



**UNIMORE**

UNIVERSITÀ DEGLI STUDI DI  
MODENA E REGGIO EMILIA

Dipartimento di Scienze della Vita  
Corso di Laurea Magistrale a ciclo unico in Farmacia

**L'APPLICAZIONE DELL'INTELLIGENZA ARTIFICIALE  
NELLA SCOPERTA E NELLO SVILUPPO DEI FARMACI:  
NUOVE METODOLOGIE PER L'INNOVAZIONE  
TERAPEUTICA**

Relatore

**Prof. Fabio TASCEDDA**

*Correlatore*

**Dott. Luigi Francesco IANNONE**

*Tesi compilativa  
di Laurea Magistrale in Farmacia di:*  
**Silvia Colasanti**

**Anno Accademico  
2025/2026**

# INDICE

|                                                                                                                                             |           |
|---------------------------------------------------------------------------------------------------------------------------------------------|-----------|
| <b>ABSTRACT.....</b>                                                                                                                        | <b>1</b>  |
| <b>1. INTRODUZIONE.....</b>                                                                                                                 | <b>3</b>  |
| 1.1 Contesto e rilevanza del problema: limiti, costi e tempistiche del paradigma tradizionale di ricerca e sviluppo farmaceutico .....      | 3         |
| 1.2 L'IA come strumento trasformativo per il settore farmaceutico .....                                                                     | 5         |
| 1.3 Analisi farmacologica delle applicazioni dell'IA e valutazione critica delle prospettive future.....                                    | 6         |
| <b>2. FONDAMENTI DI INTELLIGENZA ARTIFICIALE PER IL SETTORE FARMACEUTICO.....</b>                                                           | <b>8</b>  |
| 2.1 IA, machine learning e deep learning .....                                                                                              | 8         |
| 2.2 Le metodologie del <i>machine learning</i> .....                                                                                        | 11        |
| 2.2.1 Apprendimento supervisionato .....                                                                                                    | 11        |
| 2.2.2 Apprendimento non supervisionato e approcci correlati .....                                                                           | 13        |
| 2.3 Reti neurali artificiali e deep learning .....                                                                                          | 13        |
| 2.3.1 Le principali architetture .....                                                                                                      | 14        |
| 2.3.2 Applicazioni chimico-farmaceutiche .....                                                                                              | 18        |
| 2.4 Fondamenti farmacologici rilevanti per l'IA .....                                                                                       | 20        |
| 2.5 <i>Big Data</i> e qualità dei dati.....                                                                                                 | 22        |
| <b>3. L'INTELLIGENZA ARTIFICIALE NEL PROCESSO DI RICERCA E SVILUPPO DI NUOVI FARMACI .....</b>                                              | <b>27</b> |
| 3.1 Identificazione e validazione dei target farmacologici .....                                                                            | 28        |
| 3.1.1 <i>Target validation</i> biologica e approcci computazionali.....                                                                     | 30        |
| 3.1.2 Reti di pathway cellulari e reti di interazione proteica .....                                                                        | 31        |
| 3.2 Progettazione razionale di farmaci assistita da IA ( <i>AI-assisted Drug Design</i> ) ...                                               | 33        |
| 3.2.1 <i>Structure-Based Drug Design</i> (SBDD): AlphaFold, Molecular Docking, Free Energy Perturbation .....                               | 34        |
| 3.2.2 <i>Ligand-Based Drug Design</i> (LBDD) e modelli QSAR ( <i>Quantitative Structure-Activity Relationship</i> ) potenziati dall'IA..... | 36        |
| 3.2.3 Ottimizzazione del lead compound assistita da IA.....                                                                                 | 38        |
| 3.3 Generazione di nuove molecole ( <i>De Novo Drug Design</i> ) .....                                                                      | 39        |
| 3.3.1 Ottimizzazione multiparametrica: ADMET, biodisponibilità, tossicità .....                                                             | 42        |
| 3.3.2 Modelli PK/PD integrati con IA e simulazioni in silico di efficacia terapeutica .....                                                 | 44        |
| 3.4 Ottimizzazione del processo di screening: virtual screening e supporto all'High-Throughput Screening (HTS).....                         | 47        |

|                                                                                                                                                               |           |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------|
| 3.5 <i>Drug Repurposing</i> : algoritmi per l'identificazione di nuove indicazioni terapeutiche per farmaci già approvati .....                               | 50        |
| <b>4. L'INTELLIGENZA ARTIFICIALE NEL PROCESSO DI DRUG DEVELOPMENT</b> .....                                                                                   | <b>54</b> |
| 4.1 Studi preclinici .....                                                                                                                                    | 55        |
| 4.1.1 Predizione delle proprietà ADMET: metabolismo epatico, interazioni farmacologiche e tossicità .....                                                     | 56        |
| 4.1.2 Sviluppo di modelli di tossicologia predittiva per ridurre i test su animali.....                                                                       | 59        |
| 4.2 Innovazione negli studi clinici .....                                                                                                                     | 62        |
| 4.2.1 Ottimizzazione del disegno dello studio, stratificazione dei pazienti, medicina personalizzata e farmacogenomica.....                                   | 64        |
| 4.2.2 Identificazione di biomarcatori digitali e molecolari .....                                                                                             | 68        |
| 4.2.3 Monitoraggio remoto ed analisi dei dati in tempo reale.....                                                                                             | 72        |
| 4.3 Farmacovigilanza e sicurezza post-marketing: utilizzo dell'IA per l'analisi di <i>Real-World Data</i> e la rapida identificazione di eventi avversi ..... | 75        |
| 4.3.1 L'analisi dei Real-World Data .....                                                                                                                     | 75        |
| 4.3.2 Identificazione degli eventi avversi e sistemi di segnalazione automatizzata....                                                                        | 76        |
| 4.3.3 Implicazioni per il farmacista clinico .....                                                                                                            | 77        |
| <b>5. ASPETTI REGOLATORI E STANDARD DI QUALITÀ NELLO SVILUPPO DI FARMACI BASATI SULL'IA</b> .....                                                             | <b>79</b> |
| 5.1 Il panorama regolatorio: le posizioni di FDA ed EMA .....                                                                                                 | 79        |
| 5.1.1 L'approccio della FDA: il framework di valutazione della credibilità.....                                                                               | 80        |
| 5.1.2 L'approccio dell'EMA: il <i>reflection paper</i> e la strategia di rete .....                                                                           | 82        |
| 5.1.3 Convergenze e differenze tra i due approcci .....                                                                                                       | 83        |
| 5.2 Principi guida per le “Good AI Practices”: i dieci principi comuni EMA-FDA                                                                                | 84        |
| 5.3 Le sfide regolatorie: validazione, trasparenza, data privacy e cybersecurity....                                                                          | 86        |
| 5.3.1 La validazione degli algoritmi.....                                                                                                                     | 86        |
| 5.3.2 La trasparenza e il problema della “ <i>black box</i> ” .....                                                                                           | 87        |
| 5.3.3 Il data drift e il mantenimento lungo il ciclo di vita .....                                                                                            | 87        |
| 5.3.4 Tutela dei dati personali e cybersecurity .....                                                                                                         | 88        |
| 5.3.5 La questione della responsabilità.....                                                                                                                  | 88        |
| <b>6. CASI DI STUDIO E APPLICAZIONI PRATICHE</b> .....                                                                                                        | <b>90</b> |
| 6.1 Aziende biotech “ <i>AI-native</i> ”: modelli di business e pipeline .....                                                                                | 90        |
| 6.2 L'IA nell'identificazione dei target: il caso dell'oncologia .....                                                                                        | 91        |
| 6.3 Farmaci scoperti con l'IA .....                                                                                                                           | 93        |
| 6.4 Collaborazioni tra industria farmaceutica e aziende tecnologiche.....                                                                                     | 94        |
| 6.5 Caso studio 1: inibitore di TNF per la fibrosi polmonare idiopatica.....                                                                                  | 94        |
| 6.6 Caso studio 2: il riposizionamento del Baricitinib nella COVID-19.....                                                                                    | 96        |

|                                                                                                 |            |
|-------------------------------------------------------------------------------------------------|------------|
| 6.7 L'IA a supporto delle decisioni terapeutiche: un esempio dalla diagnostica oncologica ..... | 97         |
| <b>7. SFIDE, PROSPETTIVE FUTURE E CONSIDERAZIONI ETICHE.....</b>                                | <b>98</b>  |
| 7.1 Sfide tecniche e scientifiche.....                                                          | 99         |
| 7.2 Prospettive future .....                                                                    | 103        |
| 7.3 Considerazioni etiche .....                                                                 | 106        |
| 7.4 Il ruolo del farmacista clinico nell'era dell'IA .....                                      | 108        |
| <b>8. CONCLUSIONI.....</b>                                                                      | <b>111</b> |
| <b>BIBLIOGRAFIA .....</b>                                                                       | <b>116</b> |

## ABSTRACT

Lo sviluppo di nuovi farmaci è un processo lungo, costoso e caratterizzato da elevati tassi di fallimento, che possono limitarne o ritardarne l'introduzione sul mercato. In questo contesto, l'intelligenza artificiale rappresenta uno strumento potenzialmente utile per aumentare l'efficienza del processo, contribuendo a ridurre tempi, costi e rischi lungo l'intero ciclo di vita del farmaco.

I recenti progressi nel *Machine Learning* e nel *Deep Learning* stanno mostrando un potenziale rilevante nell'affrontare alcune delle principali sfide dello sviluppo farmacologico. Algoritmi come le *Convolutional Neural Networks* (CNN) consentono di analizzare dati biologici complessi e di identificare potenziali nuovi bersagli terapeutici, con particolare impatto nelle fasi iniziali della scoperta farmacologica. Analogamente, il *Deep Learning* ha favorito la predizione della struttura tridimensionale delle proteine, come dimostrato da strumenti quali *AlphaFold*, contribuendo alla progettazione razionale di nuovi farmaci. I modelli generativi trovano applicazione nel *de novo drug design*, poiché permettono di generare nuove strutture molecolari. I modelli QSAR, invece, consentono di predire l'attività biologica e le proprietà ADMET dei composti con maggiore rapidità. Inoltre, le tecniche di *virtual screening* permettono di analizzare ampie librerie di composti mediante *fingerprint* molecolari, rappresentando un'alternativa in silico allo screening sperimentale ad alta resa, o *high-throughput screening* (HTS).

Nell'ambito della farmacocinetica e della farmacodinamica, i *Digital Twins* consentono di simulare interventi terapeutici prima della somministrazione, favorendo la personalizzazione del dosaggio, la previsione della variabilità interindividuale e l'ottimizzazione dei trial clinici. In questo ambito, l'uso combinato di Machine Learning e big data può migliorare la stratificazione dei pazienti in sottogruppi omogenei, supportando terapie più personalizzate e una previsione del rischio basata sul profilo individuale. Parallelamente, i biomarcatori digitali raccolti tramite sensori wearable consentono una valutazione non invasiva e un monitoraggio continuo in tempo reale.

