

MACHINE LEARNING: LA CAPACITÀ DI PREVEDERE APPLICATA ALLA RICERCA E ALLA PRATICA CLINICA

Machine learning: the ability to predict applied to research and clinical practice

Ornella Colpani^{1,2}

¹Servizio di Epidemiologia e Farmacologia Preventiva (SEFAP), Dipartimento di Scienze Farmacologiche e Biomolecolari (DiSFeB), Università degli Studi di Milano

²Fondazione Bruno Kessler, Trento

Keywords

Machine learning
Medicine
Big data
Stratification

Abstract

Modern medicine is experiencing nowadays the era of -omic approaches, smartphones and health apps, generating a huge amount of data, the so-called big data. Such context inevitably requires new strategies and tools to store, integrate, and analyze data, and machine learning seems to be a promising framework to address these issues. Machine learning consists of algorithms that allow the machine to learn from data, without being explicitly programmed. The great potential of machine learning algorithms lies in their ability to make prediction or give a more synthetic and effective version of data, through dimensionality reduction and clustering methods. Machine learning can be a key tool in many branches of the biomedical research field, for instance for clustering genome regions according to similarities in epigenomic properties, or in identifying putative miRNA target genes, or to predict drug-target interactions. Nevertheless, the multiplicity of machine learning applications also embraces the clinical practice: patient stratification and prediction of risk scores are crucial for more accurate diagnosis and better therapeutic strategies. However, when training machine learning models, especially when applied to medical field, it is essential to pay attention to both the quality of data and to the effectiveness of the chosen algorithms for the investigated task. Finally, reproducibility and robustness of the learning procedures are crucial for their translation into clinical practice, forcing the FDA in introducing new guidelines in their approval protocols to deal with such brand new scenario.

Introduzione

Sebbene intelligenza artificiale, machine learning e deep learning siano spesso usati come sinonimi, esistono differenze importanti da chiarire. Infatti, l'intelligenza artificiale è un ramo dell'informatica che consente la creazione di sistemi che permettono di dotare le macchine di caratteristiche considerate proprie dell'intelligenza umana. Machine learning e deep learning fanno parte dell'intelligenza artificiale.

Intelligenza artificiale (AI), *machine learning* (ML) e *deep learning* (DL) stanno acquisendo una rilevanza sempre maggiore in ambito medico. Tuttavia, c'è parecchia confusione sull'utilizzo di questi tre termini, che vengono spesso usati impropriamente come sinonimi. È quindi necessario fare un po' di chiarezza su cosa siano esattamente AI, ML e DL e in cosa differiscano.

L'AI si riferisce a un campo dell'informatica dedicato alla creazione di sistemi che svolgono compiti che normalmente richiederebbero l'intelligenza umana. Nell'intelligenza artificiale le macchine completano l'attività in base alle regole e agli algoritmi stabiliti. Il termine "intelligenza artificiale" è quindi generico e comprende anche ML e DL. Il ML rientra nell'AI e si concentra su tutti quegli approcci che permettono alle macchine di imparare dai dati, senza che queste siano programmate esplicitamente. Un aspetto caratterizzante del ML è la dinamicità. Le macchine, infatti, ricevono una serie di dati e sono capaci di apprendere, modificando e migliorando le predizioni man mano che ricevono più informazioni su quello che stanno elaborando. Dunque, con l'apprendimento, gli algoritmi di ML tenderanno di minimizzare gli errori e massimizzare la probabilità che le loro previsioni siano vere [1]. Il DL è un sottogruppo del ML e incorpora modelli e algoritmi computazionali che imitano l'architettura delle reti neurali biologiche nel cervello (reti neurali artificiali, ANN) [1, 2].

Corrispondenza: Ornella Colpani. Servizio di Epidemiologia e Farmacologia Preventiva (SEFAP), Dipartimento di Scienze Farmacologiche e Biomolecolari (DiSFeB), Università degli Studi di Milano, via Balzaretti, 9 – 20133, Milano. E-mail: ornella.colpani@unimi.it

Box 1 Definizioni.

Intelligenza artificiale: campo dell'informatica dedicato alla creazione di sistemi in grado di risolvere problemi e di riprodurre attività proprie dell'intelligenza umana.

Machine learning: applicazione dell'intelligenza artificiale che consente alle macchine di imparare dai dati, senza che queste siano programmate in maniera esplicita.

