* inizieranno la presa dati: Hilumi-Lhc (HI-Lhc), la fase 2 di Lhc, e Ska (Square Kilometre Array Telescope), il più grande radiotelescopio mai costruito al mondo, che sarà localizzato in Sud Africa e in Australia.

Questo gruppo di esperimenti di nuova generazione produrrà dati al ritmo degli exabyte (vd. fig. b per il significato dei prefissi) all'anno, che equivale grosso modo alla capacità di *storage* attuale di tutti i centri di calcolo di Lhc distribuiti nel mondo. Anche l'analisi di quei dati richiederà potenze di calcolo ragguardevoli e, per il caso di Hilumi-Lhc, dell'ordine di 10 milioni di *core* di calcolo, circa 20 volte quelli disponibili attualmente per Lhc. Le Cpu (Central Processing Unit, v. anche p. 18, ndr) di un calcolatore, infatti, non sono più "monolitiche" (ossia costituite da un unico circuito integrato), ma contengono all'interno del proprio chip (o "*die*") più unità di calcolo indipendenti

tra loro che vengono chiamate *core*: una moderna Cpu può essere composta da qualche decina di *core*.

Queste risorse di calcolo pongono problemi sia tecnologici che di sostenibilità dei costi. La fig. b ci propone un interessante confronto, anche se solo indicativo, tra i dati prodotti da questi grandi esperimenti scientifici e quelli prodotti dai principali social network o dall'archivio Internet di Google. È alquanto impressionante notare come gli ordini di grandezza siano molto simili e che i casi scientifici abbiano a volte dimensioni anche maggiori. Nella figura si fa distinzione tra dati grezzi (in inglese "*raw data*") e dati elaborati (in inglese

“*physics data*”). I primi sono i dati come vengono raccolti dai rivelatori dell’esperimento e contengono informazioni su quale elemento del rivelatore è stato colpito, sull’energia rilasciata, sul tempo, sulla posizione. Sono dati unici e non più riproducibili e per questa ragione ne vengono eseguite più copie preservate poi al Cern e negli altri centri di calcolo di Lhc. I dati elaborati derivano invece dal processamento dei dati grezzi, al fine di ricostruire l’impulso delle particelle, la loro origine, la loro energia, il tipo di particella rivelata, ecc., e contengono anche tutti i risultati finali prodotti dall’analisi dei canali di fisica studiati che possono portare a eventuali scoperte.

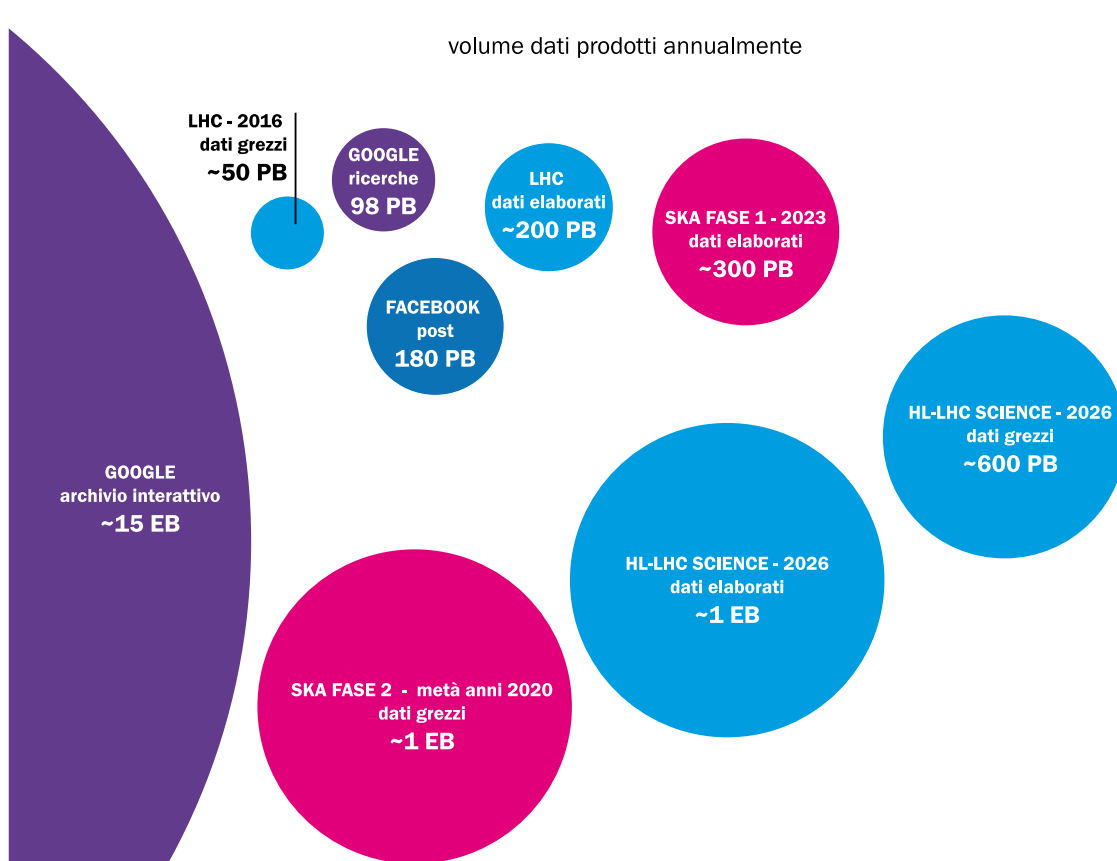
La *data science*, uno dei settori tecnologici a più elevato tasso di crescita, è nata sotto la spinta a creare sistemi efficienti di analisi dei *big data* dell’ordine di grandezza di quelli rappresentati in figura. Le tecnologie sviluppate in questo ambito utilizzano efficienti tecniche di memorizzazione e lettura dei dati, sofisticate tecniche statistiche per effettuare analisi predittive e, sempre più massicciamente, tecniche di intelligenza artificiale. Elemento comune a queste tecnologie è la possibilità di effettuare operazioni massicciamente parallele e avere architetture fortemente scalabili, anche a dimensioni elevate. L’analisi dati nella fisica delle alte energie utilizza già tecniche di *machine learning* (vd. p. 15, ndr) e sarà cruciale poter includere queste tecniche di “*big data analytics*” per far fronte alla complessità degli eventi di Hilumi-Lhc.

Uno degli aspetti che limitano la scalabilità dei software di analisi attuali, in particolare dei grandi esperimenti, è l’incapacità di sfruttare al meglio il parallelismo offerto dalle

nuove generazioni di processori (*multi-core* Cpu, acceleratori massicciamente paralleli, *graphic processor*, ecc.). Per Hilumi-Lhc questo comporta una re-ingegnerizzazione del codice di analisi esistente, ottimizzandolo e integrandolo con le tecniche di analisi menzionate più sopra, spesso già predisposte al calcolo parallelo. Il software, inteso in senso lato, è quindi uno dei punti nevralgici dell’evoluzione dei sistemi di analisi verso l’era degli “*exa-esperimenti*”, forse l’aspetto che più di altri avrà impatto sulle prestazioni e sulla scalabilità del sistema complessivo.

Evoluzione che nel settore dei processori spinge sempre più verso sistemi a elevato numero di *core* di calcolo per singolo Cpu, avendo ormai raggiunto limiti fisici alla potenza del singolo *core*: la famosa legge di Moore (secondo cui la potenza di calcolo del *core* raddoppia ogni 18/24 mesi), infatti, da qualche anno mostra evidenti segni di rallentamento. Questo approccio *multi-core* viene estremizzato nei *graphic processor*, sempre più utilizzati come acceleratori matematici in supporto alle Cpu convenzionali (vd. approfondimento). La tendenza sempre più marcata sarà quindi quella di sistemi ibridi formati da Cpu convenzionali, combinati con la potenza di calcolo straordinaria degli acceleratori.

Questo stesso approccio è seguito dai supercomputer di nuova generazione (detti anche “*computer exascale*”, perché possono eseguire “*exa*”, cioè miliardi di miliardi, di operazione al secondo). Questa comune linea di tendenza favorirà l’uso di queste gigantesche macchine di calcolo da parte della fisica sperimentale. È un approccio già in atto, che verrà consolidato in futuro e che potrà contribuire in modo significativo alle necessità



b.  
La rivoluzione dei *big data*:  
i social e i motori di ricerca  
vi contribuiscono in maniera  
predominante, ma anche i grandi  
esperimenti scientifici agiscono  
da protagonisti.

**kB** kilobyte =  $10^3$  byte  
**MB** megabyte =  $10^6$  byte  
**GB** gigabyte =  $10^9$  byte  
**TB** terabyte =  $10^{12}$  byte  
**PB** petabyte =  $10^{15}$  byte  
**EB** exabyte =  $10^{18}$  byte  
**ZB** zettabyte =  $10^{21}$  byte  
**YB** yottabyte =  $10^{24}$  byte



di calcolo dei grandi esperimenti.

L'Infn ha già fatto passi significativi in questa direzione, stringendo un accordo strategico con il consorzio interuniversitario per il supercalcolo Cineca grazie a cui verrà finanziato “Leonardo”, una macchina di calcolo con una capacità grosso modo pari a un quarto di un exascale. Leonardo verrà installato presso il Tecnopolo di Bologna, dove sarà trasferito anche il principale centro di calcolo dell'Infn, il Cnaf (vd. fig. c). In questo modo si realizzerà un grande centro di calcolo unico a livello europeo per potenza complessiva di calcolo e per capacità di memorizzazione e di analisi di dati sperimentali: una struttura fondamentale per il progetto Hilumi-Lhc.

Potenze di calcolo enormi sono e saranno sempre più disponibili anche dai *cloud provider* commerciali con una linea di tendenza dei costi di utilizzo marcatamente al ribasso.

Questo scenario ibrido, dai sistemi privati ai supercomputer e ai sistemi *cloud* commerciali, che disegna la possibilità di un sistema di calcolo a “*plug in*”, dove, a seconda del bisogno e della convenienza economica, potrà essere utilizzato uno o più di questi sistemi, sottolinea la necessità di una radicale evoluzione dei modelli di calcolo che includa anche e soprattutto l'evoluzione del software di analisi tracciato qui sopra. Al centro di questo disegno c'è un sistema robusto e veloce di gestione unica dei dati, basato su una decina di grandi “*storage center*” distribuiti a livello planetario, interconnessi tra loro da reti ad altissima velocità e progettati in modo da essere visti come un'unica entità.

Non sono prevedibili, ad oggi, svolte tecnologiche nei prossimi dieci anni tali da sconvolgere le previsioni qui mostrate, anzi si nota un certo rallentamento degli sviluppi industriali dovuti a

problemi di saturazione del mercato. La tecnologia del “computer quantistico”, che per la sua natura intrinsecamente parallela potrebbe invece avere un effetto dirompente per le nostre applicazioni, è ancora a livello sperimentale. L'interesse però è fortissimo e numerosi progetti di ricerca e sviluppo sono stati proposti dalle comunità, spesso in collaborazione con le industrie del settore.

È comunque interessante notare come la rivoluzione digitale che ci sta portando a una economia e quindi a una società *data driven*, sia progredita e si stia sviluppando quasi contemporaneamente ai grandi esperimenti scientifici di ultima generazione, in particolare di fisica delle alte energie, di astrofisica, ecc., per definizione anch'essi *data driven* e ora capaci di produrre quantità enormi di dati. L'industria informatica sta investendo ingenti capitali in ricerca e sviluppo in questo settore per facilitare decisioni aziendali strategiche e quindi accrescere la competitività delle aziende che utilizzano queste nuove tecnologie. Se sostituiamo la parola “aziende” con “esperimenti”, diventa subito evidente come questi sviluppi possano essere di fondamentale supporto alla ricerca scientifica laddove c'è una massiccia produzione di dati. A riprova di ciò va citato il recente accordo tra il Cern e Google, che avrà “l'obiettivo di realizzare progetti congiunti in ambito del *cloud computing*, *machine learning* e *quantum computing*”.

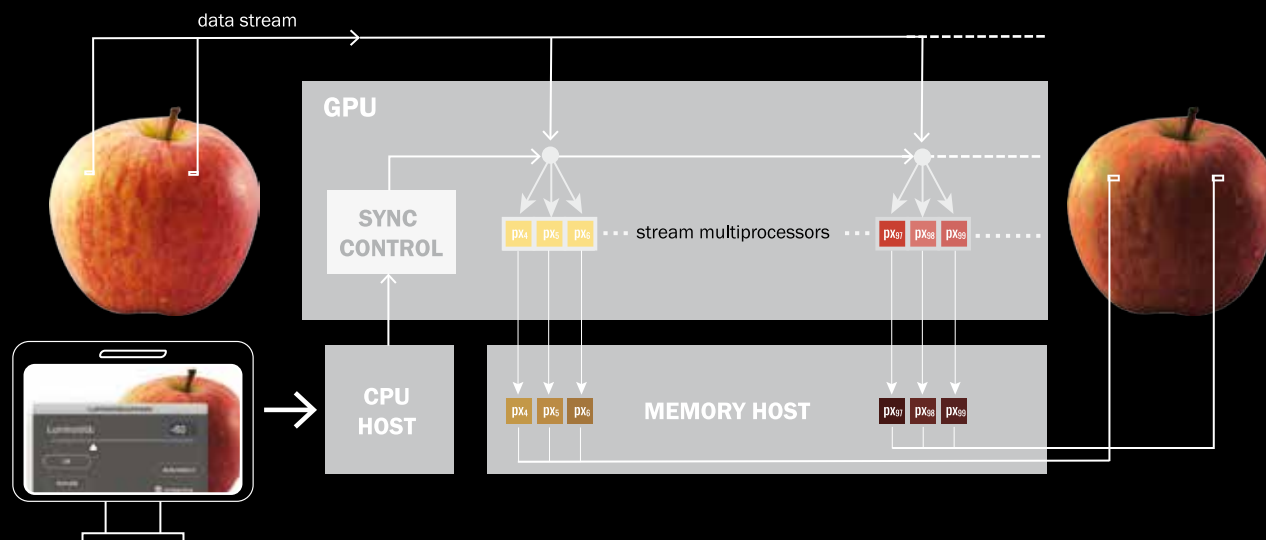
Com'è successo per lo sviluppo del web e in innumerevoli altre situazioni, anche in questo caso la scienza ha l'occasione di “adottare” tecnologie già esistenti o in corso di sviluppo per poi rielaborarle inventando per i propri fini una nuova tecnologia che spesso sarà oggetto a sua volta di trasferimento di conoscenza e ricaduta verso il mondo dell'industria e della società.



c.  
Il robot del Cnaf (una specie di silos orizzontale lungo 10 m e largo 2,5 m) che gestisce un archivio di 10.000 “*tape*”, dove sono memorizzati i dati di Lhc e dei principali esperimenti dell'Infn. Ogni *tape* (in italiano “nastro”) è una cassetta magnetica con dimensioni circa il doppio di quelle delle cassette che si utilizzavano molti anni fa per ascoltare musica.

# Processori grafici

1.  
Rappresentazione schematica del flusso dei dati in una Gpu, in risposta al comando di ridurre la luminosità dell'immagine di una mela.



Le “*graphics processing units*”, in breve Gpu, furono progettate principalmente per velocizzare gli effetti 3D da fornire a oggetti che dovevano essere rappresentati sulla superficie 2D dello schermo di un computer, aggiungendo anche effetti visivi (come gli effetti di sorgenti luminose che illuminano l'oggetto, le ombre, la consistenza delle superfici, ecc.) per rendere l'oggetto sempre più simile a quello reale. Questi effetti richiedono numerosi calcoli matematici eseguiti con un alto grado di parallelismo in modo da poter indirizzare velocemente e simultaneamente i milioni di pixel che compongono lo schermo di un computer. La matematica necessaria si limita a trasformazioni geometriche come rotazioni, traslazioni, cambio di sistemi di coordinate, quindi tipicamente delle operazioni tra matrici e tra vettori. Questo limitato set di operazioni da eseguire simultaneamente su più dati permette di architettare un processore specializzato in questo compito e in grado di eseguire la stessa istruzione su più dati contemporaneamente. La figura mostra uno schema semplificato di una Gpu, composto da più “*streaming multiprocessor*”, che possono

lavorare in parallelo su *stream* di dati diversi. All'interno dello stesso *streaming multiprocessor* agiscono più unità aritmetiche di processamento che eseguono in parallelo lo stesso set di istruzioni sullo *stream* di dati in ingresso. Le moderne Gpu possono avere fino a un centinaio di *stream multiprocessor*, ognuno contenente decine di processori matematici, e quindi complessivamente possono avere migliaia di core e possono raggiungere le migliaia di miliardi di operazioni al secondo.