L'analisi dei *Real World Data* permette di generare *Real World Evidence*, utile per valutare efficacia, sicurezza e impatto economico delle cure nel lungo periodo in popolazioni reali ed eterogenee. Queste evidenze assumono un ruolo crescente nei processi regolatori e nelle valutazioni post-marketing. L'intelligenza artificiale applicata alla farmacovigilanza consente inoltre di identificare, analizzare e prioritizzare le sospette reazioni avverse in modo più rapido ed efficiente. Tali applicazioni trovano concreta realizzazione nelle aziende

biotech AI-native, nella scoperta di molecole guidata dall'intelligenza artificiale e negli approcci di *drug repurposing* basati su network medicine e *graph neural networks*.

Le agenzie regolatorie, tra cui FDA ed EMA, stanno adottando un approccio proattivo e basato sul rischio, volto a definire principi condivisi per l'uso dell'intelligenza artificiale nel ciclo di vita dei medicinali. Restano tuttavia sfide tecniche, metodologiche ed etiche rilevanti, tra cui l'effetto black box, l'assenza di standard di validazione pienamente condivisi, la necessità di dati rappresentativi e di alta qualità, oltre ai problemi di bias, privacy, equità di accesso e responsabilità in caso di errore algoritmico. Nonostante tali criticità, le prospettive future indicano l'intelligenza artificiale come uno strumento sempre più centrale nella scoperta, nello sviluppo e nella produzione dei farmaci. L'integrazione tra l'intelligenza artificiale generativa, automazione dei laboratori e analisi avanzata dei dati potrà favorire una medicina più personalizzata e di precisione, ridefinendo anche il ruolo del farmacista, chiamato a interpretare gli output algoritmici e a integrare questi strumenti nella farmacovigilanza, nella gestione terapeutica e nel counseling del paziente.

# 1. INTRODUZIONE

La scoperta e lo sviluppo di un nuovo farmaco rappresentano uno dei processi più lunghi, complessi e incerti della ricerca biomedica. Nonostante i progressi della chimica farmaceutica, della biologia molecolare e della farmacologia, il percorso che porta una molecola dall'identificazione di un bersaglio terapeutico fino all'approvazione clinica continua a essere caratterizzato da tempi lunghissimi e da un tasso di fallimento molto elevato (1). In questo contesto, l'intelligenza artificiale (IA) si sta sviluppando negli ultimi anni come una tecnologia potenzialmente trasformativa, capace di intervenire in quasi tutte le fasi del processo.

La presente tesi si propone di analizzare, da una prospettiva farmacologica, il ruolo dell'IA nella ricerca e nello sviluppo dei farmaci: le sue applicazioni, le sue basi metodologiche, il quadro regolatorio che ne sta accompagnando l'adozione e le prospettive future, con particolare attenzione alle implicazioni per la figura del farmacista.

## 1.1 Contesto e rilevanza del problema: limiti, costi e tempistiche del paradigma tradizionale di ricerca e sviluppo farmaceutico

Il modello tradizionale di ricerca e sviluppo (R&S) farmaceutico segue una sequenza ben definita: identificazione e validazione di un bersaglio biologico, individuazione di composti capostipite, loro ottimizzazione fino al candidato clinico, sperimentazione preclinica e, infine, le tre fasi della sperimentazione clinica sull'uomo che precedono la domanda di autorizzazione agli enti regolatori e l'immissione in commercio. Si tratta di un percorso dalla durata in genere di più dieci anni e che richiede risorse enormi. Le stime del costo medio per ogni nuovo farmaco approvato variano sensibilmente a seconda dei dati e dei metodi impiegati. DiMasi et al. (2), utilizzando dati aziendali riservati, lo collocano intorno ai 2.6 miliardi di dollari, mentre Wouters et al. (3), basandosi esclusivamente su dati finanziari pubblici, riportano una stima nettamente inferiore con una media di circa 1.6 miliardi di dollari (1). Gran parte di questi costi deriva dall'alta percentuale di candidati farmacologici che non arrivano all'approvazione e commercializzazione (1).

Il dato più critico del modello tradizionale è infatti il tasso di fallimento. Analizzando un'ampia base di dati relativa a oltre ventimila composti, Wong et al. (4), hanno stimato una probabilità complessiva di successo, dall'ingresso in Fase I fino all'approvazione, intorno al

14%: circa l'86-90% dei programmi di sviluppo non raggiunge il mercato. Non si tratta di un fenomeno recente, ma di una caratteristica strutturale e di lungo periodo del settore; già Kola e Landis (5) stimavano una probabilità complessiva di successo dalla prima somministrazione nell'uomo alla registrazione intorno all'11%, ovvero circa un composto su nove. La fase più critica è la Fase II, che presenta la probabilità di superamento più bassa (intorno al 30%, contro circa il 63% della Fase I e il 58% della Fase III), poiché è in questa fase che si valuta per la prima volta, su un numero ampio di pazienti, se il farmaco è realmente efficace (ed anche ben tollerato in una più ampia casistica). Superare le fasi iniziali, dedicate soprattutto alla sicurezza e alla tollerabilità, non garantisce dunque che la molecola sia poi capace di produrre il beneficio terapeutico atteso. Si tratta di un dato di grande significato farmacologico, perché mostra che il vero collo di bottiglia non è tanto la tossicità, quanto la capacità di individuare bersagli e molecole realmente attivi sulla patologia umana (4,5).

Le ragioni di questo fenomeno non sono soltanto organizzative, ma dipendono da limiti scientifici e biologici intrinseci al modello tradizionale. In primo luogo, la validazione dei bersagli rimane incerta: un bersaglio molecolare può apparire promettente sulla base di evidenze precliniche e rivelarsi poi non realmente determinante nella patologia umana. L'analisi retrospettiva dei progetti di sviluppo di un'importante azienda farmaceutica condotta da Cook et al. (6) quantifica questo limite, mostrando che in una quota rilevante dei fallimenti per inefficacia mancava una chiara dimostrazione del legame tra il bersaglio e la malattia, o un modello sperimentale capace di prevedere in modo affidabile la risposta nell'uomo. In secondo luogo, lo spazio chimico esplorabile è immenso, mentre la capacità di sintetizzare e saggiare sperimentalmente i composti è limitata, il che rende la ricerca del capostipite un processo dispendioso e in parte affidato alla casualità. In terzo luogo, la traslazione dai modelli preclinici all'uomo è spesso problematica: le differenze tra le specie nella farmacocinetica e nella farmacodinamica fanno sì che risultati promettenti negli animali non si confermino nei pazienti (6). Infine, molte molecole vengono abbandonate tardivamente per problemi di tossicità o per proprietà ADMET (ossia assorbimento, distribuzione, metabolismo, eliminazione e tossicità) sfavorevoli, emerse solo nelle fasi avanzate (1). In un'analisi congiunta condotta su dati di più grandi aziende farmaceutiche, Waring et al. (7) hanno individuato proprio nella tossicità rilevata negli studi preclinici (in vitro e sugli animali) la principale causa di abbandono dei composti, a conferma del peso che le proprietà di sicurezza dei composti continuano ad avere sull'esito dello sviluppo.

L'insieme di questi fattori ha prodotto, nel corso degli ultimi decenni, una progressiva diminuzione del rendimento della ricerca farmaceutica, nonostante l'aumento delle risorse investite, il numero di nuovi farmaci approvati non è cresciuto in egual misura (1). Questo andamento è stato descritto da Scannell et al. (8) come "*legge di Eroom*", secondo cui emerge come il numero di nuovi farmaci a parità di investimento si è progressivamente e drasticamente ridotto dal 1950, con una riduzione di circa 80 volte in termini reali. È anche in risposta a questa sfida, scientifica ed economica, che il settore ha iniziato a guardare con crescente interesse a strumenti capaci di rendere più razionali, rapide e predittive le fasi iniziali del processo (8).

## **1.2 L'IA come strumento trasformativo per il settore farmaceutico**

L'impiego di strumenti computazionali nella ricerca farmaceutica non è una novità assoluta. Già con l'affermarsi della bioinformatica e della chemo-informatica, metodi come le relazioni quantitative struttura-attività (QSAR), il *docking* molecolare e la modellistica farmacoforica avevano introdotto un approccio razionale alla progettazione dei farmaci. La novità degli ultimi anni risiede nel passaggio da questi metodi, in gran parte basati su regole definite a priori, a tecniche di apprendimento automatico (*machine learning*, ML) capaci di apprendere direttamente dai dati e di individuare relazioni complesse e non lineari difficilmente formalizzabili in modo esplicito (9).

In una review di riferimento, Vamathevan et al (9) hanno mostrato come il ML possa intervenire in tutte le fasi della *pipeline*, dall'identificazione e validazione dei bersagli alla scoperta di biomarcatori prognostici, fino all'analisi delle immagini di patologia digitale negli studi clinici, favorendo un processo decisionale guidato dai dati. La condizione perché questi metodi siano efficaci è però precisa: la disponibilità di dati abbondanti, di alta qualità e ben caratterizzati (9), un requisito che, come si vedrà, costituisce al tempo stesso il principale punto di forza e il principale limite dell'approccio.

Il salto qualitativo più recente è legato al cosiddetto apprendimento profondo (*deep learning*) e all'IA generativa. Wang et al (10) hanno descritto come tecniche quali l'apprendimento auto-supervisionato e il *deep learning* geometrico, che sfrutta la conoscenza della struttura tridimensionale dei dati biologici, abbiano esteso enormemente le capacità predittive nei diversi ambiti della scoperta scientifica, consentendo ad esempio di generare nuove molecole o di prevedere la struttura delle proteine (10). Per la farmacologia, questo significa poter

modellare con precisione le interazioni farmaco-bersaglio, prevedere le proprietà ADMET e di tossicità, stimare l'affinità di legame ed esplorare in silico porzioni dello spazio chimico altrimenti inaccessibili sperimentalmente (9,11).

Una sintesi aggiornata è riportata in Zhang et al., (11), che ripercorrono lo stato dell'arte delle applicazioni dell'IA lungo l'intero sviluppo dei farmaci a piccole molecole: dall'identificazione del bersaglio alla sintesi chimica, fino al disegno e alla conduzione degli studi clinici, sottolineando il ruolo emergente dei grandi modelli linguistici e dell'IA generativa. Ne emerge un quadro in cui l'IA non si limita ad accelerare singoli passaggi, ma ridefinisce il modo stesso in cui si imposta la ricerca farmacologica, rendendola più predittiva e meno dipendente dall'approccio per tentativi ed errori (11).

### **1.3 Analisi farmacologica delle applicazioni dell'IA e valutazione critica delle prospettive future**

Se il potenziale dell'IA è indiscutibile, una valutazione rigorosa del reale impatto dell'IA impone di distinguere la tecnologia già esistente da sviluppi futuri. L'evidenza disponibile indica che l'IA è particolarmente efficace nelle fasi precoci della ricerca: è in grado di comprimere in modo sostanziale i tempi dell'identificazione dei bersagli e dell'ottimizzazione del capostipite, portando alla selezione di una molecola candidata in tempi nettamente inferiori rispetto a quelli del percorso tradizionale (11).