Deep learning: sottogruppo del *machine learning* basato su modelli e algoritmi computazionali che imitano l'architettura delle reti neurali biologiche nel cervello.

Supervised learning: approccio nel quale alla macchina viene fornito un dataset di esempio (*training set*) per la generazione di diversi algoritmi capaci di correlare le caratteristiche dei dati con un'etichetta (per esempio la relazione tra determinate caratteristiche cliniche con l'etichetta «soggetto sano/malato»). Questi algoritmi vengono poi validati sul *validation set* per poter selezionare il più performante. Infine, l'algoritmo selezionato viene testato su un terzo set di dati (*test set*). La macchina impara quindi a fare previsioni su nuovi dati non etichettati.

Unsupervised learning: approccio utilizzato quando le etichette delle variabili input sono ignote. La macchina impara solo dai pattern nelle caratteristiche dei dati di input. L'obiettivo è quindi quello di fornire una versione sintetica dei dati tramite il loro raggruppamento sulla base di caratteristiche simili o la riduzione delle loro dimensioni.

Reinforcement learning: approccio nel quale la macchina riceve in input un obiettivo da raggiungere: per ogni azione che la porta verso l'obiettivo riceve un premio, per ogni azione che la allontana dallo scopo ottiene una penalità.

Approcci omici: approcci che consentono la produzione di grandi quantità di dati utili per la descrizione e l'interpretazione del sistema biologico studiato. Ne sono esempio la genomica, la metabolomica, la proteomica, la trascrittomica e l'epigenomica.

Epigenomica: studio del set completo di modifiche epigenetiche (modificazioni della cromatina sia a livello del DNA che delle proteine) nel materiale genetico di una cellula, noto come epigenoma. Le modifiche epigenetiche sono estremamente importanti in quanto regolano l'espressione genica.

High-throughput imaging: flusso di operazioni automatizzate volte a testare gli effetti di agenti perturbanti sul fenotipo cellulare. Tale flusso prevede l'incubazione delle cellule con l'agente perturbante, l'acquisizione di immagini tramite microscopi automatizzati e l'analisi delle immagini.

Machine learning: concetti chiave

Nel machine learning la macchina apprende dai dati. Esistono due tipi principali di apprendimento: quello supervisionato e quello non supervisionato.

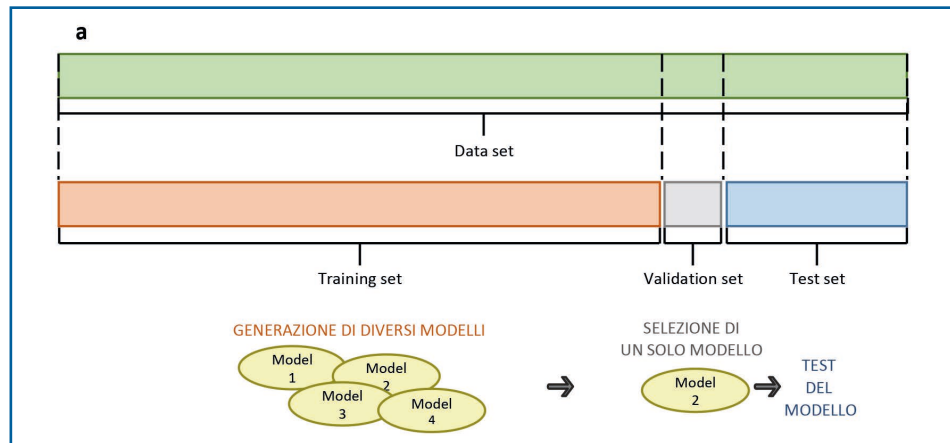
Il ML è la disciplina scientifica che si concentra su come i computer apprendono dai dati. È il risultato dell'incontro tra la statistica (che cerca di individuare relazioni dai dati) e l'informatica ed ha lo scopo di costruire modelli statistici da set di dati molto vasti. Esistono diverse categorie di apprendimento: qui menzioneremo quello supervisionato (*supervised learning*) e quello non supervisionato (*unsupervised learning*). Non parleremo invece dell'apprendimento per rinforzo (*reinforcement learning*), del quale è possibile trovare una breve spiegazione nel Box1.