Ogni singolo *stream multiprocessor* ha quindi un'architettura Simd (*Single Instruction Multiple Data*) specializzata in calcolo matriciale e questo estende l'utilizzo delle Gpu al di là della *computer graphics* e dei videogiochi. I due supercalcolatori più potenti al mondo (e altri tre tra i primi dieci) utilizzano Gpu come acceleratori matematici. Inoltre, la capacità di manipolare matrici e, in alcuni esemplari, anche tensori, consente a questi processori di essere usati in molteplici applicazioni che vanno dall'algebra lineare alla statistica, ma soprattutto nei sistemi di intelligenza artificiale come *machine learning* e *deep learning*.

## Biografia

**Gaetano Maron** attualmente è direttore del Cnaf, il Centro Nazionale dell'Infn per la ricerca e lo sviluppo nel campo delle tecnologie informatiche. Collabora all'esperimento di fisica delle alte energie Cms dell'acceleratore Lhc al Cern di Ginevra.

## Link sul web

<https://www.scientificamerican.com/article/how-close-are-we-really-to-building-a-quantum-computer/>

DOI: 10.23801/asimmetrie.2019.27.4

# Scavare nei dati

## Il *data mining* nella fisica delle particelle

di Cristiano Bozza

Il modello standard, nonostante i suoi successi, è una rete nella quale siamo intrappolati. Cerchiamo una smagliatura per venirne fuori e avere una visuale più ampia, ma per ora non si è trovata. Con il progredire delle tecniche sperimentali abbiamo accesso a misure di precisione che in passato non erano nemmeno proponibili, ma trovare l'inconsistenza, l'anomalia, il fenomeno imprevisto è tutt'altro che banale. Così come accade nello studio dell'infinitamente piccolo, andando lontano nell'infinitamente grande stiamo assistendo a una esplosione di dati. L'astronomia multimessaggera, che combina telescopi ottici, radiotelescopi, onde gravitazionali e neutrini, dà la caccia a correlazioni tra segnali flebili e nascosti nel rumore dei fenomeni banali. Nessuna strada può essere lasciata non battuta. Non è soltanto un imperativo metodologico, ma anche economico: gli esperimenti costano sia in termini di risorse finanziarie che umane. Sono vere e proprie fabbriche di dati, di una qualità e abbondanza mai viste prima, eppure cerchiamo l'ago nel pagliaio a volte senza sapere nemmeno che aspetto abbia.

Non siamo soli: per quanto la fisica delle particelle elementari abbia avuto un ruolo trainante in passato nell'evoluzione dell'informatica, oggi in molti campi i ricercatori fronteggiano montagne di dati. Si pensi ad esempio alle redditizie attività di profilazione degli utenti, per ottimizzare le campagne pubblicitarie, messe in campo dalle maggiori aziende informatiche, i cui introiti sono per oltre il 90% dovuti alla capacità di analizzare i dati e intraprendere azioni conseguenti. Se è vero che l'informazione e la conoscenza sono il petrolio del XXI secolo, il *data mining* (che potremmo tradurre con "scavare nei dati") è una delle attività da intraprendere tanto per le applicazioni quanto per la ricerca pura.

Come si "scava nei dati"? Evidentemente non si possono seguire solo logiche preordinate. Esistono correlazioni ovvie di cui tenere conto nella costruzione delle basi di dati, come le condizioni operative della strumentazione. Tuttavia, l'analisi dei dati può richiedere decenni dopo la fine della presa dati

a.  
Se l'informazione è il petrolio del XXI secolo, i *big data* degli esperimenti sono le miniere della nuova fisica.



di un esperimento. I database devono rispecchiare questa flessibilità. Nell'approccio relazionale, le connessioni logiche sono precostituite, anche se possono evolvere con il tempo, ma sono di tipo predicativo-binario. Un dato "è in relazione" oppure "non è in relazione" con un altro. Invece gli algoritmi di analisi, così come le esigenze, cambiano nel tempo.

E l'approccio statistico allo studio non consente quasi mai di utilizzare logiche di tipo binario, per cui ogni dato ha un valore se è parte di una popolazione, ma non è individualmente distinguibile dal fondo (si pensi ad esempio agli eventi in cui si è osservato il bosone di Higgs). Per quanto si possa





tentare di restringere e rendere efficiente la ricerca, il grosso del lavoro è un “attacco a forza bruta” agli archivi di dati. Pur non potendosi prescindere dalla catalogazione dei *dataset*, almeno un sottoinsieme deve passare al vaglio di algoritmi di tipo statistico. La disponibilità di elaboratori più veloci e più economici, che consentono di costruire centri di supercalcolo sempre più potenti, non è l'unica strada. Rileggere i dati costa molto sia in termini di tempo che di energia. Ogni nuova sessione di analisi, con nuove selezioni, deve essere preparata con cura per ottimizzare i risultati. Non è affatto detto che i ricercatori riescano a “sistemare tutte le trappole” per catturare fenomeni elusivi, attraverso quello che tradizionalmente viene chiamato “*analytical processing*”. Le tecniche di intelligenza artificiale, anzi, offrono numerosi vantaggi da questo punto di vista, poiché lo sperimentatore, almeno in teoria, deve “solo” preparare una simulazione, definire la tipologia di algoritmo e confrontare i risultati dell'addestramento su dati simulati con i dati reali. Tuttavia, un “albero decisionale” (Bdt) non prenderà in considerazione quantità che non siano state identificate in

precedenza da chi dirige l'analisi. Similmente accade per reti neurali a pochi livelli (in uso da tempo nelle applicazioni più semplici). Significativi vantaggi possono venire dall'uso di tecniche più aggressive come le reti neurali convoluzionali (Cnn) che trattano i dati grezzi, opportunamente addestrate e rese robuste con “reti avversarie” (Gan) che le abitua a non farsi ingannare da segnali irrilevanti, rumore, disturbi, eventi occasionali (vd. p. 15, ndr). Rispetto alle tecniche di analisi tradizionali, l'intelligenza artificiale in genere fornisce maggiore velocità di elaborazione, perché in definitiva funzioni complicate dei dati sono approssimate da “mattoni” semplici. Si sposta l'onere della scoperta dei criteri ottimali di classificazione dal ricercatore alla macchina. Tuttavia, si introducono due problemi nuovi, ossia quello della preparazione dei dati di addestramento (“*training set*”), che devono essere rappresentativi del fenomeno cercato, e quello dell'interpretazione statistica dei risultati. In altri termini si deve riguadagnare il controllo sulla significatività delle decisioni prese dalle macchine. Quanto possiamo fidarci di un classificatore automatico?

b.  
Il centro di calcolo del Cern: centinaia di migliaia di processori negli armadi refrigerati ricercano instancabilmente preziosi indizi che possano consentire alla scienza di aprire nuove prospettive.

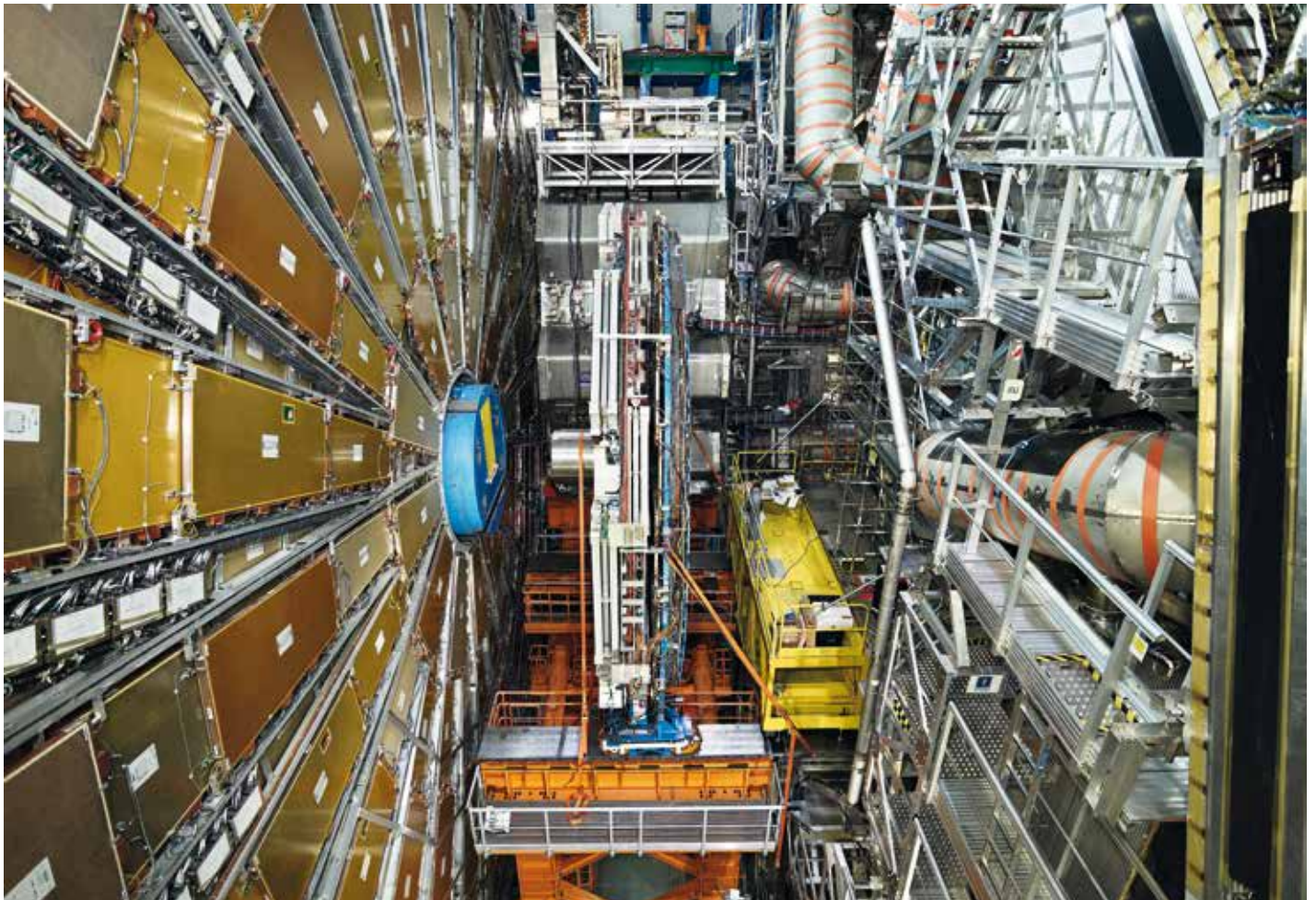


Quanto spesso si sbaglia, e di quanto? Lo stiamo usando in “*overfitting*”, cioè ricorda perfettamente l’addestramento ma non è riuscito ad astrarre quantità rilevanti? La maggior parte di queste tecniche nasce al di fuori della fisica delle alte energie, in ambiti originali di applicazione, ad esempio nel commercio, nelle scienze sociali o nella sicurezza informatica, nei quali è possibile che vi sia una reale volontà di ingannare l’algoritmo. Si affrontano e si risolvono problematiche che nella fisica non si pongono: noi possiamo credere all’ “onestà” dei nostri dati, altri no. D’altra parte possono esservi cause ignote o sottovalutate di malfunzionamento di un apparato o semplici correlazioni inesplorate che finiscono per mimare segnali statisticamente rilevanti. Questo porta un cambiamento rilevante nel profilo culturale del fisico: è relativamente meno importante la capacità di costruire programmi molto ottimizzati dal punto di vista delle prestazioni, perché strumenti aperti come TensorFlow, Keras o altri, sviluppati

e supportati da grandi multinazionali, consentono di avere accesso allo stato dell’arte dal punto di vista tecnologico. Invece, è progressivamente più importante che il fisico si addentri in aspetti matematici, statistici e più astratti della teoria dell’informazione. Immagazzinare dati, rileggerli e rielaborarli può avere un impatto notevole in termini di tempo (sia umano che macchina) e di costo della memoria di massa. Una delle più forti limitazioni negli algoritmi di selezione degli eventi sperimentali (il cosiddetto *trigger* online, vd. in Asimmetrie n. 17 p. 33, ndr) viene dalla necessità di “tagliare” il più possibile, sacrificando a volte intere linee di ricerca. Tuttavia, la potenza di calcolo oggi disponibile è tale che è possibile far girare durante l’acquisizione parti della ricostruzione degli eventi che normalmente si eseguono in fase di analisi, o versioni accelerate di esse. Si sfuma così la distinzione tra online e offline, tra presa dati e analisi.

Nell’esperimento Cms del Cern questo è stato chiamato “*data scouting*”, ossia

c.  
La spettacolare macchina di Cms, in preparazione per una nuova presa dati. La complessità della strumentazione è seconda solo a quella delle informazioni prodotte.





d.  
Quattro Detection Unit di Km3net pronte per essere deposte sul fondo del Mediterraneo. Negli abissi marini, dove migliaia di metri d'acqua schermano i raggi cosmici, cercheranno neutrini provenienti dal cosmo o dall'alta atmosfera.

andarsi a cercare già durante la presa dati ciò che può essere interessante prima che sia perso per sempre. Anche altri due esperimenti del Cern, Atlas e Lhcb, hanno implementato approcci simili. A ciò si affianca il “*data parking*”, ossia il salvataggio temporaneo di un sottoinsieme di dati grezzi, necessari in caso di scoperta (e che diversamente potranno essere cancellati). Queste politiche di gestione attiva dei dati, utilizzate nel *run 2* di Lhc dal 2015 al 2018 per la ricerca di fisica esotica (come lo *Z'*), rendono percorribili canali di indagine che con i metodi “tradizionali” (“salva subito e analizza in seguito”) richiederebbero costi enormi, a fronte di esigue speranze di scoperta. D'altra parte,

introducono fattori di rischio, poiché errori di valutazione nelle selezioni o nell'implementazione portano a una perdita secca di dati, e pertanto possono essere usate solo per incrementare il campione, senza potersi sostituire ai trigger convenzionali. L'ottimizzazione della presa dati è da tempo una realtà nel campo dell'astronomia osservativa (vd. p. 31, ndr), e in tempi di astronomia multimessaggera, che richiede l'interazione di telescopi ottici, radio, gravitazionali e per neutrini, operativi in diversi luoghi della Terra, la variazione dinamica dei trigger in conseguenza di allarmi globali è ormai prevista già in fase di sviluppo dei sistemi di acquisizione, come nel caso di Km3net che dev'essere

pronto a segnali di supernova e altri transitori. La collaborazione Ska, che costruirà e gestirà la più grande rete di radiotelescopi del mondo, ha estremizzato i concetti di acquisizione adattiva e di elaborazione “al volo”: con più di 10 exabyte di dati raccolti al giorno, di cui “solo” 1 petabyte al giorno sarà immagazzinato, la riduzione dei dati deve avvenire più vicino possibile all'antenna, possibilmente utilizzando Cpu economiche ma relativamente potenti e non “avide” di energia. I partner commerciali non perdono l'occasione di contribuire a questi sviluppi: come accaduto spesso in passato, la ricerca pura sta già incubando le tecnologie che cambieranno il nostro modo di vivere domani.

#### Biografia

**Cristiano Bozza**, Ph.D., si occupa dal 1996 di acquisizione, ricostruzione dati e database, prima in Chorus e poi in Opera, esperimento che ha scoperto l'oscillazione dei neutrini muonici in neutrini tau. Attualmente lavora in Km3net, sui neutrini da astrosorgenti e atmosferici. Nel progetto europeo Asterics, per cui è stato responsabile nazionale Infn, ha contribuito a simulazioni, supercalcolo e applicazioni di intelligenza artificiale, attività che continua nel nuovo progetto Escape.