Un esempio concreto e particolarmente significativo è rappresentato dal rentosertib, un inibitore di TNK *first-in-class*, cioè il primo nella sua categoria a sfruttare questo meccanismo d'azione per il trattamento della fibrosi polmonare idiopatica, in cui sia il bersaglio sia la molecola sono stati individuati mediante IA generativa: la molecola ha completato uno studio clinico di Fase IIa con risultati incoraggianti sulla funzione polmonare (11). Si tratta di uno dei primi casi in cui un farmaco con bersaglio e struttura interamente generati dall'IA raggiunge questo stadio di sperimentazione; tuttavia, lo studio ha coinvolto un numero limitato di pazienti e i risultati preliminari di efficacia, seppur promettenti, necessitano di conferma in coorti più ampie. L'IA ha perciò dimostrato di saper accelerare la scoperta di nuovi farmaci, ma la prova che questi farmaci funzionino davvero nei pazienti deve ancora essere effettuata.

Proprio qui si colloca il principale limite, di natura propriamente farmacologica: l'accuratezza tecnica di un modello non garantisce la sua validità clinica. Un bersaglio predetto con precisione non è necessariamente targettabile con un farmaco, e la capacità di accelerare le fasi iniziali non si è finora tradotta in un miglioramento dimostrabile dei tassi di successo clinico, che restano governati dalla complessità della biologia umana (4). La validazione sperimentale e clinica rimane perciò insostituibile e ad essa si affiancano altre questioni aperte, come la qualità e la rappresentatività dei dati, l'interpretabilità dei modelli (il cosiddetto problema della "scatola nera"), la necessità di un quadro regolatorio adeguato che le principali autorità, come la *Food and Drug Administration* (FDA) e l'*European Medicines Agency* (EMA), hanno iniziato ad affrontare solo di recente (12,13).

In questa prospettiva si inserisce anche il ruolo del farmacista, un professionista che sarà coinvolto sempre più spesso ad interpretare e validare gli output di strumenti basati sull'IA, sia nella ricerca sia nella pratica clinica e nella farmacovigilanza.

## 2. FONDAMENTI DI INTELLIGENZA ARTIFICIALE PER IL SETTORE FARMACEUTICO

L'applicazione dell'IA alla ricerca e allo sviluppo farmaceutico richiede, come premessa, la comprensione di alcuni concetti fondamentali appartenenti a due ambiti distinti ma complementari: da un lato le metodologie computazionali su cui si basa l'IA, dall'altro i principi farmacologici che descrivono il modo in cui un farmaco agisce nell'organismo. Questo capitolo introduce tali fondamenti, con l'obiettivo di fornire gli strumenti concettuali necessari per interpretare le applicazioni discusse nei capitoli successivi. Verranno in primo luogo definiti i rapporti tra IA, apprendimento automatico ed apprendimento profondo; verranno descritte quindi le principali metodologie di apprendimento e le architetture di reti neurali con le relative applicazioni in ambito chimico-farmaceutico; verranno poi discussi i concetti farmacologici di base, che costituiscono le grandezze che i modelli computazionali cercano di prevedere; infine si affronterà il tema dei Big Data, della loro qualità e gestione, che rappresenta al tempo stesso il presupposto e uno dei principali limiti dell'intero settore.

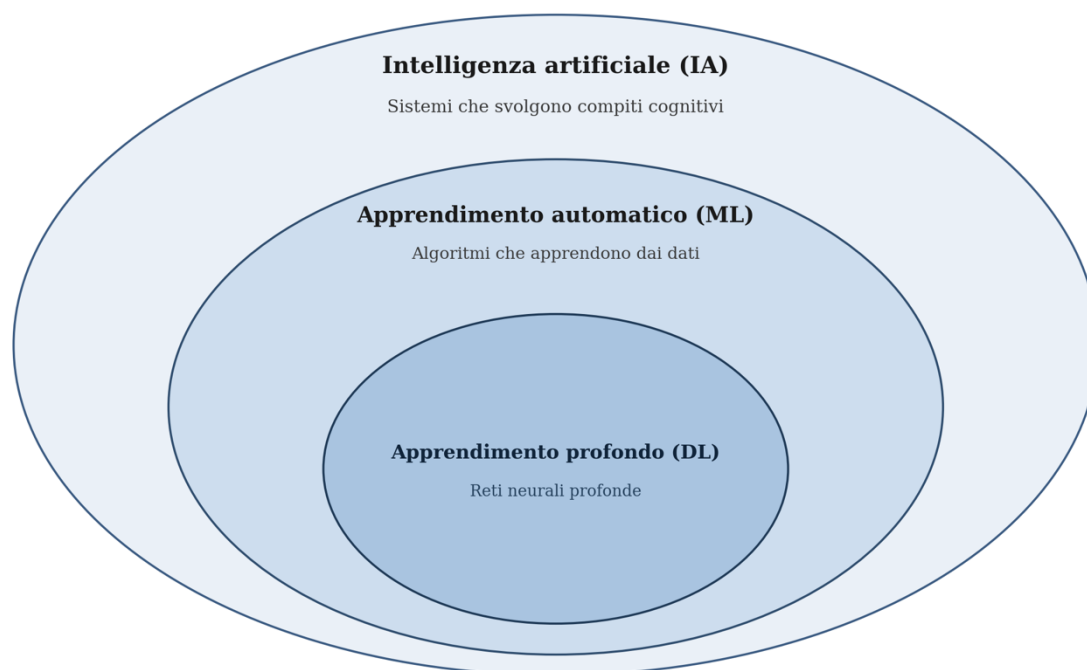
### 2.1 IA, machine learning e deep learning

Con l'espressione IA si indica un ampio campo multidisciplinare che si occupa dello sviluppo di sistemi capaci di svolgere compiti che normalmente richiederebbero l'intelligenza umana, come il riconoscimento del linguaggio, l'interpretazione di immagini, il processo decisionale e l'apprendimento dall'esperienza (14). All'interno di questo campo si colloca il *machine learning* (ML), o apprendimento automatico, un sottoinsieme dell'IA che riunisce algoritmi in grado di apprendere dai dati: a differenza della programmazione tradizionale, in cui è il programmatore a definire passo per passo le regole che il sistema deve seguire, nel machine learning le regole non vengono fornite dall'esterno ma vengono ricavate dal modello stesso a partire dagli esempi. Il modello analizza un insieme di dati e ne estrae relazioni e schemi ricorrenti, cioè le relazioni che legano le caratteristiche di ciascun esempio al risultato osservato; tali relazioni vengono poi applicate a dati nuovi per formulare previsioni (14).

In questo contesto, un modello è una rappresentazione matematica di un sistema o di un processo, definita da un insieme di parametri che vengono regolati durante una fase di addestramento (*training*) sulla base dei dati disponibili (14). Per valutare in modo affidabile la capacità di un modello di generalizzare a dati mai osservati, l'insieme dei dati viene

generalmente suddiviso in una porzione di addestramento, usata per costruire il modello, e una porzione di test, impiegata esclusivamente per verificarne le prestazioni (14). Questa suddivisione (di solito 80:20 di tutti i dati) risponde a un'esigenza precisa: se le prestazioni del modello venissero misurate sugli stessi dati impiegati per costruirlo, il risultato sarebbe ingannevolmente positivo, perché il modello potrebbe essersi limitato a "memorizzare" gli esempi visti anziché coglierne le relazioni generali. La porzione di test, tenuta da parte e mai utilizzata durante l'addestramento, simula invece dati nuovi e permette di stimare in modo realistico come il modello si comporterà su casi mai incontrati prima, ossia la sua capacità di generalizzare (14).

Il deep learning (DL), o apprendimento profondo, costituisce a sua volta un sottoinsieme del machine learning, basato su reti neurali artificiali dotate di numerosi strati intermedi (deep neural networks). In una rete neurale i dati attraversano una successione di strati di unità di calcolo: ogni strato trasforma l'informazione ricevuta dal precedente e la trasmette al successivo. L'aggettivo "profondo" (deep) si riferisce proprio all'elevato numero di questi strati intermedi, ed è da questa profondità che deriva la capacità del modello di apprendere automaticamente rappresentazioni dei dati a complessità crescente, senza che le caratteristiche rilevanti debbano essere selezionate manualmente da un esperto. Negli approcci tradizionali di machine learning, l'analisi richiede generalmente una fase preliminare di selezione, calcolo o trasformazione delle caratteristiche rilevanti per il problema, come specifici descrittori molecolari. Questa fase, nota come feature engineering, può essere guidata dall'esperienza dell'operatore, da procedure computazionali o da una combinazione di entrambe. Rispetto a questi approcci, il deep learning consente in molti casi di apprendere rappresentazioni più direttamente dai dati, riducendo la dipendenza dalla definizione manuale delle caratteristiche. Infatti, sono gli stessi strati della rete a costruire progressivamente queste caratteristiche a partire dai dati grezzi: gli strati iniziali apprendono caratteristiche semplici, quelli successivi le combinano in rappresentazioni via via più astratte e complesse. Questa capacità di apprendere automaticamente le rappresentazioni utili è all'origine dei risultati ottenuti dall'IA in ambiti come la visione artificiale e l'elaborazione del linguaggio naturale (14). Si delinea così una relazione di inclusione progressiva, (**Figura 1**) in cui il deep learning è una specializzazione del machine learning, che a sua volta è una branca dell'IA (14).



**Figura 1.** Relazione di inclusione tra IA, apprendimento automatico (machine learning) e apprendimento profondo (deep learning). L'apprendimento profondo costituisce un sottoinsieme dell'apprendimento automatico, che a sua volta rappresenta una branca dell'IA. Elaborazione propria, adattata da (14).

Tra i modelli di machine learning è utile distinguere due grandi categorie. I *modelli discriminativi* apprendono direttamente la relazione tra i dati in ingresso e una variabile in uscita, e vengono impiegati in compiti di classificazione o regressione. I *modelli generativi*, invece, non si limitano a separare le categorie ma apprendono come sono fatti i dati nel loro complesso, ovvero la distribuzione di probabilità da cui gli esempi di addestramento sembrano provenire; una volta appresa tale distribuzione, possono campionarne nuovi dati, simili agli originali ma non identici. È questa capacità che consente loro di generare, ad esempio, nuove strutture molecolari (14). Questa distinzione è particolarmente rilevante in ambito farmaceutico, poiché i modelli generativi sono alla base della progettazione di nuove molecole, di cui si discuterà più avanti.

Negli ultimi anni queste tecnologie hanno trovato applicazione crescente nella scoperta e nello sviluppo dei farmaci, dove vengono impiegate per estrarre informazioni utili da grandi quantità di dati biologici e chimici per supportare le decisioni lungo le diverse fasi del processo (11,15,16).

## 2.2 Le metodologie del *machine learning*

Gli algoritmi di machine learning possono essere classificati in base al tipo di dati su cui vengono addestrati e al modo in cui apprendono. Le due categorie principali sono l'*apprendimento supervisionato* e l'*apprendimento non supervisionato*; a queste si affiancano approcci intermedi, come l'apprendimento semi-supervisionato e auto-supervisionato, che hanno acquisito crescente importanza nel settore farmaceutico (17).