Supervised learning

Nel *supervised learning* si parte da un set di dati che viene diviso in tre parti: *training set*, *validation set* e *test set* (**Figura 1**). Per cominciare, viene fornito alla macchina un dataset di esempio (*training set*) nel quale i dati sono etichettati. In altre parole, gli esempi sono composti da una coppia di dati contenenti il dato originale e il risultato atteso (variabile di risposta). Se pensassimo all'ambito clinico, questo significherebbe fornire alla macchina una dataset composto da casi clinici (che includono le caratteristiche di ogni paziente, per esempio analisi biochimiche, parametri vitali, ecc.) in cui l'esito per ogni paziente (sano o malato) è noto. Il compito della macchina è trovare la funzione che modelli la relazione tra i due, in modo tale da saper fare previsioni su nuovi dati in cui l'esito (nell'esempio precedente, sano/malato) è sconosciuto. Più formalmente, questi metodi servono ad individuare una funzione che predica la variabile di risposta (y) da un vettore di caratteristiche x che contiene M variabili input, in modo tale che $f(x)=y$. Se la variabile di risposta è di tipo categorico si parla di classificazione, nel caso di variabili di risposta continue (per esempio il costo di una casa sulla base delle sue caratteristiche) si parla invece di regressione. Esistono diversi algoritmi di *learning* in grado di creare funzioni che eseguano classificazioni o regressioni: *support vector machine*, alberi decisionali, *random forest* e *boosting* [3, 4]. I modelli/

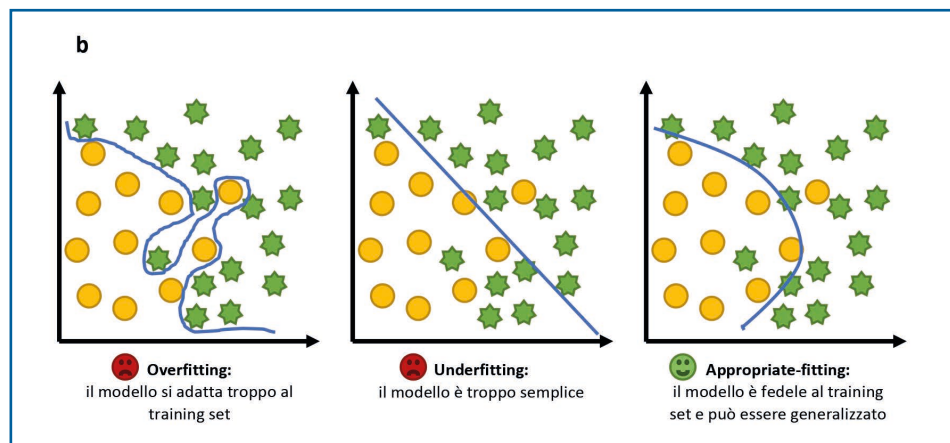
le funzioni generati durante il *training* vengono poi validati sul *validation set*. Questo permette di selezionare il modello più performante, per esempio in termini di accuratezza predittiva. L'ultimo passaggio consiste nel valutare la performance del modello selezionato su un nuovo set di dati, il *test set*.

Figura 1 *Supervised learning*: divisione del dataset in *training set*, *validation set* e *test set*.



Questi passaggi sono fondamentali per garantire di non incorrere nel cosiddetto *overfitting*, cioè la situazione in cui la funzione generata sia troppo sensibile alle caratteristiche particolari del *training set* piuttosto che alle caratteristiche generali del problema in analisi. In caso di *overfitting*, la funzione avrebbe performance di predizione altissime nel *training set* ma la sua performance risulterebbe molto inferiore sul *test set*. Nella **Figura 2** si può trovare una rappresentazione grafica di questo concetto [4, 5].

Figura 2 Rappresentazione schematica dei concetti di *overfitting*, *underfitting* e *appropriate-fitting*. Con bias/variance dilemma si intende la ricerca di un compromesso per ottenere un modello che non si adatti eccessivamente al training set e che non sia troppo semplice.



Unsupervised learning

Contrariamente a quanto visto per il *supervised learning*, gli approcci di *unsupervised learning* sono utilizzati quando le etichette delle variabili input non sono note e, dunque, questi metodi imparano solo i pattern delle caratteristiche dei dati di input. Gli obiettivi dell'*unsupervised learning* sono quindi: capire la distribuzione dei dati, il raggruppamento dei dati sulla base di caratteristiche simili o la riduzione delle loro dimensioni, al fine di fornirne una versione più sintetica ed efficace [6, 7]. Un'importante considerazione riguardo i metodi *unsupervised* è che misurare quanto l'algoritmo sia performante non è sempre semplice, poiché le prestazioni sono spesso soggettive e *domain-specific* [7].