#### Link sul web

<https://www.km3net.org>

<https://www.skatelescope.org/>

DOI: 10.23801/asimmetrie.2019.27.5



# Il casinò della fisica

## Tecniche di simulazione per l'analisi dei dati

di Luciano Pandola e Giada Petringa

La nascita di ogni esperimento di fisica moderna è preceduta da un periodo di progettazione, nel quale si determina l'equilibrio ottimale fra la sensibilità, le prestazioni e i costi. Gli esperimenti sono solitamente concepiti per cercare un tipo specifico di “segnale” (il bosone di Higgs, la materia oscura ecc.). Il rivelatore utilizzato deve perciò essere sensibile al segnale di interesse, ma allo stesso tempo ben protetto da altri tipi di eventi (il cosiddetto “fondo”) che potrebbero imitarlo. Data la complessità degli attuali esperimenti, gran parte di questo lavoro di progettazione deve essere fatto simulando con opportuni modelli la risposta del sistema nel suo complesso, sia a eventi di segnale che di fondo. È possibile in questo modo definire, per esempio, il tipo e le dimensioni dei rivelatori da utilizzare – per riconoscere meglio il segnale, oppure le schermature necessarie a sopprimere il fondo a un livello tollerabile. La simulazione è importante anche dopo che l'esperimento è stato costruito, quando si affronta l'analisi dei dati. Le simulazioni offrono infatti l'enorme vantaggio di poter “prevedere” la risposta del rivelatore al variare delle caratteristiche del segnale e

del fondo. Il confronto fra i dati reali e le simulazioni è un passaggio spesso obbligato per analizzare e interpretare i risultati di un esperimento.

Le simulazioni degli esperimenti sono un esempio particolare di problema che può essere affrontato mediante il cosiddetto “metodo Monte Carlo”. Si tratta di un approccio di tipo statistico, il cui ingrediente primario sono i numeri casuali. Il nome è stato coniato proprio traendo ispirazione dal famoso casinò di Monte Carlo, regno indiscusso della casualità. Il campo di applicazione del metodo Monte Carlo è molto vasto e include sistemi complessi, il cui comportamento è determinato dall'influenza reciproca di un gran numero di componenti elementari. Ne sono un esempio gli studi sul flusso del traffico veicolare, l'andamento del mercato finanziario o, per l'appunto, l'interazione di un fascio di particelle in un rivelatore. In pratica si utilizzano numeri casuali per simulare il comportamento delle unità elementari, così da predire il comportamento globale del sistema. Ripetendo molte volte l'intera procedura si può quindi determinare la risposta media del sistema. Il metodo Monte Carlo presuppone



a.  
La roulette del casinò di Monte Carlo, da cui trae il nome il “metodo Monte Carlo”, a rimarcare così il ruolo centrale ricoperto dai numeri casuali.



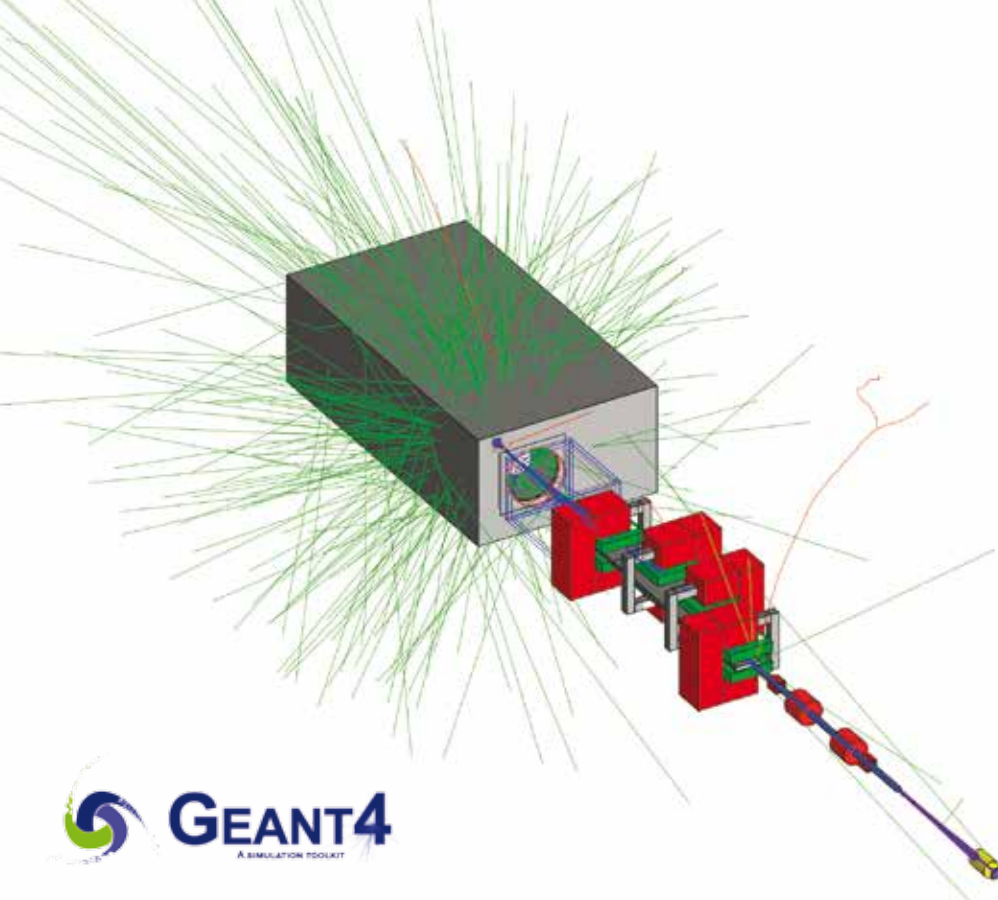
b.  
John von Neumann (1903-1957), nato in Ungheria ed emigrato negli Stati Uniti negli anni '30, è stato l'indiscusso padre fondatore, insieme a Stanislaw Ulam, del moderno utilizzo scientifico del metodo Monte Carlo.

sia la generazione di numeri casuali sia l'iterazione di una stessa procedura: si coniuga perciò perfettamente con le caratteristiche dei computer, che sono progettati per eseguire velocemente calcoli ripetitivi.

Tuttavia, le idee alla base del metodo sono dovute al conte Georges-Louis Leclerc de Buffon (1777) e a Pierre Simon Laplace (1786), in anticipo di secoli rispetto all'invenzione dei computer. Le applicazioni ebbero poi un balzo in avanti a partire dal secondo dopoguerra, grazie a Stanislaw Ulam e John Von Neumann, che sfruttarono i primi proto-computer legati allo sviluppo delle prime armi termonucleari.

L'applicazione più comune del metodo Monte Carlo nell'ambito della fisica particellare consiste nella simulazione del passaggio di radiazione attraverso la materia, per poter prevedere la risposta dei rivelatori. Negli ultimi decenni sono stati sviluppati diversi software adatti allo scopo: partendo da quelli pionieristici, limitati a un ambito specifico, si è lentamente arrivati a progetti più ampi, concepiti per rispondere a molte esigenze differenti. I programmi ad oggi disponibili differiscono soprattutto per i modelli di fisica adottati per simulare le singole interazioni. Uno dei presupposti del metodo Monte Carlo è infatti quello di saper descrivere il comportamento di ciascuna delle unità elementari che compongono il sistema, ossia, nel caso specifico, le interazioni di una data particella con il mezzo attraversato. "Regole di comportamento" leggermente diverse di queste unità possono portare a differenze più o meno visibili nei risultati per l'intero sistema. Talvolta, le "regole di comportamento" delle particelle elementari non

sono note con precisione e si è costretti a ricorrere all'uso di modelli fenomenologici o teorici, che impongono delle approssimazioni o semplificazioni: questo è il caso soprattutto per gli adroni, ossia per quelle particelle formate da quark e soggette all'interazione di tipo nucleare forte. In alcuni casi, i modelli non sono descritti mediante formulazioni analitiche, ma si basano su dati sperimentali, che vengono riportati in forma tabulata. Questo approccio è inevitabilmente limitato alle sole applicazioni per le quali sono disponibili dati precisi e affidabili. La simulazione si sviluppa "tracciando" particelle nel mezzo di interesse (per esempio, un apparato sperimentale): si parte da una particella primaria, che ha una certa energia e una certa direzione, e se ne segue l'evoluzione all'interno dei vari elementi che compongono il sistema. La particella è soggetta a interazioni successive con il mezzo, ciascuna descritta da ben precise "regole di comportamento". Ognuna di queste interazioni produce una variazione sia dell'energia che della traiettoria e può causare l'emissione di nuove particelle, che vengono dette secondarie. La procedura di tracciamento continua in modo iterativo, finché la particella primaria e tutte le eventuali secondarie vengono assorbite oppure fuoriescono dai confini del volume simulato (il cosiddetto "volume mondo"). Poiché il metodo Monte Carlo è fondato su un approccio statistico, è necessario che la procedura venga ripetuta molte volte: si traccia cioè un gran numero di particelle primarie (insieme a tutte le eventuali secondarie), in modo da determinare l'evoluzione media di tutto il sistema. La precisione della simulazione è tanto migliore, quanto



c.  
Immagine della simulazione Geant4 della linea di trasporto Elimed (Eli-Beamlines Medical and multidisciplinary application) installata a Eli-Beamlines (Praga, vd. in *Asimmetrie* n. 12 p. 34, ndr). La linea serve al trasporto dei fasci di particelle cariche prodotte dall'interazione di un laser di alta potenza su un bersaglio sottile. Il progetto è stato verificato e ottimizzato anche mediante simulazioni Monte Carlo.

più è grande il numero di particelle primarie tracciate. La natura intrinsecamente probabilistica della fisica quantistica fa sì che ogni ripetizione avrà un'evoluzione diversa, pur partendo sempre dalle stesse condizioni iniziali: le interazioni delle particelle con il mezzo avverranno generalmente in posizioni differenti e producendo effetti differenti. Così si può determinare, ad esempio, la risposta dei moderni complessi rivelatori di Lhc, ma anche la propagazione dei raggi cosmici nella Galassia.

Uno dei programmi Monte Carlo che negli ultimi anni si è affermato nella comunità dei fisici per la simulazione Monte Carlo dell'interazione fra radiazione e materia si chiama Geant4. Si tratta di un software *open source*, sviluppato da una collaborazione internazionale che include l'Infn. La prima versione è stata rilasciata nel 1998, ma Geant4 continua tuttora a essere aggiornato con regolarità al fine di migliorarne sempre più la precisione, l'affidabilità e le prestazioni. Geant4 è una sorta di "cassetta degli attrezzi" (un "*toolkit*") che contiene tutti gli strumenti necessari per simulare un rivelatore reale: gestione dei volumi e dei materiali, modelli di fisica, generazione e tracciamento di particelle, recupero delle informazioni da salvare. Agli utenti spetta il compito di scegliere dalla "cassetta" gli strumenti più adatti al proprio scopo specifico e poi utilizzarli. La precisione e la flessibilità di Geant4 lo rendono un prodotto affermato in diversi ambiti: dalla fisica delle alte energie alla fisica nucleare, fisica spaziale e astrofisica particellare.

Negli ultimi anni, poi, si è sempre più diffusa l'applicazione del metodo Monte Carlo nel campo della fisica medica. Le moderne tecniche diagnostiche, la radioterapia, l'*imaging* in medicina nucleare (vd. p. 39, ndr), così come la descrizione del danno biologico, si sono avvalse con successo di software nati originariamente per simulare esperimenti di fisica. Ciò ha permesso di risolvere problemi complessi che erano prima

affidati a calcoli analitici, spesso soggetti a errori o notevoli approssimazioni. Tra questi c'è il calcolo della dose, quantità fondamentale in radioterapia, che racchiude sia gli aspetti fisici legati alle caratteristiche delle radiazioni adottate (energia, tipo di particella) sia le caratteristiche biologiche inerenti la risposta di uno specifico tessuto alla radiazione. Esiste anche un'intera branca di Geant4 dedicata alla simulazione del danno biologico indotto da radiazioni ionizzanti, chiamata Geant4-Dna. Questa estensione del software è continuamente aggiornata da un gruppo di ricercatori che comprende fisici, informatici e biologi, così da poter migliorare sempre più i piani di trattamento dei pazienti soggetti alla radioterapia.

#### Biografia

**Luciano Pandola** è un ricercatore Infn presso i Laboratori Nazionali del Sud di Catania. Si è occupato di simulazioni Monte Carlo e di analisi dati nell'ambito di esperimenti di astrofisica particellare e di ricerca di eventi rari. È membro della collaborazione Geant4 dal 2002.

**Giada Petringa** è un'assegnista di ricerca presso i Laboratori Nazionali del Sud di Catania. Membro della collaborazione Geant4 si occupa di simulazioni Monte Carlo applicate alla fisica medica e alla modellistica radiobiologica.

#### Link sul web

<http://geant4.org>

<http://www.mathisinthair.org/wp/2015/10/i-metodi-monte-carlo-prima-parte/>

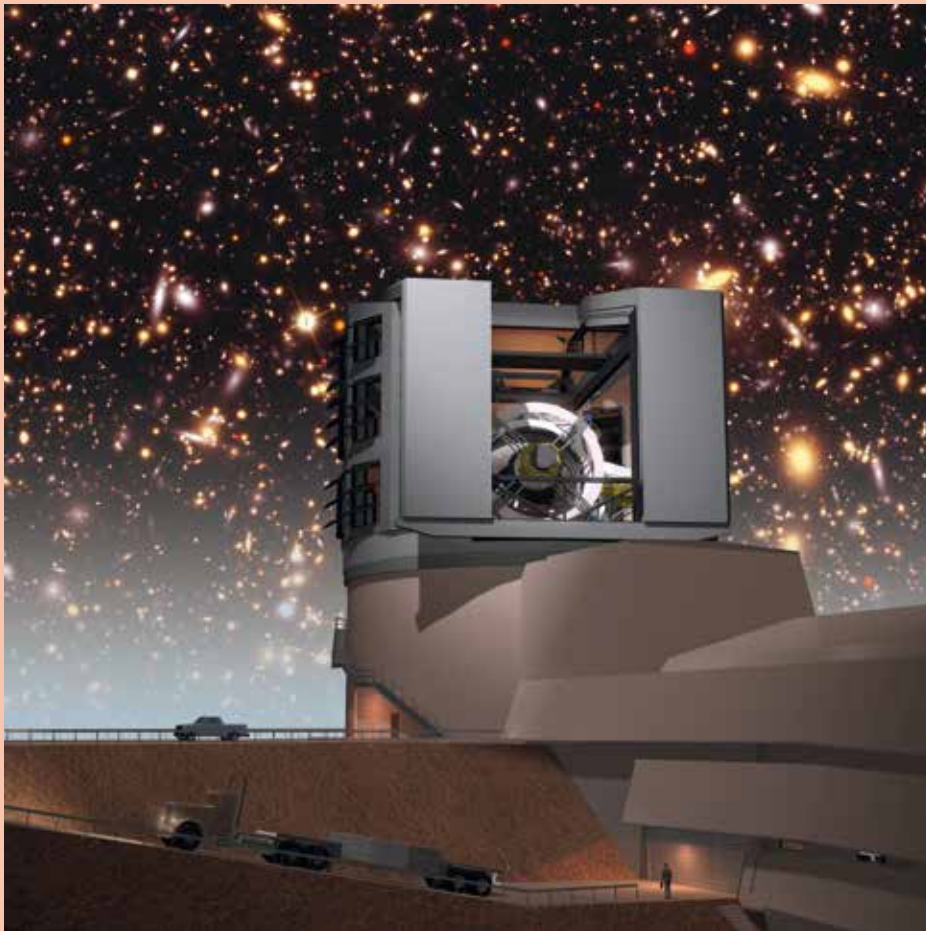
DOI: 10.23801/asimmetrie.2019.27.6



# Nello zoo galattico

## L'apprendimento profondo dell'universo

di Viviana Acquaviva



a.  
Il telescopio Lsst, attivo a partire dal 2023, produrrà immagini per miliardi di galassie, con un flusso di dati di decine di terabyte per notte.