### 2.2.1 Apprendimento supervisionato

Nell'apprendimento supervisionato il modello viene addestrato su un insieme di dati etichettati, in cui a ogni esempio in ingresso è associato un valore in uscita noto detto etichetta. In ambito farmaceutico, ad esempio, l'ingresso può essere costituito dalla descrizione di una molecola e l'etichetta dall'informazione, verificata sperimentalmente, che quella molecola sia tossica oppure non tossica. È a partire da questi accoppiamenti già noti tra ingresso ed etichetta che il modello apprende la relazione da applicare poi a molecole nuove, di cui l'etichetta non è disponibile. All'interno dell'apprendimento supervisionato si distinguono due compiti principali, a seconda della natura della risposta che si vuole prevedere: la classificazione e la regressione. Nella classificazione il modello assegna ciascun esempio a una tra un numero finito di categorie distinte: la risposta è qualitativa, come nel caso della distinzione tra un composto tossico e uno non tossico. Nella regressione, invece, il modello prevede un valore numerico continuo, che può assumere qualsiasi valore entro un intervallo, come la concentrazione di un farmaco capace di produrre un determinato effetto o l'affinità di legame di una molecola per il proprio bersaglio (17).

Numerosi algoritmi rientrano in questa categoria. Le *support vector machine* (SVM) cercano di tracciare un confine che separi nel modo più netto possibile le diverse classi di composti, ad esempio quelli attivi da quelli inattivi. Ogni composto viene rappresentato come un punto le cui coordinate corrispondono ai suoi descrittori molecolari; poiché i descrittori sono numerosi, questi punti si collocano in uno spazio a molte dimensioni, e il confine che li separa assume la forma di una superficie detta iperpiano. L'algoritmo individua, tra tutti i confini possibili, quello che lascia il margine più ampio tra le classi, così da distinguerle nel modo più affidabile anche su composti nuovi (17). Gli alberi decisionali (*decision tree*) classificano i dati attraverso una sequenza di domande successive, ciascuna delle quali divide i composti in base al valore di un descrittore (ad esempio, se una certa proprietà sia al di

sopra o al di sotto di una soglia). Il loro vantaggio è l'interpretabilità: il percorso di decisioni che porta a una previsione può essere ripercorso e letto come una sequenza di regole comprensibili, rendendo trasparente e quindi facilmente interpretabile il motivo di ciascuna classificazione; Il *random forest* combina invece numerosi alberi decisionali, ciascuno costruito su una porzione diversa dei dati: la previsione finale non è affidata a un singolo albero, ma è il risultato di una sorta di votazione tra tutti gli alberi, che concorrono alla decisione (la previsione "per consenso"). Poiché i singoli alberi tendono a commettere errori diversi tra loro, mediane i risultati attenua gli errori dei singoli e riduce il rischio di sovra-adattamento (*overfitting*), cioè la situazione in cui un modello, anziché cogliere soltanto le relazioni generali presenti nei dati di addestramento, ne apprende anche le fluttuazioni casuali e irripetibili; ne risulta un modello molto accurato sugli esempi già visti ma poco affidabile sui dati nuovi, perché ha perso capacità di generalizzare; ne deriva una previsione complessivamente più accurata e stabile (17). L'algoritmo *k-nearest neighbors* (kNN) classifica un composto in base alla classe dei composti a esso più vicini nello spazio dei descrittori, mentre i metodi di apprendimento d'insieme (*ensemble learning*) combinano più modelli individuali per ottenere prestazioni superiori a quelle di ciascun modello singolo (17).

Molti di questi algoritmi confluiscono in quella che è forse l'applicazione più classica dell'apprendimento supervisionato in ambito farmaceutico: i modelli di relazione quantitativa struttura-attività (*Quantitative Structure-Activity Relationship*, QSAR). Questi modelli mettono in relazione i descrittori molecolari di un composto, cioè grandezze numeriche che ne riassumono le caratteristiche strutturali e chimico-fisiche, con una sua proprietà o attività biologica misurata sperimentalmente. Una volta addestrato su composti di attività nota, un modello QSAR può stimare l'attività di nuove molecole sulla sola base della loro struttura, riducendo il ricorso a saggi sperimentali; si tratta dello stesso principio che, come si vedrà, il *deep learning* porta oggi a un livello di complessità molto superiore (17).

La valutazione delle prestazioni di un modello supervisionato richiede particolare attenzione. Tra i metodi più diffusi vi sono la validazione su un insieme di test indipendente (*holdout*), la validazione incrociata (*cross-validation*), in cui i dati vengono suddivisi ripetutamente in sottoinsiemi usati a turno per l'addestramento e la verifica, e la validazione esterna, condotta su un insieme di dati del tutto distinto da quello di addestramento. La validazione interna tende generalmente a sovrastimare le prestazioni reali del modello. Ciò accade perché nella

validazione interna i dati impiegati per verificare il modello provengono dallo stesso insieme usato per costruirlo e ne condividono quindi le caratteristiche, comprese eventuali distorsioni sistematiche; il modello risulta perciò avvantaggiato rispetto a quanto accadrebbe su dati realmente nuovi. La validazione esterna, condotta su un insieme di dati di origine del tutto indipendente, riproduce in modo più fedele le condizioni d'uso reali ed è per questo considerata la più rigorosa (17).

### **2.2.2 Apprendimento non supervisionato e approcci correlati**

Nell'apprendimento non supervisionato il modello viene addestrato su dati privi di etichette, con l'obiettivo di individuare schemi nascosti e raggruppamenti (*cluster*) all'interno dei dati (17). Questo approccio è utile, ad esempio, per identificare sottogruppi di pazienti o di composti con caratteristiche simili senza disporre di una classificazione predefinita.

Tra le due categorie principali si collocano alcuni approcci intermedi. L'apprendimento *semi-supervisionato* utilizza contemporaneamente dati etichettati e non etichettati, mentre l'apprendimento *auto-supervisionato* (*self-supervised learning*) ricava le etichette direttamente dai dati non etichettati, senza intervento umano. L'idea è far apprendere al modello un compito ausiliario in cui la risposta corretta è già contenuta nei dati stessi: ad esempio, nascondere una parte di una molecola o di una sequenza biologica e chiedere al modello di ricostruirla, oppure oscurare alcuni elementi di una sequenza e farli prevedere a partire dal contesto. In questo modo il modello impara rappresentazioni utili dei dati senza bisogno di annotazioni fornite da un esperto, riducendo la necessità di disporre di grandi quantità di esempi etichettati manualmente (17,18). Questi metodi risultano particolarmente vantaggiosi in ambito farmaceutico, dove i dati disponibili sono spesso abbondanti per quantità ma poveri di annotazioni affidabili: si dispone cioè di moltissime molecole o sequenze, ma per relativamente poche di esse è nota in modo certo e verificato sperimentalmente la proprietà di interesse (ad esempio l'attività o la tossicità). Gli approcci che sfruttano anche i dati non etichettati permettono di valorizzare questo patrimonio altrimenti in gran parte inutilizzato.

## **2.3 Reti neurali artificiali e deep learning**

Le reti neurali artificiali sono modelli di apprendimento automatico ispirati alla struttura del cervello umano, costituiti da unità elementari (neuroni) organizzate in strati e collegate tra loro. Tra lo strato di ingresso e quello di uscita si collocano uno o più strati intermedi, detti

strati nascosti, che effettuano trasformazioni non lineari dei dati. Ciò significa che ogni strato non si limita a sommare o riscalare proporzionalmente i valori ricevuti, ma li combina secondo relazioni più complesse; è proprio questa non linearità a permettere alla rete di rappresentare relazioni intricate tra i dati, che una semplice combinazione proporzionale (lineare) non riuscirebbe a descrivere. Quando una rete possiede molti strati nascosti si parla di rete neurale profonda, alla base del deep learning: maggiore è il numero di strati, maggiore è la capacità della rete di apprendere caratteristiche complesse direttamente dai dati, senza necessità di una selezione manuale. Questo accade perché ogni strato costruisce le proprie rappresentazioni a partire da quelle elaborate dallo strato precedente: i primi strati colgono caratteristiche elementari, gli strati successivi le combinano in caratteristiche progressivamente più astratte e articolate. Aumentando il numero di strati si allunga questa catena di combinazioni, e la rete può così rappresentare relazioni sempre più complesse a partire direttamente dai dati grezzi, senza che sia un esperto a indicare in anticipo quali caratteristiche considerare (14,17).

### **2.3.1 Le principali architetture**

A seconda del tipo di dato da elaborare e del compito da svolgere, sono state sviluppate diverse architetture di reti neurali, alcune delle quali si sono affermate come riferimento nella progettazione dei farmaci (19).

La forma più semplice è il *perceptron* multistrato (*multilayer perceptron*, MLP), in cui ogni unità di uno strato è connessa a tutte le unità dello strato successivo; si tratta dell'architettura di base, da cui derivano per specializzazione le reti descritte di seguito, ciascuna adattata a un particolare tipo di dato (19). Le reti neurali convoluzionali (*convolutional neural networks*, CNN) sono particolarmente efficaci nell'elaborazione di dati con struttura a griglia, come le immagini. Attraverso operazioni di convoluzione, queste reti fanno scorrere piccoli filtri sull'immagine alla ricerca di motivi ricorrenti, individuando dapprima elementi semplici come bordi e contorni; gli strati di aggregazione (*pooling*) riassumono via via l'informazione, riducendone il dettaglio e trattenendo ciò che è essenziale. Strato dopo strato, la rete combina questi elementi semplici in motivi sempre più complessi, passando dai dettagli locali agli schemi d'insieme: è in questo senso che si dice che apprenda le caratteristiche in modo gerarchico. Due accorgimenti rendono inoltre queste reti computazionalmente efficienti. Da un lato, lo stesso filtro viene riutilizzato su tutta l'immagine (pesi condivisi): poiché un motivo come un bordo può comparire ovunque, non

serve apprendere separatamente per ogni posizione. Dall'altro, ciascuna unità è collegata solo a una piccola regione dell'immagine e non all'intera immagine (connettività sparsa). Entrambi gli accorgimenti riducono drasticamente il numero di parametri che la rete deve apprendere, e quindi la quantità di calcolo e di dati necessari (17,19).