I metodi *unsupervised* più comunemente utilizzati includono il *clustering* e la riduzione della dimensionalità dei dati [6]. Il *clustering* consiste nel raggruppamento dei dati in modo tale che dati appartenenti allo stesso cluster condividano caratteristiche simili, mentre dati facenti parte di cluster differenti siano diversi secondo una data

metrica. Esistono numerosi approcci di *clustering*, tra i quali alcuni dei più comuni sono il *k-means clustering* e il *clustering* gerarchico (Figure 3 e 4). Il *k-means clustering* divide i dati in k gruppi basando la partizione sulla creazione di un centroide per ogni cluster, restituendo come output una funzione che assegna ogni dato ad uno dei cluster [8]. Il *clustering* gerarchico mira alla formazione di una gerarchia di cluster: gli algoritmi stimano quindi le partizioni dei dati con una procedura di ottimizzazione sequenziale. Esistono due diverse strategie: quella agglomerativa (*bottom-up*) e quella divisiva (*top-down*). L'approccio agglomerativo parte da una serie di cluster separati che vengono sequenzialmente appaiati sulla base della similarità, l'approccio divisivo procede invece nella direzione opposta [9].

La riduzione della dimensionalità (*dimensionality reduction*) consiste nella riduzione della complessità dei dati, mantenendone il più possibile la struttura, e si avvale di strumenti sia classici, come l'analisi delle componenti principali (*principal component analysis*, PCA) e la decomposizione ai valori singolari (*singular value decomposition*, SVD), che più moderni, quali la *uniform manifold approximation and projection* (UMAP) e la *topological data analysis* (TDA).

Gli approcci di *unsupervised learning* sono frequentemente utilizzati per pre-processare i dati, comprimendoli prima di fornirli a reti neurali o ad altri algoritmi di *supervised learning* come dati di input [7].

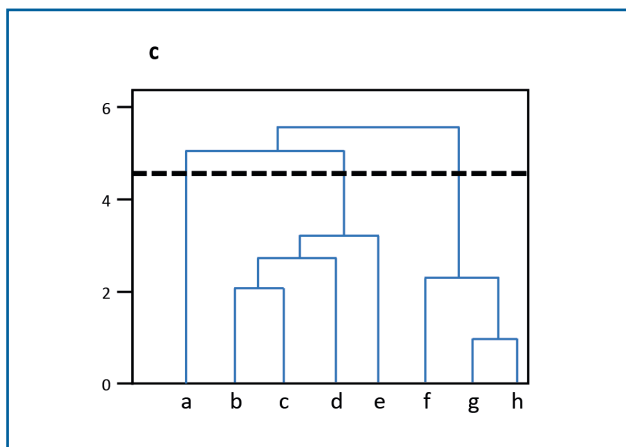


Figura 3 Clustering gerarchico.

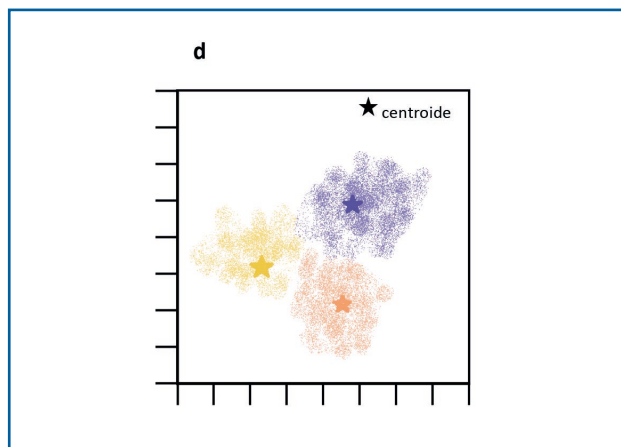


Figura 4 k-means clustering.

Machine learning e ricerca

L'avvento degli approcci omici ha generato una moltitudine di dati che devono essere integrati ed analizzati. Il machine learning, grazie ad approcci di supervised e unsupervised learning, si sta rivelando molto utile per rispondere a queste necessità.