I termini “intelligenza artificiale” e “*machine learning*” sono divenuti ormai ubiqui in qualunque campo scientifico, dalla biologia alla chimica alla medicina, passando, ovviamente, per la fisica. Ma è difficile immaginare una disciplina che abbia subito una rivoluzione metodologica così rapida e sostanziale come l'astrofisica. Le ragioni di questo cambiamento sono in parte da ricercarsi nella natura stessa dell'indagine sull'universo. I

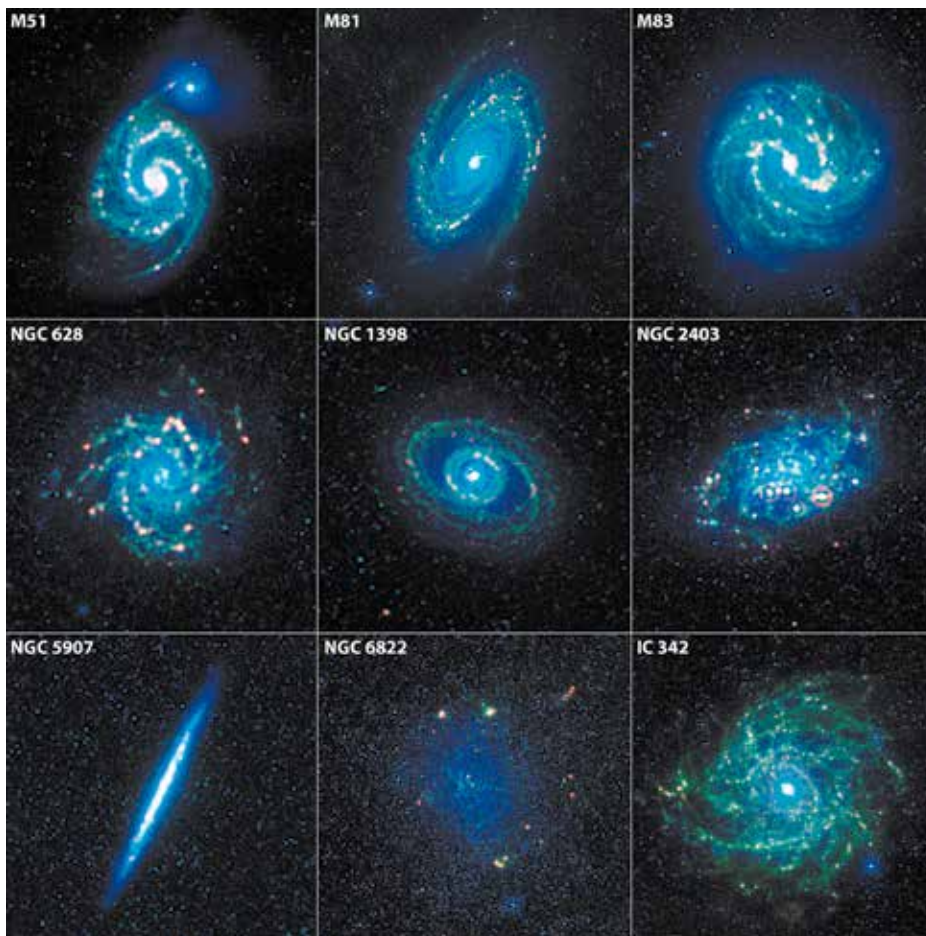
metodi dell'intelligenza artificiale, dal *machine learning* al *data mining*, sono ideali per descrivere e investigare fenomeni i cui modelli fisici non sono sufficientemente sofisticati per catturarne la complessità. Tipicamente, questo accade quando abbiamo a che fare con oggetti estremamente complicati, che sono descritti da un numero di variabili molto elevato, o quando la nostra comprensione e interpretazione dei fenomeni fisici è limitata. E questa

descrizione calza a pennello a tantissimi problemi in ambito astrofisico, dalla formazione delle galassie alle dinamiche stellari, dalle simulazioni cosmologiche che mostrano la distribuzione della materia oscura nell'universo, all'individuazione di segnali variabili in un flusso di dati. Considerando anche il fatto che missioni come il Large Synoptic Survey Telescope (Lsst) genereranno nel giro di pochi anni una quantità di dati centuplicata rispetto agli osservatori delle generazioni precedenti, rendendo alcuni dei metodi correnti, come ad esempio l'ispezione manuale, completamente obsoleti, non c'è da stupirsi che l'astrofisica sia uno dei settori che ha adottato le nuove tecniche in modo rapido e probabilmente irreversibile. Ma quali sono le applicazioni del *machine learning* che hanno avuto più successo in astrofisica e cosmologia? La lista è certamente troppo lunga per essere completa. Un grandissimo punto di forza di queste tecniche, e specialmente del *deep learning*, è la loro efficacia nell'analizzare dati multidimensionali, come ad esempio le immagini, che contengono informazioni spesso difficili da descrivere in modo

analitico. Per esempio, identificare una galassia a spirale, vista sotto diverse possibili angolazioni, è un compito relativamente semplice per l'occhio (e il cervello) umano, ma scrivere una formula matematica - dunque, un modello - che possa "cercare" questa galassia nell'immagine e tenere conto di tutte le possibili prospettive è estremamente complicato. Qui le reti neurali vengono in nostro soccorso, perché la loro architettura permette facilmente di catturare correlazioni spaziali in due o tre dimensioni, purché il sistema abbia a disposizione una serie di esempi rappresentativi da cui imparare, un po' come un bambino impara facilmente a riconoscere gli animali dopo averli visti alcune volte in fotografia. Non a caso, una delle prime applicazioni del *machine learning* ai dati astronomici è stata quella di classificare le galassie in base alla loro morfologia, in collaborazione con il team di Galaxy Zoo, in cui il pubblico è stato

chiamato in causa per formare, sulla base del riconoscimento "umano", una collezione di elementi rappresentativi che potessero guidare l'algoritmo. Più in generale, l'abilità delle reti neurali nell'analizzare immagini permette al *machine learning* di migliorare la qualità dell'analisi dati rispetto ai metodi tradizionali, per cui le medesime immagini di stelle e galassie appaiono più nitide e ricche di informazioni se analizzate con i nuovi metodi. Inoltre, permette di catturare informazioni "nascoste" in strutture complicate, come i dettagli delle proprietà statistiche delle mappe della radiazione del fondo cosmico a microonde o del "*lensing* gravitazionale debole" (cioè la distorsione subita dalle immagini delle galassie lontane dovuta alla deflessione della luce da parte della materia incontrata nel percorso tra la sorgente e la Terra). Questo apre anche una nuova promettente direzione nel campo delle simulazioni cosmologiche, che

in generale richiedono l'uso di risorse computazionali molto elevate. Alcuni tipi di reti neurali, come le Gan, nate solo nel 2014 come algoritmi per riconoscere i falsi, sono state usate con successo per generare nuovi dati e di fatto ampliare il volume delle simulazioni, con un costo computazionale assai ridotto. Perfezionare questi metodi consentirà, in futuro, di ottenere simulazioni cosmologiche ad alta risoluzione su grandissimi volumi, che oggi risultano proibitive per via degli eccessivi tempi richiesti. Le Rnn si prestano a un'analisi molto efficace di dati in sequenza temporale (*time series*), che sono fondamentali per l'individuazione tempestiva di sorgenti variabili, come supernovae o raggi cosmici, e saranno certamente fondamentali sia nell'analisi dei dati di Lsst che nel campo nascente dell'astronomia multimessaggera, l'analisi combinata di segnali di origine elettromagnetica e gravitazionale. Ma non ci sono solo le reti neurali. Da un lato, l'astrofisica non è fatta solo di immagini. Dall'altro, tutti i metodi di *deep learning* tendono a sacrificare l'interpretabilità in favore della precisione. In altre parole, è difficile capire perché una rete neurale prende una certa decisione e - di conseguenza - imparare nuova fisica. Per questo motivo, mentre lo sviluppo di strumenti che possano decodificare l'*output* delle reti neurali diventa un argomento



b.  
Forme e colori delle galassie nell'universo svelano la storia dell'evoluzione cosmica: le tecniche di *machine learning* consentono di classificarle correttamente in tempi molto rapidi.



c.  
La figura del *support scientist*, a metà tra il ricercatore e il programmatore, diventerà sempre più cruciale nelle scienze ad alto contenuto di dati, come l'astrofisica.

sempre più discusso, altri metodi più semplici e trasparenti hanno potenti applicazioni in campi in cui i modelli lasciano a desiderare. Un buon esempio è il caso del calcolo delle distanze delle galassie basato sulla fotometria (*photometric redshifts*). In questo caso, metodi di *machine learning* come Random Forest (una particolare combinazione di Bdt) o Principal Component Analysis Decomposition, o anche soltanto combinazioni lineari di spettri empirici, risultano spesso più precisi dei metodi basati sui modelli dell'emissione energetica delle galassie.

Per continuare, non possiamo non citare l'importanza di tecniche di *machine learning* nell'alleviare il problema sempre più pressante della rappresentazione e compressione delle informazioni contenute nei dati. Il volume dei dati prodotti dalle nuove generazioni di telescopi (per ritornare al caso di Lsst, dell'ordine di decine di terabytes per notte) è impossibile da maneggiare per un utente che lavora sul suo computer personale, e ovviamente la situazione peggiora quando si considerano i prodotti derivati come cataloghi, distribuzioni di probabilità di parametri, e così via. La splendida e doverosa pratica della condivisione pubblica dei dati, originali e derivati, tipica dell'astrofisica e pilastro della riproducibilità dei risultati richiesta dal metodo scientifico, rischia di arenarsi per un problema tecnico di accesso ai dati stessi. Mentre alcuni di questi problemi potranno essere risolti solo con l'allocatione di nuove risorse, il *machine learning* può aiutare tramite l'uso di tecniche specializzate, note come "riduzione della dimensionalità", che consentono di individuare le proiezioni più interessanti nello spazio dei dati e ridurne di fatto la dimensione, con un sacrificio minimo del contenuto di informazione.

Da ultimo, è bene ricordare che la rivoluzione creata

dal diffondersi dei metodi del *machine learning* ha delle implicazioni profonde a livello umano. Cambia la figura del ricercatore, che diventa forse meno specializzato e deve saper parlare il linguaggio della programmazione, della visualizzazione dei dati e diventare, ora più che mai, un "narratore" di storie complesse e affascinanti. Aumentano le possibilità di ricerca interdisciplinare, in cui la *data science* fornisce quel linguaggio comune che forse finora è un po' mancato a statistici, informatici e astronomi. Poi si aprono nuove possibilità di carriera, a cavallo tra l'astrofisica e l'informatica, come dimostrato dalle nuove figure professionali come il "*support scientist*", sempre più comuni negli annunci di lavoro. E più di ogni altra cosa, è importante tener presente questi cambiamenti nella formazione dei fisici di domani, così che possano sfruttare al meglio questa eccezionale occasione di analizzare dati in maniere finora impossibili o sconosciute, e scoprire nuove verità sull'universo che ci circonda.

#### Biografia

**Viviana Acquaviva**, astrofisica, utilizza i metodi dell'intelligenza artificiale nelle sue attività di ricerca sull'evoluzione dell'universo. È professore presso la City University of New York, membro del Centro di Astrofisica Computazionale del Flatiron Institute, *visiting scientist* presso il Museo di Storia Naturale di New York e ambasciatrice (Harlow Shapley Lecturer) per la American Astronomical Society. Nel 2018 è stata premiata come una delle 50 donne italiane più influenti nel campo della tecnologia da "InspiringFifty".

DOI: 10.23801/asimmetrie.2019.27.7

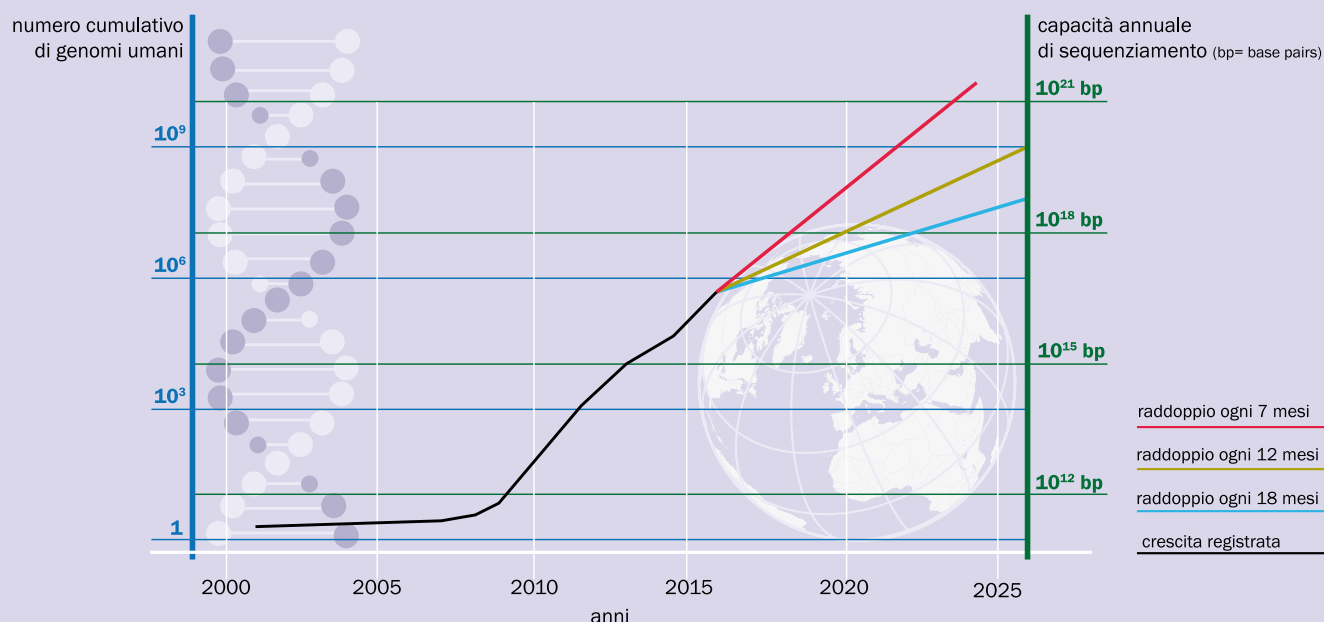


# Le reti della vita

## La bioinformatica per una medicina di precisione

di Michele Caselle

### CRESCITA DEL SEQUENZIAMENTO DEL DNA



La biologia molecolare sta attraversando in questi ultimi anni una vera e propria rivoluzione scientifica. Si tratta di un processo per certi versi simile a quello sperimentato in fisica negli anni '20 del secolo scorso quando nacque la meccanica quantistica. Grazie alle nuove tecnologie di sequenziamento e di manipolazione del genoma i biologi hanno oggi accesso a una mole di dati e informazioni inimmaginabile fino a qualche anno fa (vd. fig. a) e questo ha portato a un radicale cambiamento di paradigma, alla nascita di nuove idee e anche di nuove discipline scientifiche. Per cercare di controllare questo diluvio di dati ed estrarne tutte le informazioni nascoste i ricercatori hanno dovuto adattare alla biologia molecolare le idee e gli strumenti tipici del *data mining*. Termini come bioinformatica, biologia computazionale o *systems biology*, che erano sconosciuti vent'anni fa, sono diventati ormai la norma nei laboratori di tutto il mondo. Si tratta di discipline giovani,