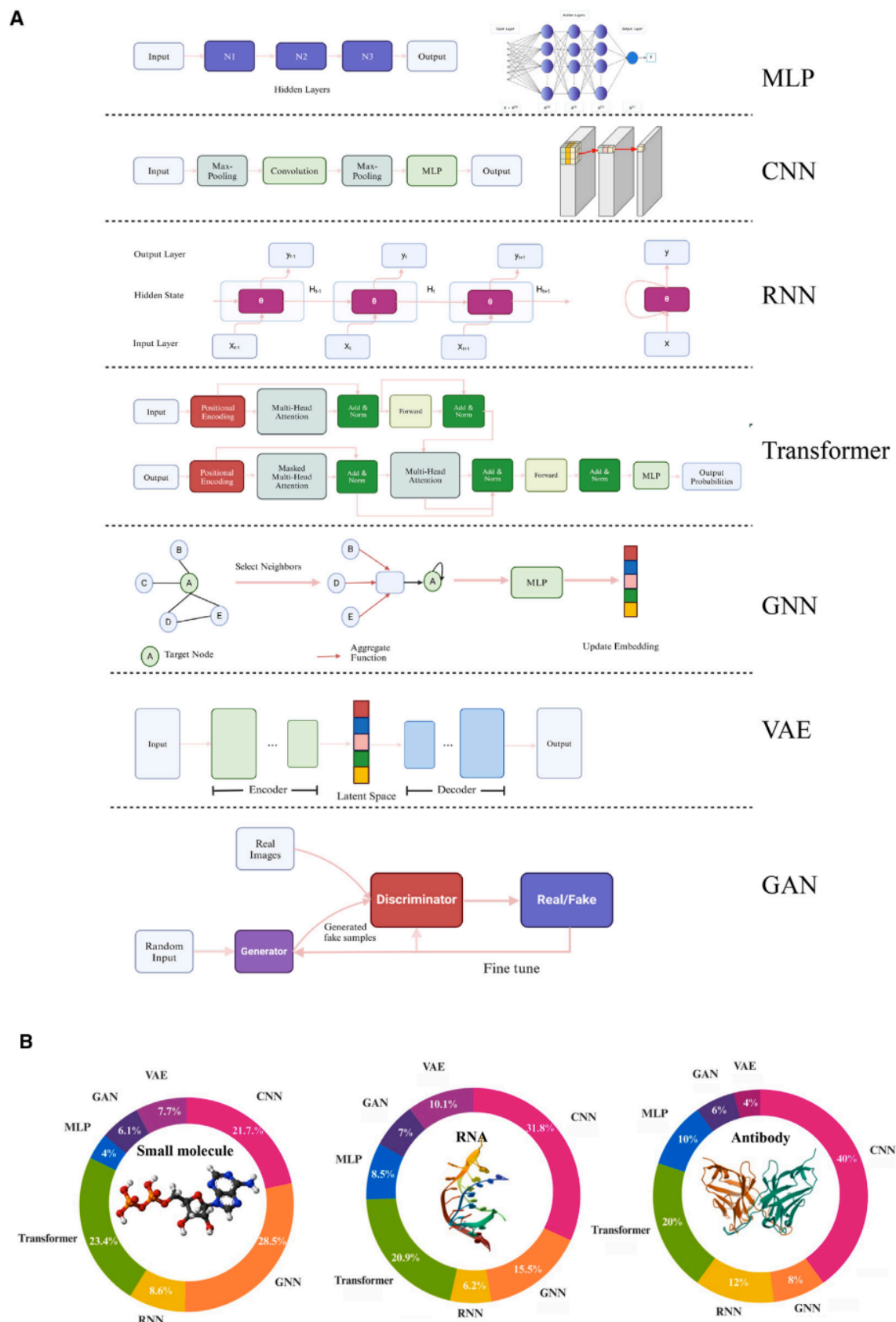
Le reti neurali ricorrenti (*recurrent neural networks*, RNN) sono progettate per elaborare dati sequenziali, come testi o sequenze biologiche, e mantengono una sorta di memoria degli ingressi precedenti. Le RNN tradizionali incontrano tuttavia difficoltà nel gestire sequenze lunghe, ciò dipende dal modo in cui elaborano le sequenze: le RNN procedono un elemento alla volta, aggiornando a ogni passo una memoria interna che riassume quanto visto fino a quel momento. Su sequenze lunghe, però, l'informazione proveniente dagli elementi iniziali tende a diluirsi progressivamente man mano che nuovi elementi vengono elaborati, fino quasi a svanire; la rete fatica così a cogliere relazioni tra parti molto distanti della sequenza. Per attenuare questo limite sono state sviluppate varianti come le reti *long short-term memory* (LSTM) e *gated recurrent unit* (GRU), dotate di meccanismi che permettono di trattenere più a lungo le informazioni rilevanti (19).

I transformer rappresentano un'architettura più recente, pensata anch'essa per i dati sequenziali ma con un funzionamento diverso dalle RNN. Il loro elemento caratteristico è il meccanismo di auto-attenzione (*self-attention*): nell'elaborare ciascun elemento di una sequenza, la rete "pesa" quanto ognuno degli altri elementi sia rilevante per interpretarlo, mettendo così in relazione diretta anche parti molto distanti tra loro. A differenza delle RNN, che procedono un elemento alla volta, i transformer considerano l'intera sequenza simultaneamente: questo li rende più veloci da addestrare e particolarmente efficaci nel cogliere relazioni a lunga distanza, come quelle presenti nelle sequenze di DNA, RNA e proteine, dove elementi lontani possono influenzarsi reciprocamente (19).

Una menzione particolare meritano i modelli generativi. Gli *autoencoder* variazionali (*variational autoencoder*, VAE) sono costituiti da un codificatore, che proietta i dati in uno spazio latente a dimensione ridotta, e da un decodificatore, che ricostruisce i dati a partire da tale rappresentazione; questa struttura consente di generare nuovi dati, ad esempio nuove strutture molecolari (19). Le reti generative avversarie (*generative adversarial network*, GAN) mettono invece in competizione due reti, un generatore e un discriminatore: il primo produce dati artificiali, il secondo cerca di distinguerli da quelli reali, e questo confronto spinge il generatore a produrre dati sempre più realistici (19).

Infine, le reti neurali a grafo (*graph neural network*, GNN, e la loro variante convoluzionale, GCN) elaborano dati strutturati come grafi, in cui le entità sono rappresentate come nodi e le loro relazioni come archi. Questa struttura è particolarmente adatta a rappresentare le piccole molecole, i cui atomi corrispondono ai nodi e i legami chimici agli archi, motivo per cui le GNN sono ampiamente impiegate nella previsione delle proprietà molecolari (18,19).

La scelta dell'architettura non è arbitraria, ma dipende dal tipo di dato biologico e dall'obiettivo: le GNN risultano più adatte alle piccole molecole, le CNN all'analisi di dati strutturati e di immagini, i transformer alle sequenze biologiche complesse. Questa corrispondenza non è casuale: ciascuna architettura è costruita per la struttura del dato che deve elaborare. Le GNN operano su grafi, e una piccola molecola è naturalmente rappresentabile come un grafo, con gli atomi come nodi e i legami come archi; le CNN sono progettate per dati organizzati a griglia, come i pixel di un'immagine; i transformer eccellono nel cogliere relazioni tra elementi distanti di una sequenza, come quelle delle lunghe catene biologiche. La scelta dell'architettura segue quindi la natura del dato e l'obiettivo dell'analisi (19). Questa preferenza trova riscontro anche nella letteratura: un'analisi degli studi pubblicati tra il 2014 e il 2024 mostra che, tra tutte le architetture, le reti convoluzionali, le reti a grafo e i transformer si sono affermate come le scelte più diffuse nella progettazione di farmaci mediante IA, con frequenze d'uso che variano a seconda che si tratti di piccole molecole, terapie a base di RNA o anticorpi (19) (**Figura 2**). A queste architetture si affianca spesso l'apprendimento per rinforzo (*reinforcement learning*), in cui il modello viene guidato verso un obiettivo attraverso un sistema di ricompense e penalizzazioni, impiegato in particolare per ottimizzare le molecole generate rispetto a proprietà desiderate (20).



**Figura 2.** Principali architetture di reti neurali impiegate nella progettazione di farmaci (A) e loro diffusione negli studi pubblicati tra il 2014 e il 2024 per piccole molecole, terapie a base di RNA e anticorpi (B) (19).

### 2.3.2 Applicazioni chimico-farmaceutiche

Le architetture appena descritte trovano numerose applicazioni nelle fasi iniziali della ricerca farmacologica. Un primo ambito è quello della rappresentazione delle molecole e della previsione delle loro proprietà. Una molecola può essere rappresentata in forme diverse: come stringa di caratteri (ad esempio nel formato SMILES), come insieme di sottostrutture (*fingerprint*), come grafo bidimensionale o come struttura tridimensionale (**Figura 3**). A ciascuna rappresentazione corrispondono architetture più adatte, come reti ricorrenti e transformer per le sequenze, reti convoluzionali per le immagini, reti a grafo per i grafi molecolari, e la combinazione di più rappresentazioni può migliorare la qualità delle previsioni (18). Questo perché ciascuna rappresentazione coglie aspetti diversi e complementari della stessa molecola: la sequenza di caratteri ne descrive la composizione, il grafo ne cattura la connettività tra atomi, la struttura tridimensionale ne riflette la forma nello spazio. Combinando più rappresentazioni, il modello dispone di un quadro più completo della molecola rispetto a quello offerto da una singola descrizione, e può formulare previsioni più accurate (18). Su queste basi si comprende l'evoluzione dei modelli QSAR introdotti in precedenza. I QSAR tradizionali si fondavano su descrittori molecolari scelti dall'esperto e su relazioni matematiche relativamente semplici tra questi e l'attività biologica. Il deep learning supera entrambi questi limiti: da un lato apprende automaticamente dai dati le caratteristiche rilevanti delle molecole, dall'altro è in grado di rappresentare relazioni molto più complesse tra struttura e proprietà. Ne derivano modelli capaci di prevedere con maggiore accuratezza proprietà come l'attività biologica o altre caratteristiche dei composti, a partire dalla sola struttura molecolare (18).

### Imatinib Mesylate (C<sub>30</sub>H<sub>35</sub>N<sub>7</sub>O<sub>4</sub>S)

#### (a) 1D Representation

SMILES:

CC1=C(C=C(C=C1)NC(=O)C2=CC=C(C=C2)CN3CCN(CC3)C)NC4=NC=CC(=N4)C5=CN=CC=C5.CS(=O)(=O)O

ECFP:

216545222: 2, 337875000: 1, 353395765: 2,  
368350036: 1, 847957139: 1, 847961216: 2,  
...  
3275683399: 1, 3918336191: 2, 3975275337: 1]

MACCS:

000000000.....10100110101111111  
1110

Mathematical Representation:

Differential geometry  
Algebraic graph  
Algebraic topology

#### (b) 2D Representation

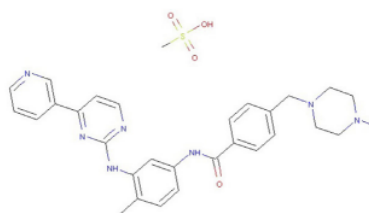
$$\begin{bmatrix} C & 0 & 1 & 0 & & 1 & 1 & 0 \\ C & 1 & 0 & 0 & \dots & 1 & 1 & 0 \\ C & 0 & 0 & 0 & & 0 & 0 & 0 \\ \vdots & \vdots & & \ddots & & \vdots & & \\ O & 1 & 1 & 0 & & 0 & 0 & 1 \\ O & 1 & 1 & 0 & \dots & 0 & 0 & 0 \\ O & 0 & 0 & 0 & & 1 & 0 & 0 \end{bmatrix}$$

Adjacent Matrix

$$\begin{bmatrix} C & 1 & 0 & \dots \\ C & 1 & 0 & \dots \\ C & 1 & 0 & \dots \\ \vdots & \vdots & \vdots & \\ O & 0 & 1 & \dots \\ O & 0 & 1 & \dots \\ O & 0 & 1 & \dots \end{bmatrix}$$

Feature Matrix

2D Molecular Graph



2D Molecular Image

#### (c) 3D Representation

$$\begin{bmatrix} C & 0 & 1 & 0 & & 1 & 1 & 0 \\ C & 1 & 0 & 0 & \dots & 1 & 1 & 0 \\ C & 0 & 0 & 0 & & 0 & 0 & 0 \\ \vdots & \vdots & & \ddots & & \vdots & & \\ O & 1 & 1 & 0 & & 0 & 0 & 1 \\ O & 1 & 1 & 0 & \dots & 0 & 0 & 0 \\ O & 0 & 0 & 0 & & 1 & 0 & 0 \end{bmatrix}$$

Adjacent Matrix

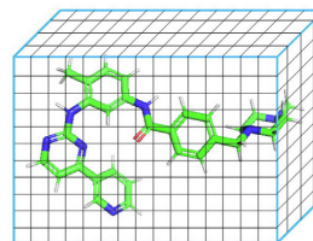
$$\begin{bmatrix} C & 6.16 & 0.32 & 0.78 \\ C & 7.85 & 0.67 & -0.90 \\ C & 7.12 & -0.56 & 1.58 \\ \vdots & \vdots & \vdots & \\ O & -0.83 & 0.25 & -1.04 \\ O & -6.61 & -0.26 & -0.56 \\ O & -1.83 & -0.23 & -0.11 \end{bmatrix}$$

3D Coordinates

$$\begin{bmatrix} C & 1 & 0 & \dots \\ C & 1 & 0 & \dots \\ C & 1 & 0 & \dots \\ \vdots & \vdots & \vdots & \\ O & 0 & 1 & \dots \\ O & 0 & 1 & \dots \\ O & 0 & 1 & \dots \end{bmatrix}$$

Feature Matrix

3D Molecular Graph



3D Molecular Grid

**Figure 3.** Diverse rappresentazioni di una stessa molecola (imatinib mesilato): stringa SMILES e fingerprint (1D), grafo e immagine molecolare (2D), grafo e griglia tridimensionali (3D) (18).