La ricerca scientifica ha compiuto passi da gigante e volge ora lo sguardo verso la generazione e l'integrazione di un'enorme quantità di dati che permetterà un approccio olistico alla comprensione del funzionamento degli organismi viventi e dei meccanismi patologici. L'avvento delle moderne tecniche di sequenziamento a basso costo e la spettrometria di massa sono esempi emblematici della mole di dati che siamo ormai in grado di produrre; basti pensare che progetti come il *1000 Genomes* (progetto nato con lo scopo di fornire una descrizione completa delle varianti genetiche più comuni tramite sequenziamento *whole-genome* [10]) porteranno dati sulla scala dei petabyte (10^{15} byte). Non solo, lo sviluppo del sequenziamento di terza generazione e l'utilizzo di tecniche di single cell sequencing non potranno che incrementare ulteriormente la produzione dei cosiddetti *big data* [11]. In questo contesto è quindi fondamentale trovare la soluzione ai problemi di *storage* e di analisi che i *big data* portano inevitabilmente con loro. Come gestire quantità di dati così rilevanti? Come integrarli con dati provenienti per esempio da tecniche di *imaging*? Come analizzarli? Il ML è sicuramente un'arma potente per ovviare a questi problemi.

Uno dei campi in cui il ML si è rivelato di grande utilità è l'epigenomica, cioè lo studio delle modificazioni epigenetiche che contribuiscono alla regolazione dell'espressione genica. I ricercatori stanno raccogliendo numerosi dati tramite approcci omici (per esempio DNase-seq, FAIRE-seq, Hi-ChIP, ecc.) e il ML può essere un valido aiuto nell'interpretazione dei risultati e nell'integrazione di dati ottenuti da diverse metodiche [12].

Gli *epigenomic sequencing assay* si basano sulla frammentazione del DNA e terminano con il sequenziamento dei frammenti che presentano una determinata proprietà (come per esempio il legame ad una specifica proteina di interesse). Per fornire un'immagine completa dello stato epigenomico di ogni parte del genoma, è necessario combinare i risultati di numerosi saggi ed esperimenti. A questo scopo, sono stati sviluppati metodi di annotazione genomica semi-automatica (SAGA) che, basandosi su approcci di *unsupervised learning*, clusterizzano le regioni del genoma sulla base di similarità in termini di proprietà epigenomiche [13]. Lo studio di Malta et al. [14] è un esempio del grande potenziale degli algoritmi di ML. Partendo da dataset molecolari di cellule staminali e della loro progenie, sono stati generati due indici di staminalità utilizzati successivamente per valutare le caratteristiche epigenomiche e trascrittomiche di cellule cancerose (12000 tumori primari di 33 tipi di cancro diversi) con diverso grado di dedifferenziamento. Questo ha permesso di rilevare pattern di eterogeneità intra-tumore e potrebbe contribuire all'individuazione di nuovi target terapeutici [14].

Tramite ML è anche possibile predire quali sequenze di DNA possono costituire un sito di legame per fattori di trascrizione (superando le limitazioni dei saggi comunemente utilizzati in laboratorio), identificare possibili geni target dei microRNA e predire possibili effetti deleteri di polimorfismi [13, 15].

In maniera del tutto simile, il ML trova applicazione anche nella farmacologia computazionale [13]. Il ML si è infatti rivelato molto utile agli approcci di *image-based profiling* per il processo di *drug discovery*. Tali approcci si basano sul *high-throughput imaging*, un flusso di operazioni automatizzate che iniziano con l'incubazione delle cellule con agenti perturbanti (*small molecules*, siRNA, reagenti CRISPR/Cas9 [crRNA:trcrRNA], vettori plasmidici), a cui seguono l'acquisizione di immagini tramite microscopi automatizzati e, ad ultima, l'analisi delle immagini. Il fine di questi approcci consiste nel testare gli effetti degli agenti perturbanti di interesse sul fenotipo cellulare. Il *high throughput imaging* genera grandi raccolte di dati; le dimensioni e la complessità e la natura multiparametrica di tali dataset richiedono quindi strategie di analisi automatizzata. Per queste ragioni molti studi di *profiling* si affidano ad approcci di *unsupervised* e *supervised learning* per valutare cambiamenti fenotipici di singole cellule o di popolazioni cellulari. Per esempio, il *clustering* di profili *image-based*, per l'identificazione di molecole che condividono meccanismi di azione, si basa proprio su approcci di *unsupervised learning* [16, 17].

Machine learning e pratica clinica

Il machine learning sta acquisendo rilevanza anche in campo medico. Infatti, metodi di unsupervised learning permettono di stratificare i pazienti mentre tramite supervised learning è possibile individuare nuovi fattori di rischio per diverse patologie, migliorando notevolmente l'accuratezza diagnostica.