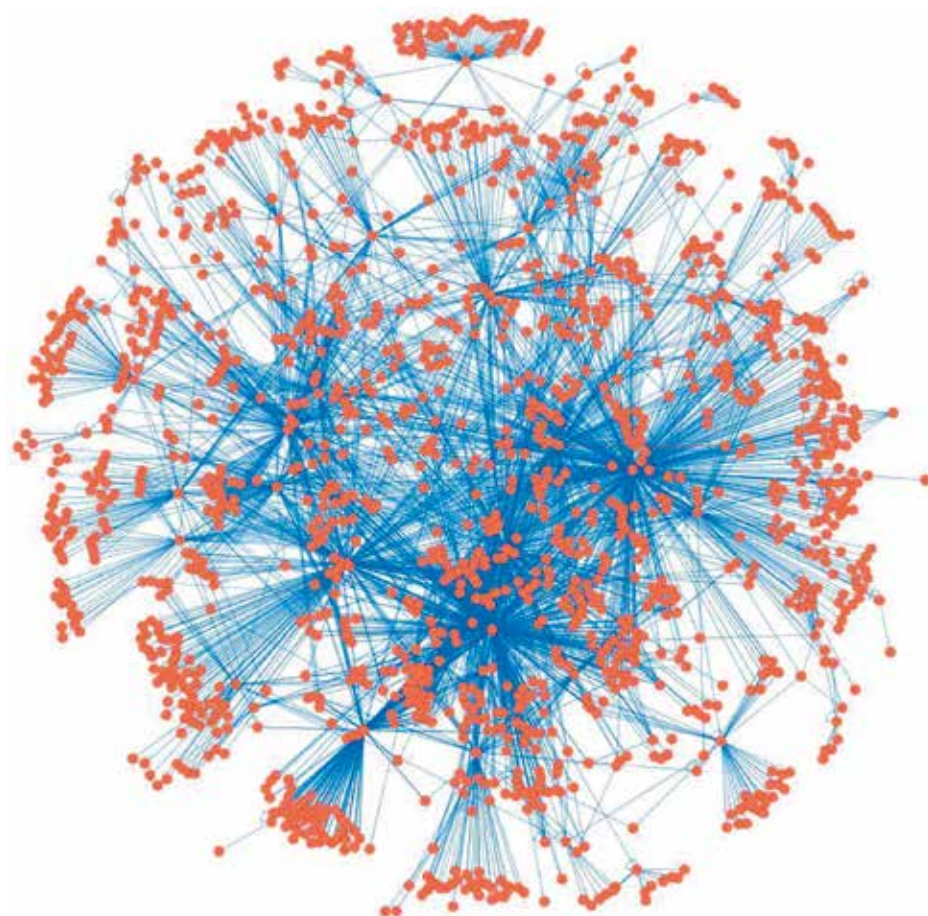
fortemente interdisciplinari, dove fisici, matematici, esperti di informatica, biologi e medici lavorano assieme per arrivare a una comprensione sempre più profonda di come funzionano le cellule e gli organismi viventi. Alla base di questo progresso sta il sequenziamento del genoma umano ottenuto all'inizio degli anni 2000. Un risultato importante non solo per le sue ricadute scientifiche, ma soprattutto per il suo valore simbolico. Con il termine "genoma" si intende la sequenza di Dna che in ogni organismo vivente codifica l'informazione genetica. Il Dna è una lunga molecola formata dalla combinazione di quattro elementi fondamentali detti "nucleotidi" o "basi": adenina, citosina, guanina e timina, solitamente indicati con l'abbreviazione A, C, G, T. La lunghezza del genoma può variare molto da specie a specie. Nell'uomo la sequenza è composta da circa 3 miliardi di basi mentre per gli organismi più semplici che esistono in natura, i batteri, può scendere fino a 500.000 basi. Nel

a. Nella figura è riportata la crescita della quantità di informazioni disponibili per i ricercatori (misurata in numero di genomi umani) al passare del tempo. Come si può vedere stiamo assistendo a una crescita esponenziale (l'asse delle ordinate è in scala logaritmica, in questa scala una crescita esponenziale corrisponde a una retta), ma la cosa impressionante è la rapidità di questa crescita. La cosiddetta "legge di Moore", che ha governato la crescita della potenza di calcolo dei computer negli ultimi anni, prevede un raddoppio ogni 18 mesi (la retta blu). In ambito genomico abbiamo assistito negli ultimi 10 anni a un raddoppio ogni 7 mesi! Estrapolando questa curva (la retta in rosso nella figura) arriveremo tra un paio di anni ad avere una quantità di dati a disposizione dei ricercatori dell'ordine dello yottabyte ( $10^{24}$  byte, un milione di miliardi di miliardi di byte), paragonabile all'intera capacità di memoria disponibile sul pianeta oggi.

Dna è codificata l'intera informazione necessaria per formare l'individuo adulto e, infatti, tipicamente la lunghezza del Dna cresce con la complessità dell'organismo. La molecola di Dna è formata da una doppia elica, in cui le basi sono sempre appaiate nella combinazione (A,T) e (C,G), e proprio questa proprietà è alla base delle due caratteristiche fondamentali del Dna: la capacità di duplicarsi in modo fedele (e quindi trasferire alla progenie l'informazione genetica) e la possibilità di essere letta dalla cellula che (mediante una serie di complessi processi biochimici indicati complessivamente con il termine di "espressione genica") usa l'informazione codificata nel Dna per produrre le proteine che sono i costituenti elementari della cellula. Possiamo pensare alle proteine come ai "mattoni" con cui sono costruite tutte le componenti cellulari e che permettono alle cellule di organizzarsi in tessuti e in un organismo completo. Ogni proteina è codificata da un'unità del genoma, cioè

da una porzione di una sequenza del Dna chiamata "gene". Dalla conclusione del Progetto Genoma nel 2000 siamo in grado di leggere l'intera sequenza del Dna dell'uomo e negli anni successivi lo stesso risultato si è ottenuto per tantissimi altri organismi viventi. Negli ultimi anni grazie al miglioramento delle tecnologie siamo riusciti anche a riconoscere le differenze nel genoma da individuo a individuo, e in prospettiva sarà possibile per ogni singolo individuo conoscere la propria sequenza di Dna. Ma leggere non significa capire! Siamo in grado di identificare tutte le proteine codificate nel genoma umano, ma siamo ancora molto lontani dal capire le loro funzioni. Alcune cose cominciamo però a intuirle. Innanzitutto, è chiaro che per capire un sistema complesso come una cellula vivente è necessario un cambiamento di paradigma, in cui al centro dello studio deve essere posto non il singolo gene o la singola molecola, ma l'intero sistema, e soprattutto la rete di interazioni tra

i vari geni. Questo cambiamento di paradigma è il primo grande regalo dei *big data* alla biologia molecolare. È ormai chiaro che è solo a livello di rete che si possono capire tutte le complesse funzioni di un organismo vivente e anche della sua unità base, la cellula. L'altra cosa ormai chiara è che le proteine con cui sono costruiti tutti gli organismi viventi sono più o meno sempre le stesse e che il loro numero non cresce in modo significativo con la complessità dell'organismo. Il genoma umano è composto da circa 20.000 geni, più o meno lo stesso numero del topo o del moscerino della frutta. Quello che cambia al crescere della complessità è il "libretto di istruzioni", l'insieme di regole che decide con quali "mattoni" costruire una particolare cellula e in che ordine disporli. Queste istruzioni sono codificate in una rete detta "rete di regolazione" (un esempio è riportato in fig. b), le cui proprietà per gli organismi complessi sono ancora largamente inesplorate. Ricostruire questa rete di regolazione dai dati sperimentali è una delle sfide più impegnative e affascinanti della biologia computazionale moderna. Esistono essenzialmente due approcci. Il primo si basa sul fatto che i principali attori di questo sistema di regolazione (detti "fattori di trascrizione") sono in grado di riconoscere certe sequenze di Dna (dette "siti di legame") e di legarsi a queste sequenze in modo da regolare l'accensione e lo spegnimento dei geni vicini. Il problema è che identificare queste sequenze all'interno del genoma è come trovare un ago in un pagliaio. Questi siti di legame sono abbastanza corti (tipicamente da 5 a 20 nucleotidi) e possono essere degeneri (cioè accettare in alcune posizioni più basi diverse), mentre il genoma, ad esempio nel caso dell'uomo, contiene più di 3 miliardi



b.  
La rete di regolazione dell'*Escherichia coli*, un piccolo batterio che vive abitualmente nel nostro intestino e che è uno dei più studiati organismi modello in biologia molecolare. Nella figura ogni nodo rappresenta una proteina e le interazioni con cui una proteina regola l'espressione di un'altra proteina sono indicate da linee blu.

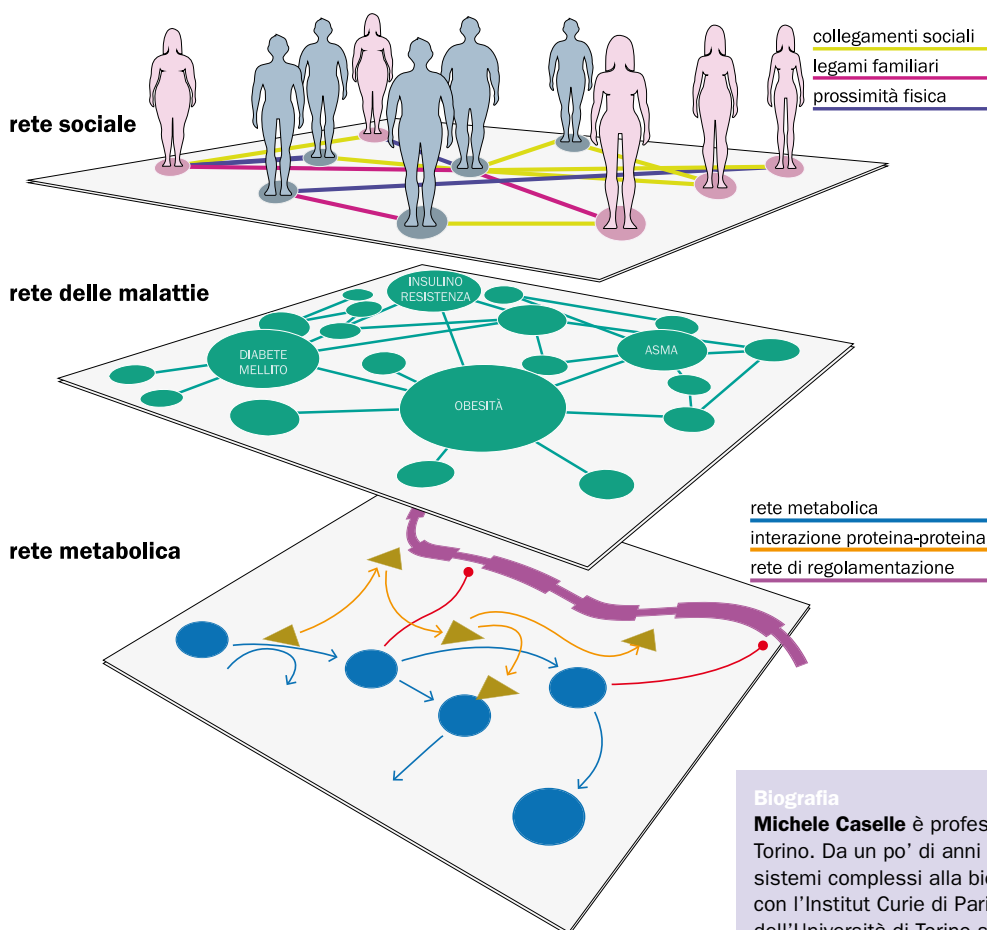
di basi. Esistono metodi estremamente ingegnosi per ridurre la complessità del problema. Ci sono tecniche sperimentali (dai nomi esotici, come Chip-seq, Dnase, Footprinting) che permettono di restringere la porzione di genoma in cui cercare, e si possono combinare queste tecniche con uno studio di conservazione evolutiva, basato sull'idea che se una porzione di genoma è particolarmente importante, ad esempio perché è un sito di legame, allora sarà conservato evolutivamente, cioè comparirà praticamente uguale in specie diverse.

Il secondo approccio è più indiretto. Una serie di altre tecniche sperimentali permette di misurare l'effetto finale di queste regolazioni, valutando la cosiddetta "espressione genica" (in termini più precisi, queste tecniche permettono di misurare per ogni gene la quantità di Rna messaggero, la molecola in cui viene trascritta l'informazione del Dna).

Da questa informazione, sofisticati metodi di inferenza permettono di ricostruire la rete di regolazione che ha la maggiore probabilità di aver generato lo schema di espressione genica osservato. Si tratta di metodi probabilistici, caratterizzati da un notevole grado di incertezza ma che se combinati con il metodo diretto possono dare ottimi risultati.

Siamo ancora molto lontani dall'aver una conoscenza esaustiva

della rete di regolazione di organismi complessi come l'uomo. Le reti che siamo in grado di costruire oggi sono ancora piene di lacune e di errori. Inoltre, è ormai chiaro che ogni tessuto ha una propria rete di regolazione, diversa da quella degli altri tessuti. Migliorare la conoscenza di queste reti è un processo che coinvolge centinaia di gruppi di ricerca in tutto il mondo. Si tratta di un obiettivo di grande importanza. Infatti, è ormai chiaro che molte delle patologie più complesse, dal cancro al diabete all'Alzheimer, sono in realtà dovute ad alterazioni di questa rete di regolazione. Una migliore comprensione della struttura delle reti di regolazione dei tessuti sani è il punto di partenza per riconoscere queste alterazioni e curarle. Un aspetto importante di questa linea di ricerca è stato il rendersi conto che non esistono due pazienti uguali. Ogni tumore è diverso dagli altri. Ogni tumore trova una via diversa per alterare la rete di regolazione e trovare la sua strada per invadere l'organismo. La sfida è quindi riuscire a trovare la cura giusta per ogni singolo paziente. È quella che si chiama "medicina di precisione", che è forse il regalo più grande che ci ha fatto l'era dei *big data*: la possibilità di costruire terapie personalizzate, su misura per ogni paziente, ottimizzando le possibilità di guarigione e minimizzando gli effetti nocivi dei farmaci.



c. Alla base della medicina di precisione è l'idea di integrare diverse fonti di informazione molecolare (nella figura sono rappresentate dal livello più basso: ogni nodo potrà essere un gene o una proteina o un metabolita) associandole alle possibili patologie (nella figura è il livello intermedio, ogni nodo rappresenta una malattia diversa) e associare queste informazioni alle caratteristiche peculiari dei singoli individui (rappresentati nella figura dal livello più alto, in cui ogni singolo individuo è caratterizzato non solo dalla sua particolare sequenza di Dna e dalla sua particolare rete di regolazione genica, ma anche dalle sue interazioni con gli altri individui e con l'ambiente).

#### Biografia

**Michele Caselle** è professore di fisica teorica presso l'Università di Torino. Da un po' di anni si interessa alle applicazioni della fisica dei sistemi complessi alla biologia molecolare e in questa veste collabora con l'Institut Curie di Parigi e con il dipartimento di Oncologia dell'Università di Torino su progetti di medicina di precisione.



# Bio-intelligenza artificiale

## Simulazioni di attività cerebrale

di Pier Stanislaò Paolucci



a.

Il sonno è un fenomeno cerebrale essenziale per l'apprendimento in tutte le specie viventi.

Quali sono i meccanismi fisici con cui il cervello percepisce, memorizza, decide e genera la coscienza di sé? È possibile creare sistemi artificiali capaci di prestazioni paragonabili a quelle delle bio-intelligenze, partendo dalla conoscenza dell'architettura e dei principi che governano il funzionamento del sistema nervoso?

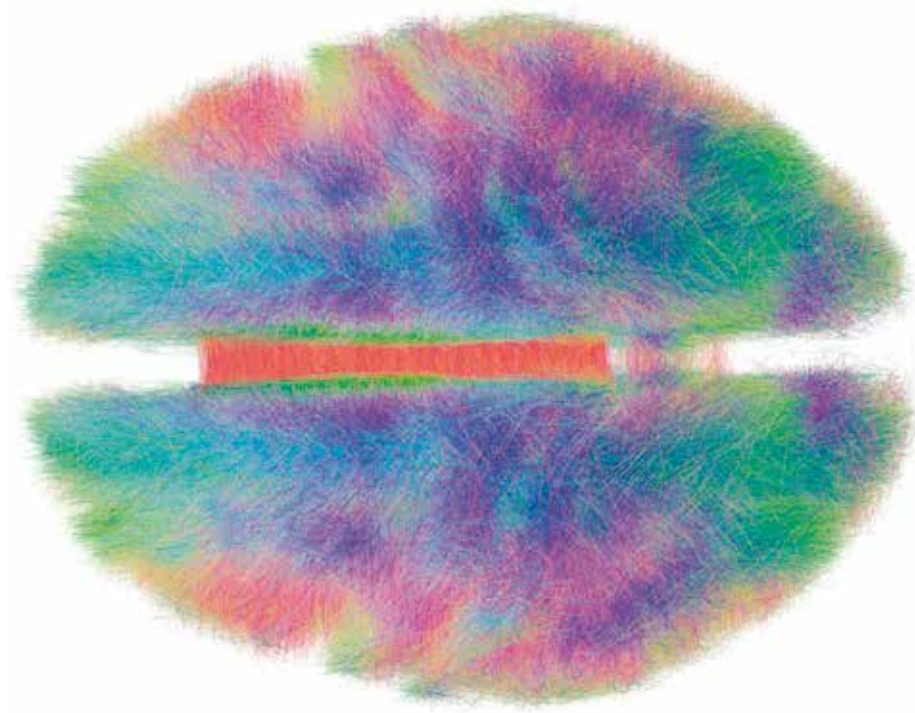
Intelligenza artificiale, robotica e macchine neuromorfe (cioè costruite incorporando meccanismi simili a quelli biologici) giocano un ruolo chiave nella rivoluzione tecnologica e industriale in atto, che avrà un radicale impatto sociale e profondi risvolti culturali. Alcuni esempi di applicazione di queste nuove tecnologie sono già apparsi, come i sistemi di riconoscimento del parlato e i prototipi di riconoscimento visivo per i sistemi di guida automatica. Inoltre, il costo legato alle patologie del sistema nervoso è stimato in circa 800 miliardi di euro per anno nella sola Europa, e quindi la comprensione della fisiologia e della patologia cerebrale

è di fondamentale importanza economica per i sistemi sanitari pubblici.