Un secondo ambito, reso possibile soprattutto dai modelli generativi, è la progettazione *de novo* di molecole, ovvero la generazione di nuove entità chimiche dotate di proprietà desiderate. A differenza dei metodi tradizionali, basati su regole esplicite definite dall'esperto, i modelli di deep learning apprendono direttamente dai dati le relazioni tra struttura e proprietà e sono in grado di esplorare in modo efficiente l'enorme spazio chimico delle molecole potenzialmente sintetizzabili (20). Per questo scopo vengono impiegate architetture come le reti ricorrenti, gli *autoencoder* variazionali e le reti generative avversarie, spesso combinate con l'apprendimento per rinforzo per orientare la generazione verso proprietà specifiche, quali l'affinità di legame, la selettività e un profilo favorevole di assorbimento, distribuzione, metabolismo, escrezione e tossicità (20).

Un terzo ambito riguarda l'analisi di immagini biologiche mediante reti convoluzionali. Tali reti sono state impiegate, ad esempio, per estrarre informazioni dalle immagini cellulari nell'ambito del cosiddetto profiling morfologico, una tecnica in cui le cellule vengono

trattate e colorate per evidenziarne diverse componenti, e le immagini così ottenute vengono analizzate per estrarne un gran numero di caratteristiche relative a forma, dimensione e organizzazione cellulare; l'insieme di queste caratteristiche costituisce una sorta di "impronta" morfologica che riflette lo stato della cellula e la sua risposta a un trattamento (21). Le stesse reti sono state inoltre impiegate per analizzare immagini istologiche tumorali: in studi sul carcinoma del colon-retto, modelli basati su reti convoluzionali si sono dimostrati in grado di prevedere l'esito clinico dei pazienti (22) e di identificare lo stato di specifiche vie molecolari e mutazioni direttamente da immagini istologiche di routine (23).

Queste applicazioni, pur promettenti, presentano limiti importanti. Le molecole generate dai modelli non sempre sono facilmente sintetizzabili e molti studi si fermano alla validazione computazionale (*in silico*), senza una verifica sperimentale (20). Inoltre, i modelli di deep learning si comportano spesso come una "scatola nera" (*black box*): pur fornendo previsioni accurate, non rendono trasparenti i passaggi e i criteri che li hanno condotti a una determinata conclusione, tanto che neppure chi li ha costruiti è sempre in grado di ricostruire il perché di una singola previsione. Questa opacità è particolarmente critica in ambito farmaceutico, dove la possibilità di comprendere e giustificare le decisioni, cioè l'interpretabilità del modello, è essenziale tanto per la fiducia degli esperti quanto per i requisiti regolatori (18).

## **2.4 Fondamenti farmacologici rilevanti per l'IA**

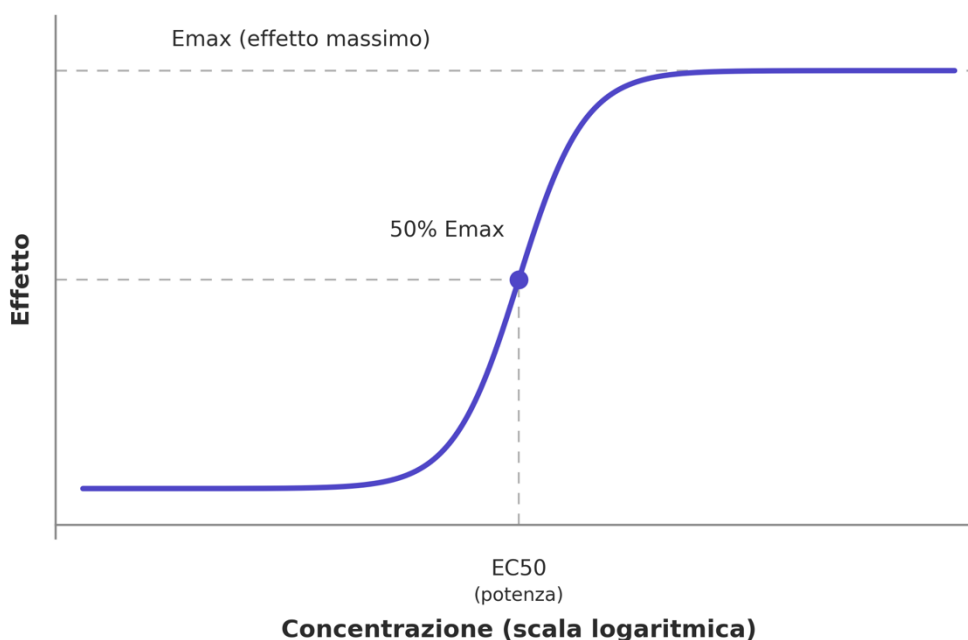
Per comprendere come l'IA si applichi alla farmacologia è utile richiamare alcuni concetti fondanti, che rappresentano le grandezze e le relazioni che i modelli computazionali cercano di apprendere e di prevedere. L'azione di un farmaco può essere descritta da due punti di vista complementari. La farmacodinamica studia ciò che il farmaco produce nell'organismo, mentre la farmacocinetica descrive ciò che l'organismo fa al farmaco. Distinguere chiaramente i due ambiti è essenziale, perché l'effetto terapeutico nasce proprio dal modo in cui essi si combinano.

Sul versante farmacodinamico, il fenomeno centrale è l'interazione farmaco-recettore, alla base dell'azione della maggior parte dei farmaci. Una molecola, detta ligando, si lega a una specifica macromolecola recettoriale, e tale legame può essere descritto in termini quantitativi mettendo in relazione la concentrazione del farmaco con la quota di recettori occupati (24). L'interazione è governata da due proprietà distinte: l'affinità, che esprime la tendenza del farmaco a legarsi al recettore, e l'efficacia, che indica la capacità del farmaco, una volta legato, di attivarlo e produrre una risposta. Da questa distinzione deriva la

classificazione dei farmaci in agonisti, che si legano al recettore e lo attivano, antagonisti, che si legano bloccandolo senza attivarlo, e agonisti parziali, che ne determinano un'attivazione soltanto parziale, con una risposta sub-massimale (24).

A queste due proprietà si collega la relazione struttura-attività (*Structure-Activity Relationship*, SAR), che descrive il legame tra la struttura chimica di una molecola e la sua attività biologica. Si tratta di un concetto centrale per l'applicazione dell'IA, perché modificazioni della struttura possono tradursi in variazioni dell'affinità e dell'efficacia: una stessa variazione può, ad esempio, aumentare l'affinità di una molecola per il recettore senza modificarne l'efficacia, o viceversa. Considerare un solo parametro fornirebbe quindi un quadro incompleto della relazione tra struttura e attività (24). È proprio su questo tipo di relazioni tra struttura e attività che si fondano i modelli predittivi come i QSAR, di cui si è discusso in precedenza (24).

L'entità della risposta, infine, può essere rappresentata dalla curva concentrazione-effetto, a sua volta descritta da due parametri quantitativi: l' $E_{max}$ , ossia l'effetto massimo che il farmaco è in grado di produrre, e l' $EC_{50}$ , la concentrazione che produce il 50% di tale effetto e che costituisce un indice della potenza del farmaco (25) (**Figura 4**).



**Figura 4.** Curva concentrazione-effetto. All'aumentare della concentrazione del farmaco (in scala logaritmica) l'effetto cresce fino a raggiungere un valore massimo ( $E_{max}$ ). La concentrazione che produce il 50% dell'effetto massimo ( $EC_{50}$ ) costituisce un indice della potenza del farmaco.

Questi parametri descrivono però soltanto la relazione tra concentrazione ed effetto. Per tradurla in un effetto reale nel paziente occorre conoscere anche quale concentrazione il farmaco raggiunge effettivamente nell'organismo nel corso del tempo, che è precisamente l'oggetto della farmacocinetica; è l'integrazione tra i due versanti a permettere di collegare l'esposizione al farmaco alla sua attività terapeutica (25).

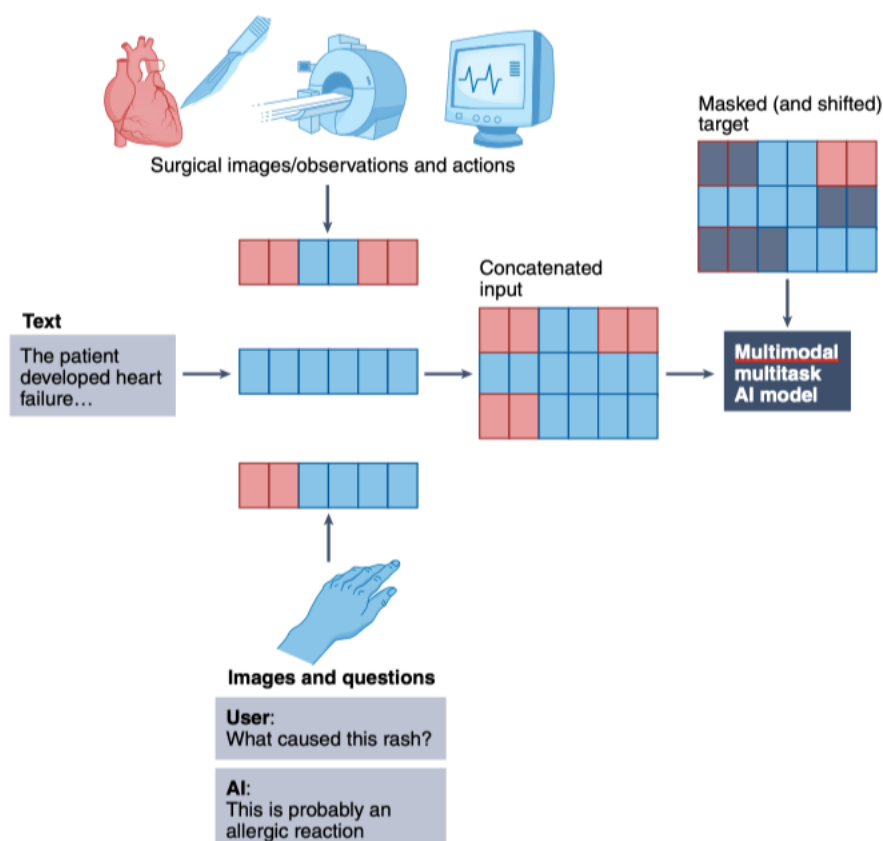
In questo passaggio assume particolare rilievo, soprattutto per i modelli predittivi, il principio della frazione libera. Solo la quota non legata alle proteine plasmatiche è in grado di interagire con il bersaglio e produrre una risposta. Ciò implica che, per stimare correttamente l'effetto di un farmaco, non basta considerarne la concentrazione totale, ma occorre tenere conto anche della frazione libera (25).