Il ML potrebbe rivelarsi molto promettente anche nell'ambito della pratica clinica, dove potrebbe essere di supporto all'attività decisionale dei medici [18]. Per esempio, algoritmi di ML potrebbero generare una stima del rischio di un paziente per uno specifico outcome, aiutando il medico a impostare terapie più consapevoli [19]. Nondimeno, approcci di ML potranno essere usati nella precision medicine per identificare i pazienti che potrebbero rispondere meglio ad una determinata terapia, indirizzando il medico verso la strategia terapeutica migliore [20]. Non solo, studi comparativi hanno mostrato che tramite ML è possibile identificare nuovi predittori di rischio, migliorando l'accuratezza diagnostica quando paragonati a metodi tradizionali [3, 19]. Per esempio, in uno studio prospettico di coorte condotto in UK [21], quattro algoritmi di ML (*random forest*, *logistic regression*, *gradient boosting machines*, *neural networks*) sono stati paragonati all'algoritmo raccomandato dall'*American Heart Association/American College of Cardiology* (ACC/AHA) per la predizione di un primo evento cardiovascolare entro 10 anni. I risultati sono molto interessanti. Il diabete, classicamente considerato di grande rilevanza per il rischio cardiovascolare, non era presente tra i fattori di rischio principali in nessuno dei quattro algoritmi di ML testati. Erano invece presenti come fattori di rischio altre condizioni quali COPD, malattie mentali severe, prescrizione di corticosteroidi orali, alti livelli dei trigliceridi. In aggiunta, i quattro algoritmi, in particolare le reti neurali, hanno mostrato capacità predittive superiori a quelle del modello ACC/AHA [21]. Un altro studio di particolare interesse è quello di Esteva et al. [22]. Tramite tecniche di deep learning è stato possibile classificare correttamente lesioni cutanee a partire da immagini cliniche, con performance paragonabili a quelle di dermatologi esperti. Questo metodo è implementabile su dispositivi mobili e ren-

denderebbe quindi accessibile ed economico tale sistema diagnostico.

Un altro campo medico in cui il ML si sta facendo strada è il *medical imaging*, grazie alla capacità degli algoritmi di riconoscere pattern che vanno oltre la percezione umana, superando le performance dei radiologi. Il ML sarà dunque un valido aiuto per i medici che potranno trovare un supporto nell'interpretazione delle immagini riducendo i tempi richiesti per una diagnosi [23-24]. Ne è esempio il lavoro di Schoepf et al., che ha mostrato l'utilità di una *computer- aided detection* (CAD) per la rilevazione di emboli polmonari segmentari e subsegmentari [25].

Gli algoritmi di ML potrebbero anche essere una risposta promettente all'annoso problema della variazione fenotipica. Lo studio dell'associazione tra genotipo e fenotipo si basa tipicamente sulla dicotomizzazione dei partecipanti in due categorie, caso o controllo, in funzione di determinati criteri clinici. Tuttavia questa suddivisione rischia di essere eccessivamente semplicistica poiché non tiene in considerazione l'eterogeneità fenotipica tra i soggetti. Questo si traduce in una riduzione del potere statistico e dell'abilità di individuare una vera associazione tra una malattia e un locus genico. Approcci di *unsupervised learning*, forti delle grandi raccolte di dati ora disponibili, permetterebbero di definire profili fenotipici in maniera più fine o di stratificare i pazienti in sottogruppi più omogenei, superando le limitazioni dovute alla variazione fenotipica [26].

Gli approcci di ML sembrano quindi essere destinati ad entrare a far parte della pratica clinica, dove contribuiranno, tra le altre cose, alla stratificazione dei pazienti, all'identificazione dei soggetti più responsivi a determinate terapie e alla definizione di nuovi indici di rischio.

Machine learning: criticità

L'applicazione del machine learning alla pratica clinica richiede la messa a punto di un percorso di approvazione che si adatti alle particolari caratteristiche dei software ML-based.

Le prestazioni del software dovrebbero essere robuste in tutta la popolazione di pazienti target, la sua implementazione dovrebbe garantire un uso corretto e le sue analisi o previsioni dovrebbero fornire un contesto che aiuti l'interpretazione [27].