L'umanità s'interroga da sempre sui meccanismi che sostengono la percezione e il pensiero, ma l'applicazione di una metodologia scientifica a queste sfide è possibile solo ora, grazie alla combinazione d'innovative metodologie sperimentali e di nuovi paradigmi teorici. Si possono osservare in dettaglio la costituzione e l'attività del sistema nervoso: una moltitudine di dati immagazzinati e catalogati utilizzando computer paralleli (cioè con una molteplicità di processori che operano simultaneamente, vd. p. 20, ndr) che consentono anche di eseguire simulazioni per paragonare i modelli fisico-matematici con i dati sperimentali.

Per affrontare questi temi strategici, l'Unione Europea sta finanziando con circa cinquanta milioni di euro l'anno lo Human Brain Project (Hbp), che coinvolge attualmente più di 120 istituzioni di ricerca. Dal 2015, il laboratorio di calcolo

parallelo Ape dell'Infn è entrato in questo progetto alla guida di un esperimento denominato Wavescales, che combina metodologie sperimentali, teoriche e di simulazione. In particolare, Wavescales studia i ritmi cerebrali nei diversi stati cerebrali e la transizione tra sonno e veglia, e in particolare i meccanismi che rendono il sonno uno stadio cerebrale indispensabile. In effetti, nonostante si tratti di un comportamento pericoloso e apparentemente improduttivo, il sonno è un fenomeno essenziale presente in tutte le specie animali. Pericoloso, perché aumenta i rischi di predazione. Apparentemente improduttivo, perché sottrae energie ad attività di più evidente utilità come la ricerca di cibo o di un partner. Eppure, tutti gli esseri viventi dotati di un sistema nervoso dormono. La privazione di sonno danneggia la cognizione (e si tratta di una delle peggiori torture che possano essere inflitte). Viceversa, un buon sonno migliora le prestazioni intellettuali, l'apprendimento e la memoria. Quando si deve prendere una decisione importante si consiglia "dormici su, prima di decidere", e in effetti è conoscenza comune il fatto che al risveglio troviamo spesso pronta la soluzione di problemi irrisolti il giorno prima. E gli infanti, per i quali ogni avvenimento è nuovo e carico di informazioni importanti, passano la maggior parte del loro tempo a dormire, mentre viceversa un problema comune



b.  
Il connettoma: una mappa sperimentale delle connessioni tra aree cerebrali. Un ingrediente essenziale per le future simulazioni dell'attività cerebrale umana su calcolatori paralleli. I colori rappresentano la direzione lungo cui agiscono le sinapsi.

negli anziani è l'insonnia. Un fatto notevole è che nel sonno le attività del cervello sono guidate dall'interno, non dall'ambiente, e questa disconnessione, in qualche modo simile agli stati meditativi, può consentire la creazione di nuove associazioni e soluzioni ai problemi, che sarebbe molto più difficile costruire mentre si è impegnati nelle attività quotidiane, distratti dalla decifrazione degli stimoli ambientali e dalla elaborazione di una reazione appropriata.

In uno studio di recente pubblicazione (vd. link sul web), l'Infn fornisce un contributo rilevante alla comprensione del fenomeno, simulando in un modello talamo-corticale gli effetti benefici del sonno profondo sulle memorie apprese durante la veglia. Lo studio dimostra il funzionamento simultaneo di due processi che avvengono durante il sonno: il primo crea nuove associazioni, il secondo riequilibra la forza delle connessioni tra neuroni. Insieme, migliorano l'accuratezza con cui la corteccia cerebrale classificherà gli stimoli al successivo risveglio. Abbiamo ancora molto da imparare dalla bio-intelligenza, ma il meccanismo identificato dall'Infn avrà interessanti effetti di trasferimento tecnologico verso i futuri sistemi di intelligenza artificiale e robotica. Molteplici però sono le sfide scientifiche e tecnologiche da affrontare per ottenere sistemi artificiali con capacità cognitive realmente paragonabili a quelle umane. Per farci un'idea delle dimensioni di una simulazione cerebrale completa,

teniamo presente che il sistema nervoso umano contiene circa 100 miliardi di cellule specializzate (i neuroni). Ogni neurone è connesso a parecchie decine di migliaia di altri neuroni per mezzo di connessioni (sinapsi) che modificano la loro forza in una rete la cui struttura dipende dalla storia dell'individuo e dalla selezione naturale, evolutiva, vissuta dal cervello (individuale e della specie). In totale, un cervello contiene più di un milione di miliardi di sinapsi: un calcolatore con un elevatissimo grado di parallelismo. Inoltre, il cervello umano consuma poche decine di Watt, come una lampadina a basso consumo, mentre un computer di pari potenza computazionale richiederebbe una centrale elettrica dedicata: un ulteriore elemento che suggerisce i vantaggi di macchine neuromorfe per eseguire applicazioni di

bio-intelligenza artificiale.

Ma l'aspetto scientificamente più rivoluzionario è la creazione di nuovi schemi epistemologici che superano in un modo controllato, nello studio dei sistemi cerebrali, la tradizionale netta separazione tra osservatore umano e sistema fisico osservato. Si parte invece dall'assioma che la coscienza e la percezione esistano, e che gli altri esseri umani (e i mammiferi) condividano con noi questo tipo di percezioni. Si combinano, utilizzando metodi statistici, le osservazioni sperimentali quantitative con i report soggettivi delle percezioni e pensieri elaborati da una molteplicità di soggetti e si tenta di riprodurre il fenomeno in simulazioni quantitative di un sistema complesso. Una metodologia che richiede il contributo della fisica sperimentale e teorica.

#### Biografia

**Pier Stanislao Paolucci** coordina dal 2015 l'esperimento Wavescales nello Human Brain Project. È tornato all'Infn nel 2010 dopo un periodo dedicato al trasferimento tecnologico, in cui è stato fondatore e Cto di un centro di progettazione leader mondiale nella produzione di micro-processori. Paolucci ha iniziato la sua attività di ricerca nel gruppo di calcolo parallelo Ape dell'Infn nel 1984 ed è autore di brevetti internazionali sulle architetture di processori e di algoritmi internazionalmente utilizzati.

#### Link sul web:

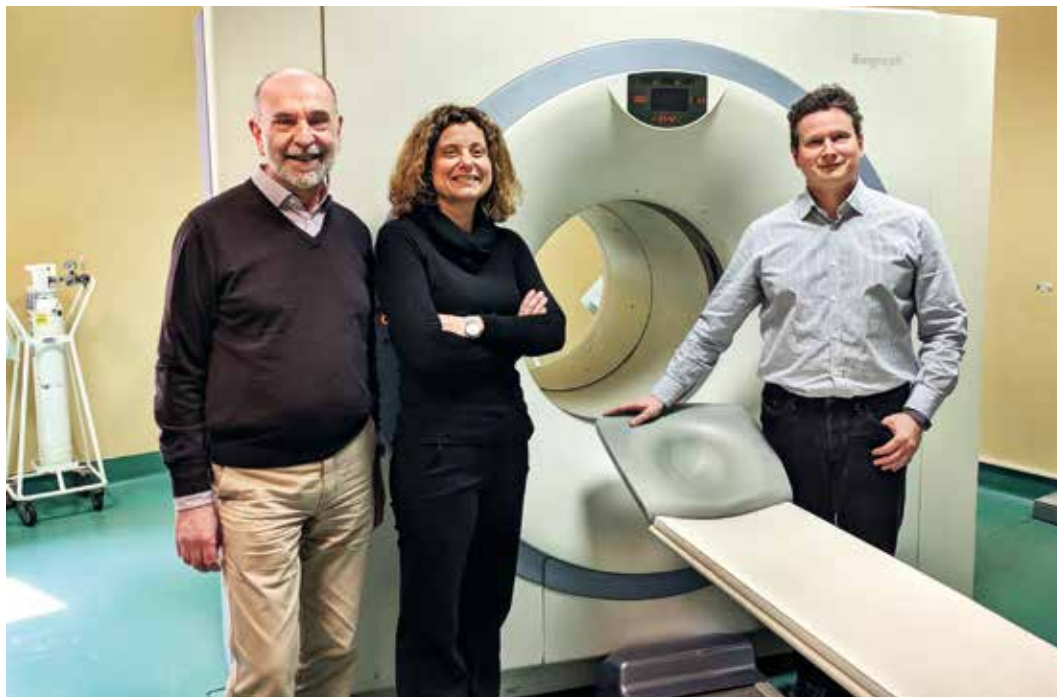
<https://www.nature.com/articles/s41598-019-45525-0>

<http://home.infn.it/it/comunicati-stampa/comunicati-stampa-2015/1648-il-progetto-wavescales-guidato-dall-infn-entra-nello-human-brain-project>

DOI: 10.23801/asimmetrie.2019.27.9

# Data mining in fisica medica.

di Francesca Mazzotta



a.

Tre dei ricercatori che hanno lavorato allo studio sulla diagnosi della demenza a corpi di Lewy, davanti a una Pet. Da sinistra a destra: Flavio Nobili e Silvia Morbelli, Università di Genova e Irccs S. Martino, Andrea Chincarini, Infn.

Orientarsi in una foresta di dati, cercando un ordine, una struttura o un significato all'interno di un gran numero di informazioni. È questo l'obiettivo del *data mining*, un insieme di tecniche matematiche che stanno entrando gradualmente in contatto con la ricerca medica: dalla classificazione dei tumori all'interpretazione di una radiografia, dalla diagnosi di alcune patologie all'individuazione automatica delle strutture anatomiche.

Il processo di collaborazione tra medicina e *data mining* risponde all'esigenza della ricerca medica di gestire campioni di dati molto eterogenei e inizia a ottenere i primi risultati. Tra questi, uno studio recente ha confermato che la tomografia a emissione di positroni (Pet) è un ottimo strumento per la diagnosi della demenza a corpi di Lewy (*Dementia with Lewy Bodies*, Dlb), una patologia simile alla malattia di Alzheimer, e ha delineato le complesse relazioni tra il metabolismo cerebrale, i sintomi della malattia e la sua gravità. Lo studio, che ha impiegato strumenti di *data mining* largamente usati in fisica, ha coinvolto fisici dell'Infn coordinati dal ricercatore Andrea Chincarini e medici del Policlinico San Martino e dell'Università di Genova.

“La ricerca nel campo del *neuroimaging* è particolarmente adatta all'applicazione del *data mining*, perché il suo oggetto di studio, il cervello, ha una struttura complessa, articolata, piena di collegamenti, quindi molto eterogenea,” spiega Chincarini. “In ogni caso, non ci fermiamo a questo importante

risultato per la diagnosi della Dlb: stiamo continuando a lavorare a stretto contatto con i medici per studiare altre patologie neurodegenerative come la malattia di Alzheimer, la sclerosi multipla o il morbo di Parkinson, utilizzando contemporaneamente più tecniche di *neuroimaging*: Pet, tomografia computerizzata (Tac) e risonanza magnetica”. Secondo Chincarini, oltre agli ottimi risultati immediati, lavorare su questo filone di ricerca, alla frontiera tra fisica, medicina e *data science*, sarà anche di grande impatto metodologico: per esempio, la medicina si sta dotando di modelli quantitativi che, tra i molti vantaggi, renderanno la ricerca medica più efficace e forniranno nuovi strumenti per arrivare a decisioni maggiormente informate sulle patologie. Si riuscirà poi a utilizzare al meglio le ricerche che coinvolgono più centri di ricerca, riuscendo a confrontare tra loro dati presi in tutto il mondo. Per quanto riguarda la ricerca in fisica di base, invece, lavorare sul *data mining* stimola la creazione di nuove figure professionali, i *data scientist*, che sanno legare il mondo variegato del *data mining* e quello della fisica. I *data scientist* conoscono a fondo gli strumenti matematici, che adattano alle caratteristiche dei dati da analizzare, così da individuare le tecniche e gli strumenti più adatti a un particolare esperimento. “All'Infn iniziamo ad avere figure professionali di questo tipo per migliorare la comprensione dei nostri esperimenti anche sotto l'aspetto delle proprietà dei dati”, conclude Chincarini.

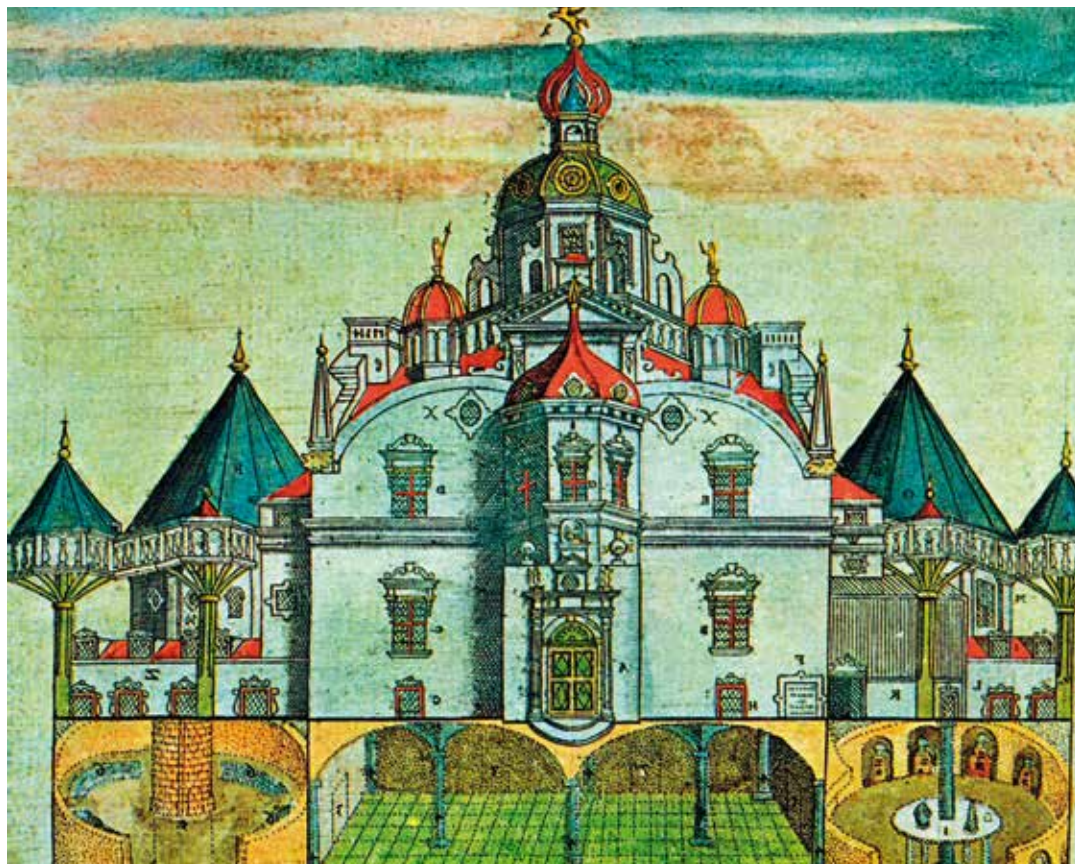


[as] radici

# I dati di Tycho.

di Paolo Rossi

*fisico teorico e storico della scienza*



a.  
L'osservatorio astronomico di Tycho Brahe a Uraniborg.

Per molti secoli è stata l'astronomia la scienza che ha maggiormente raccolto e fatto uso di grandi moli di dati, e un esempio significativo degli straordinari risultati cui ha condotto questa pratica è la scoperta delle leggi di Kepler dei moti planetari, a partire dalle osservazioni compiute dal danese Tyge Ottesen Brahe (1546-1601).