Di interesse crescente è la definizione di biomarcatore molecolare, definito come un parametro misurabile in modo oggettivo che funge da indicatore di un processo biologico normale o patologico, oppure della risposta a un trattamento (25). I biomarcatori costituiscono uno degli ambiti di maggiore interesse per l'IA: derivando da dati molecolari complessi e numerosi, si prestano all'analisi mediante algoritmi capaci di riconoscere schemi non immediatamente evidenti. Su questa base l'IA può contribuire a selezionare i pazienti più adatti a un determinato trattamento in funzione del loro profilo molecolare e a valutarne precocemente la risposta, ancor prima che l'effetto clinico sia pienamente manifesto (25).

## **2.5 *Big Data* e qualità dei dati**

I dati costituiscono il fondamento su cui si reggono tutti i modelli di IA: la quantità, la varietà e la qualità dei dati determinano in larga misura l'affidabilità dei modelli di IA che su di essi si fondano. Negli ultimi due decenni la ricerca biomedica ha generato volumi di dati senza precedenti, grazie ai progressi nelle tecnologie di sequenziamento, nella diagnostica per immagini, nei dispositivi indossabili e nella digitalizzazione delle cartelle cliniche (26).

Questi dati sono per loro natura eterogenei e multimodali: eterogenei perché provengono da fonti diverse e differiscono per formato, scala e grado di strutturazione; multimodali perché appartengono a modalità distinte, ovvero tipologie qualitativamente diverse di informazione, come una sequenza genetica, un'immagine radiologica o un valore di laboratorio, ciascuna delle quali descrive un aspetto differente dello stesso individuo. Comprendono, tra gli altri, i dati genomici e i dati cosiddetti *multi-omici* (trascrittoma, proteoma, metaboloma, microbioma), le immagini mediche e istologiche, i dati delle cartelle cliniche elettroniche e i parametri raccolti da sensori indossabili (26). L'integrazione di queste diverse fonti, cioè la cosiddetta *IA multimodale*, consente una caratterizzazione più completa dello stato di salute e di malattia di un individuo. Ciascuna tipologia di dato, presa singolarmente, ne illumina infatti solo un aspetto; la loro integrazione permette invece di metterli in relazione e di cogliere connessioni che nessuna fonte isolata rivelerebbe, restituendo un quadro più ricco e coerente, analogamente a quanto fa un clinico quando combina anamnesi, esami di laboratorio e diagnostica per immagini (26) (**Figura 5**).



**Figura 5.** Integrazione di dati multimodali in un modello di IA. Tipologie eterogenee di dato biomedico (genomici e multi-omici, immagini mediche e istologiche, cartelle cliniche elettroniche, sensori indossabili) confluiscono in un modello di IA multimodale, che le integra per ottenere una caratterizzazione più completa dello stato di salute e di malattia (26)

Nel contesto dello sviluppo del farmaco, l'integrazione dei dati multi-omici permette in particolare di identificare nuovi bersagli terapeutici e di stratificare i pazienti sulla base delle loro caratteristiche molecolari (27). I dati multi-omici descrivono i diversi livelli molecolari di funzionamento di un organismo: non solo il genoma, cioè l'insieme delle informazioni genetiche, ma anche i prodotti della loro espressione e attività, come l'insieme degli RNA trascritti (trascrittoma), delle proteine (proteoma) e dei metaboliti (metaboloma). La loro importanza per lo sviluppo del farmaco deriva dal fatto che operano al livello molecolare in cui i farmaci stessi agiscono e in cui le malattie hanno origine. Integrando questi livelli, è possibile individuare le molecole coinvolte in un processo patologico e quindi potenziali nuovi bersagli terapeutici, oltre a distinguere sottogruppi di pazienti che, pur condividendo la stessa diagnosi clinica, differiscono nel profilo molecolare e possono perciò rispondere in modo diverso a uno stesso trattamento: è ciò che si intende per stratificazione dei pazienti (27).

Trasformare questa enorme quantità di dati in conoscenza utile pone però sfide considerevoli, legate alla possibilità di accedervi, di renderli confrontabili tra loro e di garantirne la qualità. Dati raccolti da gruppi diversi, con strumenti e formati diversi, rischiano infatti di non poter essere messi insieme o riutilizzati. Per questo la loro gestione richiede regole comuni e strutture condivise. Tra le regole più diffuse vi sono i cosiddetti principi FAIR, secondo cui i dati dovrebbero essere reperibili, accessibili, interoperabili (cioè utilizzabili insieme ad altri dati, anche se prodotti altrove) e riutilizzabili.

Tra le strutture condivise figurano invece i consorzi collaborativi, che permettono a più centri di mettere in comune e curare grandi quantità di dati: in alcuni casi i dati non vengono nemmeno spostati dai centri che li custodiscono, ma resi analizzabili sul posto, una soluzione utile soprattutto quando si tratta di dati sanitari sensibili (27). A ciò si aggiungono i requisiti normativi in materia di protezione dei dati sanitari, primo fra tutti il Regolamento generale sulla protezione dei dati (GDPR) nell'Unione europea, che impone misure rigorose di tutela della riservatezza (26,27).

È tuttavia importante sottolineare che la disponibilità di grandi quantità di dati non è di per sé sufficiente a garantire il successo dei modelli di IA nel settore farmaceutico. A differenza di ambiti come il riconoscimento di immagini o del linguaggio, in cui l'IA ha ottenuto risultati notevoli, i dati chimici e biologici presentano difficoltà intrinseche su più piani (28).

In primo luogo, i dati che realmente rilevanti, ossia quelli che indicano se una molecola sia efficace e sicura nell'uomo (gli endpoint clinici, come la guarigione, la riduzione dei sintomi o la sopravvivenza), sono alquanto scarsi se confrontati con il numero estremamente elevato di molecole teoricamente possibili. Lo spazio chimico, ossia l'insieme di tutte le molecole con caratteristiche da farmaco potenzialmente sintetizzabili, è stimato intorno a  $10^{60}$ : di questa quantità enorme, solo una porzione minima è stata effettivamente studiata nell'uomo. Il modello dispone quindi di un numero molto elevato di molecole possibili ma di un numero assai ridotto di dati clinici certi su cui apprendere (28). In secondo luogo, non è sempre chiaro quale sia la modalità più adeguata per descrivere una molecola in forma computazionale. Una stessa molecola, infatti, possiede molte caratteristiche diverse, e caratteristiche diverse rilevano per effetti diversi: la parte della struttura che determina, ad esempio, il legame con un bersaglio non è necessariamente la stessa che ne influenza la solubilità o la tossicità. Scegliere quali aspetti della molecola rappresentare, e in quale modo, non è quindi un'operazione immediata (28). In terzo luogo, e soprattutto, è difficile assegnare ai dati biologici “etichette” affidabili. Etichettare un dato significa associargli l'informazione che il modello deve imparare a prevedere, ad esempio se una molecola sia tossica o efficace. Occorre tuttavia considerare che la tossicità o l'attività non sono proprietà fisse e univoche di una molecola, valide in ogni circostanza: dipendono dal contesto, ossia dalla dose somministrata, dal sistema biologico in cui la si misura, dalle condizioni dell'esperimento e dalle caratteristiche del singolo paziente. La stessa molecola può quindi risultare efficace o tossica in una situazione e non in un'altra. Di conseguenza, l'etichetta che forniamo al modello reca con sé questa incertezza, e un modello che apprende da etichette ambigue produrrà a sua volta previsioni poco affidabili (28).

Un'ulteriore difficoltà deriva dal ricorso frequente a misure indirette, dette proxy. Un proxy è una grandezza facile da misurare che viene usata al posto di un'altra, più difficile da osservare, nella speranza che la prima rifletta la seconda. In questo ambito, parametri come l'affinità di una molecola per il proprio bersaglio o la sua attività in un saggio di laboratorio (in vitro) vengono impiegati come approssimazioni di ciò che conta davvero, cioè l'efficacia e la sicurezza nell'organismo (in vivo) (28). Il limite è che questa corrispondenza non è garantita. Un saggio in vitro isola infatti un singolo aspetto, mentre nell'organismo intervengono numerosi fattori che in vitro non sono presenti, come l'assorbimento, la distribuzione nei tessuti, il metabolismo e l'eliminazione del farmaco. Per questo una molecola promettente in laboratorio può rivelarsi inefficace o tossica una volta

somministrata, e la capacità dei dati in vitro di predire gli esiti clinici reali resta spesso limitata (28). Ne consegue che, in molti casi, il limite principale non sta nell'algoritmo impiegato, ma nella natura stessa dei dati disponibili. Per un uso davvero efficace dell'IA nel settore farmaceutico, prima ancora di disporre di algoritmi più sofisticati, occorre comprendere meglio i sistemi biologici e generare dati pertinenti, ben rappresentati e correttamente etichettati: dati cioè che misurino davvero ciò che interessa prevedere (pertinenti), descritti in una forma adeguata al problema (ben rappresentati) e accompagnati da etichette affidabili (correttamente etichettati) (28).