Una prima importante considerazione sull'utilizzo degli approcci ML nella ricerca e, ancora di più, nella pratica clinica riguarda la qualità degli algoritmi che è dipendente da quella dei dati utilizzati in fase di *training*. Infatti, non è escluso incorrere in bias statistici. Questo può essere conseguenza per esempio di campioni subottimali, misure errate nei predittori ed eterogeneità degli effetti. Si pensi ad un algoritmo generato da dati ottenuti solo da soggetti caucasici: è intuitivo pensare che lo stesso algoritmo avrebbe capacità predittive ridotte se applicato a soggetti asiatici o africani. Dunque, bisogna sempre tener presente che un algoritmo potrebbe avere performance ridotte se applicato per esempio a soggetti appartenenti alle minoranze etniche (fairness) [28]. È quindi fondamentale mitigare questo tipo di bias, per esempio preferendo dati provenienti da studi randomizzati invece che da studi osservazionali, dati non dipendenti dal giudizio del medico, o utilizzando strumenti volti a valutare il rischio di bias nell'algoritmo [29]. I modelli di machine learning sono generalmente più performanti quando il training set è più ampio: un nodo fondamentale sarà quindi il bilanciamento tra le normative sulla privacy e il bisogno di accesso a dati quanto più vasti e diversificati [30].

Un secondo punto riguarda l'approvazione dei software *machine learning-based* concepiti per la pratica clinica. Quando messi a punto per l'ambito medico, tali software vengono definiti come *medical device* nell'ambito della normativa *Food, Drug, and Cosmetic Act*. Tra il 2017 e il 2018 la FDA ha approvato 14 *AI/ML-based software* come *device*. Nell'aprile 2019, la FDA ha annunciato una revisione degli approcci regolatori per i software ML: ciò si è reso necessario a causa di alcune caratteristiche intrinseche di questi software, che rendono la loro approvazione complicata rispetto ai *medical device* tradizionali. Infatti, gli algoritmi ML sono per loro stessa natura interattivi. Ciò significa che un algoritmo, col passare del tempo, potrebbe comportarsi in maniera differente rispetto al momento della sua approvazione. Inoltre, bisogna tenere in considerazione che la maggior parte dei software (11 su 14) è stata approvata tramite il 510(k) *pathway*, che si basa sulla dimostrazione dell'equivalenza sostanziale tra il nuovo software e quello già in uso. Questo creerebbe una catena di *device* sostanzialmente equivalenti tra loro ma che col passare degli anni risulterebbero essere in realtà molto

diversi dal *device* originale. Infine, è necessario un percorso di approvazione che testi le effettive efficacia e sicurezza dei *ML-based software* e non solo la loro equivalenza [31]. A tal scopo bisogna ricordare che, sebbene alcuni software siano stati approvati dalla FDA, questi tool presentano molteplici criticità anche per quanto riguarda la loro riproducibilità che sta subendo una vera e crisi [32] e ciò richiede un approccio moderno per il superamento di questi limiti [33].

Conclusioni

In conclusione, il ML ricoprirà un ruolo sempre più preponderante sia nella ricerca che nella pratica clinica. Le applicazioni sono molteplici e la comunità scientifica sta guardando con grande interesse ai progressi di questi approcci così potenti e innovativi. La raccolta di grandi quantità di dati, generati tramite approcci omici ma anche dall'esplosione dell'utilizzo di dispositivi mobili, *app* e condivisione di dati on-line, sta aprendo la strada ad una medicina caratterizzata da diagnosi più accurate [34], migliori stratificazioni dei pazienti, tempi di diagnosi ridotti. Proprio perché il ML si sta rivelando estremamente potente, è di cruciale importanza stabilire una regolamentazione adeguata alla particolare natura di questi approcci, così da garantirne l'efficacia e la sicurezza.