Di nobile famiglia, Brahe mostrò presto un forte interesse per l'astronomia. Tuttavia, sia in occasione dell'eclisse del 1560, sia per la congiunzione tra Giove e Saturno del 1563, egli notò la scarsa precisione delle predizioni basate sui dati disponibili ai suoi tempi. Si dedicò quindi a un loro sostanziale miglioramento, realizzando il

grande osservatorio di Uraniborg.

I risultati raccolti nell'arco di trent'anni gli permisero di sfatare l'immutabilità del mondo celeste, dimostrando che la supernova del 1572 e la cometa del 1577 dovevano essere molto lontane dalla Terra, in quanto la loro posizione nella volta celeste non mutava durante il giorno e quindi erano esterne al mondo sublunare, a differenza di quanto sostenuto da astronomi e filosofi.

La precisione delle osservazioni di Tycho (il nome latino usato da Brahe) superava di un ordine di grandezza tutti i risultati precedenti, riducendo l'errore a un minuto d'arco, il meglio che si potesse ottenere a occhio nudo.

Contrario all'ipotesi copernicana di una Terra in moto, Tycho cercò di conciliare i dati con il proprio pregiudizio formulando un modello in cui il Sole gira intorno alla Terra, ma i pianeti girano intorno al Sole. Il modello godette a lungo di una certa fortuna, soprattutto negli ambienti restii ad accettare la teoria eliocentrica. Un elemento "rivoluzionario" era comunque costituito dalla eliminazione delle "sfere" cristalline, incompatibili con l'esistenza di orbite destinate a incrociarsi.

La fama dell'astronomo danese, che nel frattempo si era trasferito a Praga, raggiunse il giovane matematico tedesco Johannes Kepler (1571-1630), che gli fece visita all'inizio dell'anno 1600, ma

solo a settembre fu accolto all'osservatorio, e comunque non ebbe accesso a tutti i dati che Brahe aveva accumulato e di cui era gelosissimo.

Kepler aveva già dedicato parecchia attenzione all'astronomia, con un deciso orientamento copernicano, ma nella sua prima importante opera, il *Mysterium Cosmographicum* (1596) si era concentrato sull'esposizione di un modello del sistema solare in cui le orbite dei pianeti erano legate ai solidi platonici, e i risultati di Tycho gli servivano per tentare di suffragare la propria teoria.

Al proprio maestro Michael Maestlin scrisse: "Tycho possiede le osservazioni più accurate esistenti al mondo, ma gli manca un architetto capace di costruire un edificio a partire dai suoi dati". Per (nostra) fortuna il limitato compito che Tycho affidò a Kepler fu lo studio dell'orbita di Marte nel quadro del modello semi-eliocentrico. Kepler scommise che avrebbe impiegato otto giorni, ma in realtà gli ci vollero molti anni. In ogni caso la collaborazione non era destinata a durare a lungo, perché Tycho morì nell'ottobre del 1601.

Agendo con spregiudicatezza, come egli stesso ammise, Kepler si impadronì dei dati prima ancora di essere nominato come successore di Tycho nel ruolo di astrologo imperiale. Poté così proseguire nel lavoro che era ormai il principale obiettivo della sua vita: la costruzione di un modello matematico dell'orbita di Marte che riproducesse con cura le osservazioni.

È cruciale sottolineare che nella sua ricerca Kepler era guidato dall'esigenza, assai "moderna", di individuare un meccanismo dinamico atto a giustificare la legge del moto dei pianeti. Poiché egli attribuiva la causa del moto a una "forza" originata dal Sole, il suo primo fondamentale progresso consisté nel riferire le posizioni dei pianeti al Sole stesso. Ciò gli permise, grazie ai dati di Tycho, di scoprire presto la "seconda legge", per cui il raggio che va dal Sole al pianeta "spazza" aree uguali in tempi uguali. La relazione inversa tra velocità e distanza, associata al pregiudizio di un legame diretto tra forza e velocità, lo convinsero dell'esistenza di una forza decrescente come  $1/r$ , contro la sua primitiva (e straordinaria) intuizione di una forza decrescente con  $1/r^2$ , la cui successiva dimostrazione avrebbe tuttavia richiesto la nozione di accelerazione centripeta e l'opera di Newton.

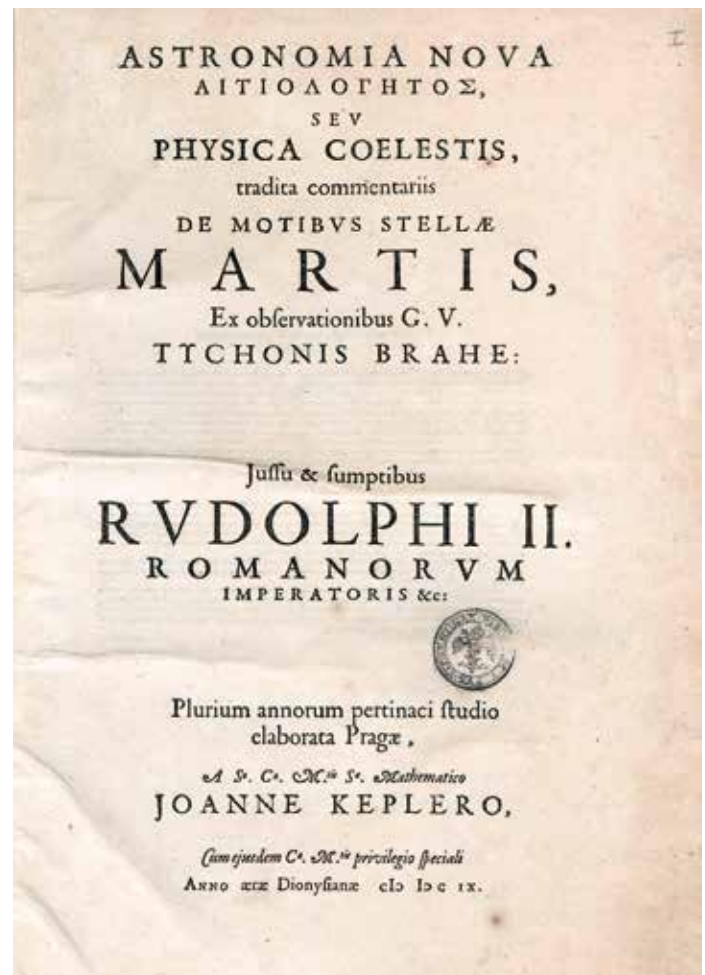
Kepler per cinque anni esplorò ogni possibilità che consentisse di ricondurre l'orbita di Marte a una circonferenza, come emerge dal suo fondamentale trattato *Astronomia nova* (1609). Il suo miglior risultato, per quanto assai aderente alle osservazioni, mostrava tuttavia deviazioni fino a otto minuti d'arco, del tutto inaccettabili alla luce della precisione dei dati.

Kepler si rassegnò quindi ad assumere che l'orbita potesse avere una forma ovoidale, e provò varie curve, arrestandosi sull'ellisse solo quando si accorse che in tal caso il Sole si sarebbe trovato esattamente in uno dei fuochi. L'accuratezza con cui quella scelta riproduceva i dati dell'orbita di Marte lo spinse a estendere senza altre verifiche questa sua "prima legge" a tutti gli altri pianeti. I dati di Tycho furono essenziali anche per giungere alla "terza legge", che lega i quadrati dei periodi dei pianeti ai cubi dei semiassemi maggiori delle orbite. Questa relazione non soddisfaceva Kepler, che dalla propria ipotesi sulla forza si sarebbe aspettato una relazione tra i periodi e i quadrati delle distanze, ma l'evidenza era inconfutabile.

Volendo trarre dall'intera vicenda una "morale" epistemologica, e una lezione per l'attualità, possiamo dire che, come spesso è accaduto nella storia della fisica, i soli dati, per quanto accurati,

non avrebbero permesso di giungere a una legge abbastanza generale e predittiva. I dati non bastano: occorre quasi sempre un solido pre-giudizio (attenti al trattino) teorico per costruire uno schema alla luce del quale interpretare i dati. Senza l'idea di una forza avente origine nel Sole, un'idea che nemmeno Galileo giunse a formulare o condividere, Kepler forse non sarebbe mai giunto a un risultato che non meno (e forse più) del cannocchiale "sgombrò primo le vie del firmamento".

b.  
Frontespizio della prima edizione di  
*Astronomia nova*.



[as] traiettorie

# Point 8.

di Catia Peduto



Kevin Dungs ha studiato fisica all'Università tecnica di Dortmund e durante la tesi inizia a lavorare all'esperimento Lhcb del Cern. Oltre al lavoro di ricerca, collabora alla creazione dello "starterkit", un progetto per introdurre i giovani fisici dell'esperimento all'uso ottimale del software di Lhcb. Per questo ha ricevuto il premio di Lhcb *Early Career scientist award*. Nel 2016 è andato a lavorare per Google a Zurigo e dall'inizio del 2018 è tra i soci di Point 8, una *startup* tedesca che offre servizi di *data science* e così chiamata in onore del punto di accesso al tunnel Lhc dove si trova l'esperimento Lhcb.

**[as]:** Kevin, com'è stata la tua esperienza al Cern?

**[Kevin]:** Il mio primo contatto con l'esperimento Lhcb è stato nel 2013 grazie alla tesi di laurea triennale in fisica sul *flavour tagging* (algoritmi per l'identificazione dei quark). Quell'estate, sono stato anche abbastanza fortunato da essere scelto come *summer student* per il gruppo Hlt (High Level Trigger) dell'esperimento Lhcb, che si occupa dell'analisi dati in tempo reale e con cui ho collaborato fino al 2016, anno in cui ho lasciato il laboratorio.

Al Cern ho avuto la possibilità di imparare molto. Innanzitutto per quanto riguarda la fisica. Poi, interagendo costantemente con persone di tutto il mondo, ho imparato di più sulle loro culture,

attitudini, abitudini, ecc. Inoltre, le persone non sono interessate solo alla fisica: quasi tutti hanno un hobby che praticano con la stessa passione che mettono nel proprio lavoro quotidiano. Il gruppo Hlt di Lhcb è stato particolarmente adatto a me. Lì, la programmazione era immensamente importante. Quindi, da una parte sono riuscito a mettere a frutto le mie particolari capacità di programmazione e *data analysis*. E dall'altra ho avuto la possibilità di lavorare con veri esperti del settore, dai quali ho imparato moltissimo.

**[as]:** Nel 2016 hai lasciato il Cern per lavorare per Google. Perché? Che cosa hai fatto in Google?

**[K]:** Come ho già accennato, la programmazione è stato un aspetto molto importante del mio percorso al Cern. Con il tempo mi sono reso conto che non mi importava tanto della fisica, quanto delle persone intorno a me. Poi, ovviamente, lavorare per Google era il sogno della mia vita! In Google ho lavorato come *software engineer*. Prima di andare lì, pensavo di essere un programmatore piuttosto bravo. Ma in Google ho imparato la differenza tra programmazione e ingegneria del software: una distinzione che tipicamente non viene fatta nel mondo della fisica. Ancora una volta, come al Cern, lavoravo in un fantastico ambiente internazionale,



circondato da persone estremamente brillanti e motivate. È stato un lavoro molto divertente, molto impegnativo che non mi sarei mai aspettato di lasciare... eppure è successo!

**[as]:** Incredibile! E cosa ti ha spinto a lasciare il lavoro dei tuoi sogni a Google?

**[K]:** Già nel 2013, con alcuni amici a Dortmund avevamo avuto l'idea di fondare un'azienda per supportare la trasformazione digitale e i progetti di automazione dell'industria (la cosiddetta industria 4.0). Ma nel momento in cui dovevamo davvero fondare la società, è arrivata l'offerta per la posizione

di dottorato al Cern. Quindi, a causa dei rischi che comportava, gli altri mi suggerirono di fare il mio dottorato e di unirmi a loro solo più tardi, se la compagnia fosse stata ancora in piedi. Hanno fatto un ottimo lavoro e alla fine del 2017 la società stava andando bene e loro volevano crescere. Per questo, ho fatto il mio ingresso in azienda. Fare questo passo è stata probabilmente la decisione più dura della mia vita. Alla fine, la cosa che mi ha convinto maggiormente è stata che avrei potuto lavorare con i miei più cari amici. Basandoci sulla nostra esperienza nel lavorare con enormi quantità di dati, in Point 8 offriamo la *data science* come servizio alle piccole e medie imprese

dell'industria manifatturiera e industriale tedesca. Quando dico "come servizio", non intendo "nel Cloud", intendo un servizio come quello di un sarto che fa vestiti su misura: lavoriamo a stretto contatto con i nostri clienti per capire i loro bisogni, il loro modo di lavorare e le loro macchine, per riuscire ad aiutarli a ricavare un valore reale dai loro dati. Nella nostra azienda, formata da 15 persone, la maggior parte dei nostri dipendenti sono fisici delle particelle o hanno un *background* scientifico simile. Quello che facciamo, quindi, è applicare le abilità di *data science* e gli strumenti di *machine learning* acquisiti al Cern alla nostra attività quotidiana con l'industria tedesca.

b.  
L'accesso esterno al Point 8 del tunnel di Lhc, dove è situato l'esperimento Lhcb, il cui schema è disegnato a grandezza naturale intorno alle sbarre di ingresso.



[as] selfie

# Blazar, che passione!

di Elettra Colonna e Chiara Neglia

studentesse del Liceo Scientifico "A. Scacchi" di Bari



Il giorno 5 aprile 2019 noi studenti del Liceo Scientifico "A. Scacchi" di Bari abbiamo avuto la possibilità di partecipare a un'esperienza unica nel suo genere: la Fermi Masterclass, organizzata dall'Infn e tenutasi presso la locale sezione di Bari, al Dipartimento Interateneo di Fisica "M. Merlin". Per noi studenti, una *"full immersion"* nel mondo della fisica astroparticellare. Il satellite Fermi, in orbita dall'11 giugno 2008, ospita un particolare rivelatore, il Lat (Large Area Telescope), che è studiato per essere sensibile alle frequenze più alte, i raggi gamma, dello spettro elettromagnetico. Il Lat è basato su una tecnologia di tracciatori a microstrip di silicio. L'emissione di raggi gamma è associata a fenomeni o sorgenti che coinvolgono energie elevatissime. Paradossalmente, alcune di queste sorgenti non potrebbero essere visibili con un telescopio convenzionale perché, soprattutto nella direzione del piano galattico,

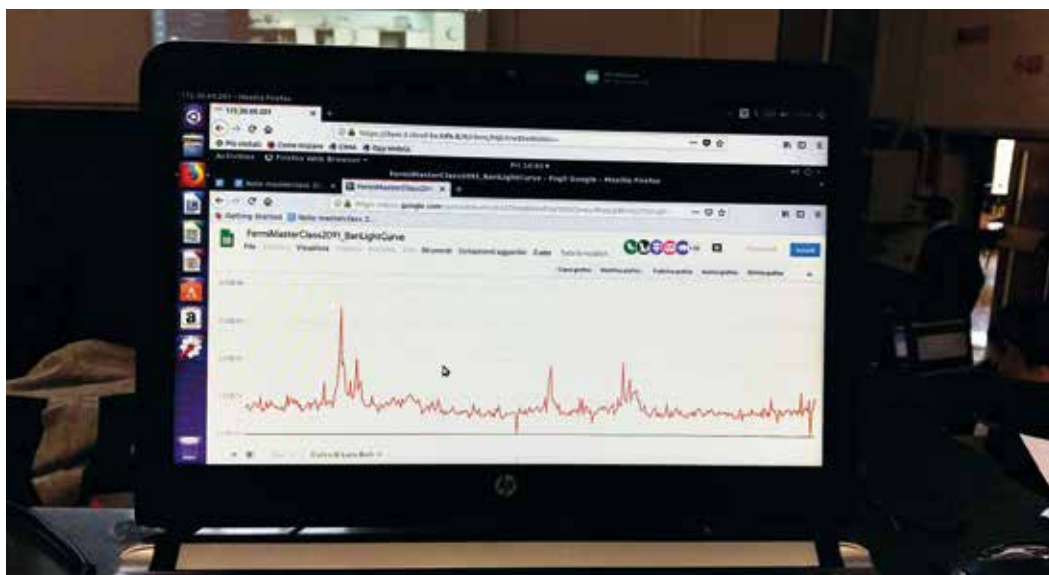
la materia stellare stessa assorbe molta radiazione visibile. La spettroscopia gamma rientra, perciò, in un progetto più ampio, denominato astronomia multimessaggera, che coinvolge simultaneamente tutta la banda di frequenze dello spettro elettromagnetico, compresi i raggi X, l'ultravioletto, l'infrarosso e le onde radio.

Durante la mattinata abbiamo seguito lezioni sul funzionamento del Lat, abbiamo scoperto che cosa è un Agn (nucleo galattico attivo), in particolare abbiamo studiato i *blazar*, e per finire abbiamo consolidato le conoscenze attraverso un istruttivo e divertente quiz a risposta multipla sulla app Kahoot.

Nella seconda parte della giornata siamo entrati nel vivo dell'esperienza analizzando un set di dati raccolti dal Lat. Il nostro obiettivo era quello di ricostruire, attraverso un opportuno software, la "curva di luce" del *blazar* denominato PKS1830-211 (coordinate

a.  
Chiara ed Elettra durante la Fermi Masterclass a Bari, lo scorso 5 aprile.





b.  
La "curva di luce", cioè come cambia la luminosità nel cielo gamma nel corso degli anni, del *blazar* AGN PKS1830-211 che è stato studiato durante la Fermi Masterclass di Bari.

equatoriali celesti: ascensione retta =  $278,41^\circ$ ; declinazione =  $-21,06^\circ$ ). A tal fine siamo stati suddivisi in 20 postazioni e a ciascun gruppo è stato assegnato un intervallo di tempo di 6 mesi per ricoprire i 10 anni di attività dell'esperimento Fermi. Un doveroso ringraziamento va ai docenti e ricercatori che ci hanno supportato, grazie ai quali siamo stati in grado di utilizzare la *software suite* necessaria per eseguire l'analisi. Infine, i dati relativi ai diversi intervalli temporali sono stati inseriti in un unico grafico, che ci ha permesso di studiare la variazione della luminosità della nostra sorgente di raggi gamma in funzione del tempo. Dall'osservazione

della curva ottenuta, abbiamo potuto rilevare la presenza di picchi che segnalavano un aumento di radiazione gamma emessa dalla sorgente presa in esame. I risultati sono stati poi confrontati con quelli ottenuti da altri studenti, contemporaneamente, in altri centri di ricerca in Italia e in Slovenia, i quali hanno effettuato uno studio simile in intervalli di tempo diversi e su altre sorgenti.

La giornata si è conclusa con il collegamento in videoconferenza con un ricercatore della Nasa, dimostratosi subito disponibile ad assecondare ogni nostra curiosità con ricche ed esaurienti risposte. Ogni sua parola

è stata in grado di trasmetterci quella stessa fervida passione che gli ha consentito di realizzare il suo sogno, superando insidie e ostacoli. Insomma, sebbene l'infinità dell'universo, i suoi innumerevoli misteri e i segreti da esso racchiusi possano violare la nostra sicurezza conoscitiva, incutendo timore e alimentando la nostra rinuncia a svelarli, l'essere umano tenderà sempre di varcare tali limiti, riuscendo a seguire quella sola traccia capace di cambiare tutto, di stravolgere i precedenti schemi, offrendo l'accesso verso nuove possibilità che accompagnino l'uomo nella scoperta di se stesso e del mondo che lo circonda.

## [as] approfondimento

# La Fermi Masterclass

La Fermi Masterclass è un'iniziativa rivolta alle scuole superiori, coordinata in Italia dall'Istituto Nazionale di Fisica Nucleare (Infn), arrivata alla terza edizione. Nel 2019 si è svolta il 5 aprile nelle università delle città di Bari, Perugia, Roma "Sapienza", Torino e Trieste, nonché all'estero nella University of North Florida (USA) e la University of Nova Gorica (Slovenia). I ragazzi si recano nelle università e nei laboratori delle città coinvolte, dove vengono accompagnati dai ricercatori in un viaggio alla scoperta dei segreti dell'Universo, grazie allo studio dei veri dati del Large Area Telescope, uno strumento che si trova a bordo del satellite Fermi. Durante la mattinata, gli studenti assistono a lezioni e seminari introduttivi sugli esperimenti nello spazio, sull'analisi dei dati e

su argomenti fondamentali della fisica delle astroparticelle. Nel pomeriggio, vengono impegnati in esercitazioni al computer con i dati di Fermi, attraverso le moderne tecniche di analisi usate dagli scienziati. Poi, come in una vera collaborazione internazionale, si collegano in videoconferenza con i coetanei delle altre sedi per discutere insieme i risultati, e con un ricercatore della Nasa. Fermi è il satellite della Nasa dedicato allo studio della radiazione gamma di alta e altissima energia, costruito e operato da un'ampia collaborazione internazionale, cui l'Italia partecipa, oltre che con l'Infn, anche con l'Agenzia Spaziale Italiana (Asi) e l'Istituto Nazionale di Astrofisica (Inaf). [Catia Peduto]



[as] spazi

# Il colore dell'invisibile.

di Eleonora Cossi

a.  
Studenti, docenti e ricercatori partecipanti allo Stage Ocra ad aprile 2019 a Campo Imperatore, assieme al rivelatore Cosmic Ray Cube per la misura del flusso dei muoni atmosferici.



Non si vede ma c'è, e catturarla è possibile, ma soltanto con occhi tecnologici molto sensibili. Anche per questo è così affascinante. Parliamo del fenomeno della radiazione cosmica, che ci parla di un universo lontano e misterioso e arriva sulla Terra sotto forma di un intenso flusso di particelle che i fisici chiamano "raggi cosmici" (vd. *Asimmetrie* n. 10, ndr). Parlare di raggi cosmici vuol dire parlare dell'esplorazione dell'universo in tutte le sue forme: a questo tema così vasto e suggestivo è dedicato il progetto Ocra che raccoglie le attività di *outreach* sul tema dei raggi cosmici (in inglese Outreach Cosmic Ray Activities) dell'Infn.

"Ocra è nato nel 2018 con l'obiettivo di raccogliere in un unico progetto nazionale le tante attività di *public engagement* nel campo della fisica dei raggi cosmici già presenti a livello locale nelle sezioni e nei laboratori dell'Infn, con una particolare attenzione al mondo della scuola", sottolinea Carla Aramo, una delle due responsabili nazionali del progetto e ricercatrice della sezione di Napoli dell'Infn. Il progetto coinvolge ogni anno quasi 100 scuole e 800 studenti e vede la partecipazione

di 21 sedi locali dell'Infn: Bari, Catania, Cosenza, Firenze, Genova, Gran Sasso Science Institute (Gssi), Lecce, Laboratori Nazionali del Gran Sasso (Lngs), Milano, Milano Bicocca, Napoli, Padova, Perugia, Pisa, Pavia, Roma "Sapienza", Laboratori Nazionali del Sud (Lns) con la sede di Sassari, Roma "Tor Vergata", Trento Institute for Fundamental Physics and Applications (Tifpa), Torino e Trieste.

Una delle attività principali del progetto Ocra è la partecipazione delle scuole italiane all'International Cosmic Day (lcd, vd. anche in *Asimmetrie* n. 23 p. 26, ndr), una giornata dedicata agli studenti delle superiori e a cui nel 2018 hanno aderito 2250 ragazzi e ragazze di 16 paesi diversi e 30 istituti di ricerca guidati da alcuni dei più importanti laboratori per la fisica delle particelle nel mondo, come il tedesco Desy e l'americano Fermilab. Questa iniziativa, che quest'anno si svolge il 6 novembre, ha lo scopo di far incontrare studenti, insegnanti e ricercatori per scoprire e approfondire le proprietà e il significato delle informazioni che ci arrivano dall'universo attraverso i raggi cosmici. "Durante l'lcd i ragazzi hanno la

possibilità di vivere una vera esperienza di ricerca” spiega Sabine Hemmer, l'altra responsabile nazionale del progetto Ocra e tecnologa della sezione di Padova dell'Infn. La parte teorica è affidata a seminari introduttivi tenuti da ricercatori e ricercatrici. Poi si passa alla parte pratica con la misura dal vivo del flusso di raggi cosmici. L'esperienza si conclude, come in una vera collaborazione internazionale, con un collegamento con altri gruppi di studenti di tutto il mondo, che stanno svolgendo le stesse misure in altri laboratori, per confrontare i dati ottenuti. Un'altra importante attività è lo Stage per studenti Ocra, a cui partecipano 50 studenti e studentesse selezionati attraverso un concorso che si svolge ogni anno tra i partecipanti all'Icd. Per l'edizione 2019, i vincitori, due per ogni sede partecipante, avranno l'opportunità di trascorrere tre giorni presso i Laboratori di Frascati ed essere coinvolti nel lancio di un pallone aerostatico per misurare il flusso dei raggi cosmici durante il volo. La selezione avviene

in base a un elaborato preparato da ciascuna coppia di studenti, che descrive l'attività svolta durante l'Icd e in cui viene discussa l'analisi dei dati raccolti. Nella passata edizione sono stati selezionati 30 studenti che hanno seguito uno stage di tre giorni presso il Gssi e i Lngs, dove sono stati coinvolti nella misura del flusso dei raggi cosmici a diverse quote (dai laboratori sotterranei fino a Campo Imperatore), utilizzando il rivelatore Cosmic Ray Cube. Il progetto Ocra viene completato da una serie di iniziative locali, come il lancio di un pallone stratosferico, di nome Mocris, organizzato dalla sezione Infn di Roma “Sapienza” e che si svolge da due anni in Calabria, in collaborazione con il Liceo Scientifico di Cariati e l'azienda AB Project. Mocris è dotato di due rivelatori di raggi cosmici del tipo Arduisipm (vd. in *Asimmetrie* n. 19 p. 48, ndr). Lo scorso giugno ha raggiunto la quota record per la Calabria di 34.111 m e poi è atterrato su un albero, dove è stato recuperato grazie a un drone. Gli studenti di Cariati hanno lavorato alla

progettazione del pallone, contribuito al lancio e all'analisi dei dati. Una delle foto scattate da Mocris nella stratosfera, raffigurante Puglia, Basilicata e Calabria, è stata pubblicata come Earth Picture of the Day dalla Universities Space Research Association. Inserito in Ocra anche il concorso “A scuola di astroparticelle”, arrivato alla sua terza edizione e organizzato dalla sezione di Napoli dell'Infn, a cui ogni anno partecipano oltre 600 studenti di 18 scuole campane. Agli studenti viene chiesto di realizzare progetti su tematiche attuali della ricerca scientifica. L'idea del concorso nasce dall'installazione nella stazione Toledo della metropolitana di Napoli di un totem multimediale collegato a un telescopio per raggi cosmici installato nella stessa stazione. L'iniziativa si svolge anche nell'ambito dei percorsi di Alternanza Scuola Lavoro (Asl). Altre iniziative locali comprendono stage e percorsi Asl per studenti delle superiori in varie sedi (Milano e Milano Bicocca, Lecce, Bari, ecc.).



b.  
Il lancio del pallone Mocris l'8 giugno 2019.

Per maggiori informazioni:  
<https://web.infn.it/OCRA/>

[as] illuminazioni

# Immergersi nell'universo.

L'universo in un palmo di mano. È questa l'idea di Big Bang AR, l'applicazione di realtà aumentata sviluppata dal Cern, in collaborazione con Google, e rivolta a chiunque sia interessato a conoscere, e a interpretare, la storia dell'universo e i meccanismi che ne regolano l'evoluzione. Il ricorso alla realtà aumentata, in questo caso, non è un semplice espediente alla moda per attrarre gli appassionati di nuove tecnologie, ma il linguaggio più appropriato per un racconto, quello sull'universo di cui siamo parte, che non può che essere immersivo.

Punto di riferimento dell'applicazione, che le permette di assumere il punto di vista di chi guarda, è la mano dell'utente: registrate le coordinate della sua posizione, Big Bang AR è l'occhio di chi guarda, capace di osservare nei dettagli che cosa lo circonda e di ricostruire la formazione e l'evoluzione dell'universo dalla sua origine. Il viaggio di 13,8 miliardi di anni dal Big Bang a oggi dura una quindicina di minuti, sempre che non ci si voglia soffermare a osservare i tanti spunti che si incontrano lungo la strada. E la scoperta più straordinaria è forse proprio questa: non importa in quale direzione rivolgiamo lo sguardo (la mano), l'universo si presenta sempre nello stesso modo: lo stesso fondo cosmico a microonde, la

stessa composizione di stelle, la stessa densità di galassie. Guidati dalla voce narrante dell'attrice Tilda Swinton (nella versione in inglese) e da musiche dal sapore "cosmico", è possibile interrogare l'universo interagendo con gli elementi che appaiono sullo schermo e assistere alla nascita delle prime stelle, del nostro sistema solare e della Terra stessa. Si può inoltre far ruotare il cellulare o il tablet per osservare stelle, galassie e pianeti da angolazioni diverse od osservare da vicino la formazione di tutte le strutture che hanno portato all'esistenza di quello che conosciamo: quark, protoni e neutroni, nuclei leggeri, atomi e molecole, cellule e organismi multicellulari.

Il viaggio attraversa le più grandi scoperte fatte nel secolo scorso e fino a oggi, per scoprire, tra le altre cose, che l'età della Terra è solo un terzo dell'età dell'universo, che le particelle del nostro corpo sono eterne e che noi abbiamo, di fatto, la stessa età dell'universo.

Presentata il 6 marzo scorso all'evento di Google Arts and Culture di Washington D.C, l'applicazione è disponibile per dispositivi Android su <https://play.google.com/store/apps/details?id=ch.cern.BigBangAR> e per Apple su <https://itunes.apple.com/app/id1453396628>. [Francesca Scianitti]







Istituto Nazionale di Fisica Nucleare

**I laboratori dell'Istituto Nazionale  
di Fisica Nucleare sono aperti alle visite.**

I laboratori organizzano, su richiesta  
e previo appuntamento, visite gratuite  
per scuole e vasto pubblico.

La visita, della durata di tre ore circa,  
prevede un seminario introduttivo  
sulle attività dell'Infn e del laboratorio  
e una visita alle attività sperimentali.

Per contattare  
i laboratori dell'Infn:

*Laboratori Nazionali di Frascati (Lnf)*

T + 39 06 94032423  
/ 2552 / 2643 / 2942  
sislnf@lnf.infn.it  
www.lnf.infn.it

*Laboratori Nazionali del Gran Sasso (Lngs)*

T + 39 0862 4371  
(chiedere dell'ufficio prenotazione visite)  
visits@lngs.infn.it  
www.lngs.infn.it

*Laboratori Nazionali di Legnaro (Lnl)*

T + 39 049 8068342 356  
direttore\_infn@lnl.infn.it  
www.lnl.infn.it

*Laboratori Nazionali del Sud (Lns)*

T + 39 095 542296  
sislns@lns.infn.it  
www.lns.infn.it

**www.infn.it**



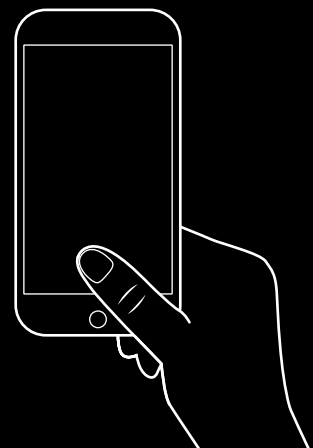
Sul sito **www.asimmetrie.it**  
vengono pubblicate periodicamente  
notizie di attualità scientifica.

Per **abbonarti gratuitamente** ad Asimmetrie  
o per **modificare** il tuo abbonamento vai su:  
<http://www.asimmetrie.it/index.php/abbonamento>

Si prega di tenere sempre aggiornato il proprio  
**indirizzo mail** per ricevere le nostre comunicazioni.

Leggi anche le nostre **faq** su:  
<http://www.asimmetrie.it/index.php/faq>

Asimmetrie è anche una app,  
ricca di contenuti multimediali.





**[www.infn.it](http://www.infn.it)**

rivista online  
[www.asimmetrie.it](http://www.asimmetrie.it)