Bibliografia

- [1] Jakhar D, Kaur I. Artificial intelligence, machine learning and deep learning: definitions and differences. Clin. Exp. Dermatol, 2019; ced. 14029.
- [2] LeCun Y, Bengio Y, Hinton G. Deep learning. Nature, 2015; 521: 436-444.
- [3] Deo RC. Machine Learning in Medicine», Circulation, 2015; 132: 1920-1930.
- [4] Schrider DR, Kern AD. Supervised Machine Learning for Population Genetics: A New Paradigm. Trends Genet. TIG, 2018; 34: 301-312.
- [5] Angermueller C, Pärnamaa T, Parts L, Stegle O. Deep learning for computational biology. Mol. Syst. Biol, 2016; 12: 878.
- [6] Camacho DM, Collins KM, Powers RK, Costello JC, Collins JJ. Next-Generation Machine Learning for Biological Networks. Cell, 2018; 173: 1581-1592.
- [7] Maini V, Sabri S. Machine Learning for Humans. pag. 97.
- [8] Brusco M, Steinley D. Psychometrics: Combinatorial Data Analysis», in International Encyclopedia of the Social & Behavioral Sciences (Second Edition), J. D. Wright, A. c. di Oxford: Elsevier, 2015; 431-435.
- [9] Kimes PK, Liu Y, Neil Hayes D, Marron JS. Statistical significance for hierarchical clustering. Biometrics, 2017; 73: 811-821.
- [10] 1000 Genomes Project Consortium et al. A global reference for human genetic variation. Nature, 2015; 526: 68-74.
- [11] Schadt EE, Linderman MD, Sorenson J, Lee L, Nolan GP. Computational solutions to large-scale data management and analysis. Nat. Rev. Genet, 2010; 11: 647-657.
- [12] Xu C, Jackson SA. «Machine learning and complex biological data», Genome Biol, 2019; 20: 76, s13059-019-1689-0.
- [13] Zitnik M, Nguyen F, Wang B, Leskovec J, Goldenberg A, Hoffman MM. Machine learning for integrating data in biology and medicine: Principles, practice, and opportunities. Inf. Fusion, 2019; 50: 71-91.
- [14] Malta TM, et al. Machine Learning Identifies Stemness Features Associated with Oncogenic Dedifferentiation. Cell, 2018; 173: 338-354.e15.
- [15] Ching T, et al. Opportunities and obstacles for deep learning in biology and medicine. J. R. Soc. Interface, 2018; 15: 20170387.
- [16] Pegoraro G, Misteli T. High-throughput imaging for the discovery of cellular mechanisms of disease. Trends Genet. TIG, 2017; 33: 604-615.
- [17] Scheeder C, Heigwer F, Boutros M. Machine learning and image-based profiling in drug discovery. Curr. Opin. Syst. Biol., 2018; 10: 43-52.
- [18] Ascent of machine learning in medicine. Nat. Mater., 2019, 18:407.
- [19] Peterson ED. Machine Learning, Predictive Analytics, and Clinical Practice: Can the Past Inform the Present?. JAMA, 2019.
- [20] Guthrie NL, et al. Achieving Rapid Blood Pressure Control With Digital Therapeutics: Retrospective Cohort and Machine Learning Study. JMIR Cardio, 2019; 3: e13030.
- [21] Weng S., Reps J, Kai J, Garibaldi JM, Qureshi N. Can machine-learning improve cardiovascular risk prediction using routine clinical data?. PLoS One, 2017; 12: e0174944.
- [22] Esteva A., et al. Dermatologist-level classification of skin cancer with deep neural networks. Nat., 2017, 542: 115-118.
- [23] Erickson BJ, Korfiatis P, Akkus Z, Kline TL. Machine Learning for Medical Imaging. RadioGraphics, 2017; 37: 505-515.
- [24] Topol EJ., High-performance medicine: the convergence of human and artificial intelligence. Nat. Med., 2019, 25: 44-56
- [25] Schoepf UJ, Schneider AC, Das M, Wood SA, Cheema JI, Costello P. Pulmonary embolism: computer-aided detection at multidetector row spiral computed tomography. J. Thorac. Imaging, 2007; 22: 319-323.
- [26] Basile AO, Ritchie MD. Informatics and machine learning to define the phenotype. Expert Rev. Mol. Diagn., 2018; 18: 219-226.
- [27] Towards trustable machine learning. Nat. Biomed. Eng., 2018, 2: 709-710.
- [28] Kelly CJ., et al. Key challenges for delivering clinical impact with artificial intelligence. BMC Med., 2019, 17, 195.
- [29] R. B. Parikh, S. Teeple, e A. S. Navathe. Addressing Bias in Artificial Intelligence in Health Care. JAMA, 2019.
- [30] Rajkomar A., et al. Machine Learning in Medicine. N. Engl. J. Med., 2019; 380:1347-1358.
- [31] Hwang TJ, Kesselheim AS, Vokinger KN. Lifecycle Regulation of Artificial Intelligence - and Machine Learning-Based Software Devices in Medicine. JAMA, 2019.
- [32] Huston M. Artificial intelligence faces reproducibility crisis. Science, 2018, 359:725-726.
- [33] Beam AL., et al., Challenges to the Reproducibility of Machine Learning Models in Health Care. 2020, Published online
- [34] Pani L. L'innovazione sostenibile. pag. 148.