### 3. L'INTELLIGENZA ARTIFICIALE NEL PROCESSO DI RICERCA E SVILUPPO DI NUOVI FARMACI

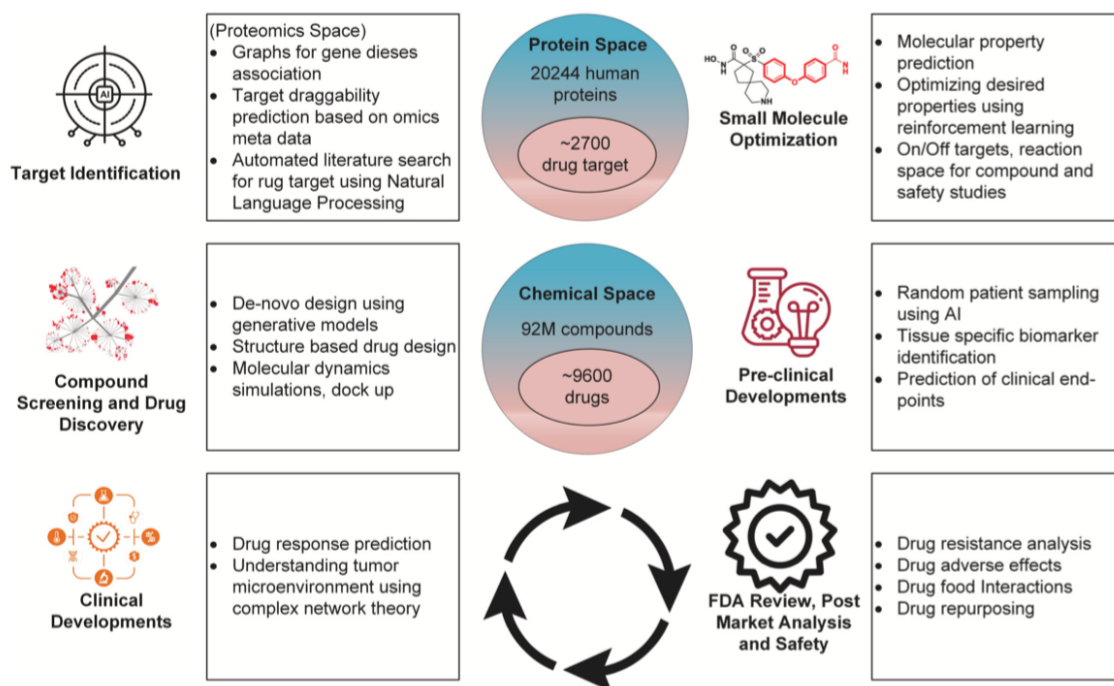
L'integrazione dell'IA nella ricerca e sviluppo di nuovi farmaci (*drug discovery*) ha determinato un profondo cambiamento nel modo in cui la ricerca farmacologica affronta le fasi precoci dello sviluppo dei farmaci, con applicazioni che spaziano dallo sviluppo di piccole molecole fino alla progettazione di anticorpi monoclonali e farmaci a base di RNA (29).

Il processo di sviluppo di un nuovo farmaco prevede tipicamente diverse fasi sequenziali: identificazione del target e dei biomarcatori, scoperta e ottimizzazione del *lead compound*, sviluppo clinico e sorveglianza post-commercializzazione. Da ciò si può evincere come lo sviluppo di un farmaco sia un processo lungo e costoso che richiede una quantità di risorse significativa, in particolare nella fase di ricerca clinica, che da sola rappresenta circa il 60% della spesa complessiva.

In questo contesto, l'IA potrebbe intervenire offrendo strumenti concreti per ridurre i costi nelle fasi precoci di ricerca e sviluppo, con stime che indicano un potenziale di risparmio compreso tra 60 e 110 miliardi di dollari all'anno per l'intero settore farmaceutico.

L'impiego di metodologie basate sull'IA per l'ottimizzazione della selezione dei pazienti, nella progettazione degli studi clinici e nell'analisi dei dati clinici consente inoltre di ottenere sperimentazioni più efficienti e successivamente di accelerare i tempi per l'approvazione da parte delle agenzie regolatorie, mantenendo al contempo elevati standard di sicurezza e qualità (29).

I recenti progressi nelle tecniche sperimentali hanno arricchito in modo significativo i database biomedici, rendendo disponibili dati di espressione genica, dati genomici e proteomici, nonché immagini di cellule e tessuti (29). Questa disponibilità di dati multi-omici costituisce il presupposto fondamentale per l'applicazione dell'IA su larga scala nella ricerca farmacologica (30). Infatti, diversi metodi basati sull'IA permettono di: analizzare grandi volumi di dati, incluse ad esempio immagini cellulari per valutare le alterazioni morfologiche e funzionali indotte dalla somministrazione dei farmaci; l'analisi di sequenze e strutture proteiche per l'identificazione di nuovi target molecolari ; creazione e ottimizzazione di strutture molecolari; infine, integrano e correlano informazioni eterogenee relative a target, molecole e patologie (ad esempio utilizzando *Knowledge Graph*), rappresentando uno strumento sempre più centrale nell'analisi dei sistemi biologici complessi (29). **(Figura 6)**



**Figura 6.** Applicazioni dei principali metodi di IA nelle diverse fasi del processo di *drug discovery*, dalla identificazione del target alla progettazione molecolare e al *virtual screening*. (31)

Un esempio paradigmatico di integrazione dell'IA lungo l'intera pipeline dello sviluppo farmacologico è rappresentato dalla piattaforma *Pharma.AI* sviluppata dall'azienda Insilico Medicine, che combina tre moduli distinti e complementari: *PandaOmics*, dedicato all'identificazione di target terapeutici e biomarcatori attraverso l'analisi integrata di dati multi-omici e knowledge graph; *Chemistry42*, orientato alla progettazione de novo e all'ottimizzazione molecolare mediante modelli generativi; e *InClinico*, finalizzato alla predizione dell'esito degli studi clinici tramite l'analisi di dati real-world (29). Questa piattaforma ha già portato ad un primo candidato, scoperto e progettato tramite IA (32), fino alla sperimentazione clinica, un caso che verrà approfondito nel capitolo 6.

Questo approccio illustra concretamente come l'IA possa operare lungo tutte le fasi del processo farmaceutico, migliorando l'efficienza complessiva e aumentando la probabilità di successo clinico

### 3.1 Identificazione e validazione dei target farmacologici

L'identificazione e la validazione del target farmacologico rappresenta il punto di partenza di qualsiasi programma di sviluppo farmacologico. In questa fase si individuano le proteine o più in generale le biomolecole, tra cui enzimi, recettori o RNA, potenzialmente coinvolte

nei meccanismi patogenetici della malattia, con l'obiettivo di selezionare bersagli terapeutici la cui modulazione possa tradursi in un effetto clinico rilevante (33).

I progressi nelle tecniche sperimentali ad alto throughput (ossia un insieme di metodi sperimentali automatizzati o semi-automatizzati che permettono di analizzare rapidamente un numero molto elevato di target in parallelo), uniti alla disponibilità crescente di dati multi-omici provenienti da ampie coorti di pazienti, hanno reso possibile la costruzione di modelli pato-fisiologici sempre più dettagliati (29,33). Il confronto tra profili genomici, trascrittomici e proteomici di soggetti malati e sani consente di descrivere le patologie non solo in termini clinici, ma sulla base dei loro meccanismi molecolari sottostanti, identificando specifiche alterazioni biologiche che ne definiscono gli endotipi (33). L'analisi e l'integrazione di questi dati spesso eterogenei, difficilmente gestibili con strumenti bioinformatici tradizionali, vengono oggi condotte attraverso algoritmi di *machine learning* sia supervisionati che non supervisionati, con l'obiettivo di stratificare i pazienti in gruppi omogenei e di associare specifici profili molecolari a risposte terapeutiche differenti (33).

Un aspetto centrale di questi approcci riguarda la possibilità di identificare possibili bersagli terapeutici attraverso l'analisi delle reti biologiche, complessi sistemi di interazioni molecolari in cui geni, proteine e altre biomolecole sono collegati tra loro da relazioni funzionali formando una rete che riflette il funzionamento cellulare nel suo insieme. Le informazioni derivanti dalle reti di interazione proteina-proteina possono essere rappresentate come network biologici complessi, nei quali i nodi corrispondono a geni o proteine e le connessioni alle loro interazioni funzionali: una rappresentazione spesso indicata come i sopramenzionati *Knowledge Graph* (33). Attraverso l'analisi della struttura di queste reti è possibile individuare le proteine o i geni che occupano una posizione centrale, i cosiddetti *hub*, i quali, proprio per la loro rilevanza nell'organizzazione della rete, rappresentano candidati particolarmente promettenti come bersagli terapeutici, in quanto una proteina altamente connessa all'interno della rete biologica di una malattia è probabile che svolga un ruolo importante in quel processo patologico (33).

Un esempio concreto è rappresentato da *GuiltyTargets*, uno strumento che analizza le reti di interazione tra proteine a livello genomico, integrando queste informazioni con i dati di espressione genica, al fine di identificare e classificare i target terapeutici più rilevanti per una determinata malattia (29). Applicato alla malattia di Alzheimer, *GuiltyTargets* ha permesso di identificare tre target candidati, la cui rilevanza nella patologia era già supportata dalla letteratura scientifica, dimostrando al contempo la sua applicabilità a diverse classi di patologie, tra cui tumori, patologie metaboliche e malattie neurodegenerative (29).

La prioritizzazione dei target così identificati richiede tuttavia una valutazione che vada oltre la sola rilevanza biologica: è necessario considerare la modulabilità farmacologica del target, il rischio di effetti off-target e la fattibilità dello sviluppo di un farmaco (*druggability*) (33). A questo proposito, la disponibilità di informazioni strutturali sulle proteine e i progressi nella predizione della loro conformazione tridimensionale ha migliorato significativamente la valutazione della *druggability* dei target candidati (30,33). Modelli predittivi come *AlphaFold* e *RoseTTAFold* consentono di ottenere la struttura tridimensionale delle proteine direttamente informatizzate, senza dover ricorrere a esperimenti di laboratorio complessi e costosi, rendendo possibile lo studio e la caratterizzazione farmacologica anche di proteine difficili da analizzare con metodiche tradizionali (30). Tuttavia, l'utilità di queste strutture predette per la progettazione razionale dei farmaci richiede una validazione rigorosa, trattata nel paragrafo 3.2.1.

### **3.1.1 Target validation biologica e approcci computazionali**

Una volta identificato un potenziale bersaglio terapeutico, la fase successiva consiste nella sua validazione biologica, ovvero nella verifica che la modulazione di quel target produca effettivamente l'effetto terapeutico desiderato nel contesto fisiopatologico rilevante. Questa fase è critica: un target identificato sulla base di dati computazionali deve essere confermato attraverso evidenze biologiche sperimentali prima che su di esso si possa avviare un programma di sviluppo farmacologico (33).

La validazione biologica del target si avvale di un insieme di approcci complementari, che includono tecniche di silenziamento genico come CRISPR-Cas9 e siRNA, modelli cellulari e animali, e analisi funzionali in vitro (33). Parallelamente, l'IA contribuisce a questo processo attraverso l'analisi integrata di dati multi-omici, che permette di correlare specifiche alterazioni molecolari con fenotipi di malattia ben definiti e di costruire modelli predittivi della risposta biologica alla modulazione del target (33). Un ulteriore aspetto riguarda lo studio dei pathway cellulari: comprendere in quale rete di segnalazione è coinvolto un potenziale target permette di valutare se la sua modulazione produrrà l'effetto terapeutico desiderato e, al tempo stesso, di anticipare possibili effetti indesiderati legati all'interferenza con altri elementi del sistema biologico (33).

Uno strumento di particolare rilievo nella validazione computazionale del target è la piattaforma *PandaOmics* dell'azienda Insilico Medicine, che integra dati di trascrittomica e proteomica con una vasta raccolta di letteratura biomedica e brevetti scientifici. Ad esempio,

applicata all'identificazione di target nella sclerosi laterale amiotrofica (SLA), *PandaOmics* ha consentito di analizzare profili di espressione genica derivati da campioni del sistema nervoso centrale (SNC) e da neuroni motori differenziati da cellule staminali indotte, fornendo candidati target con un profilo di rilevanza biologica e terapeutica misurabile (29). È importante sottolineare che gli approcci computazionali di *target validation*, per quanto potenti, non sostituiscono la sperimentazione biologica ma la orientano e la ottimizzano, riducendo il numero di esperimenti necessari e concentrando le risorse sui candidati con maggiore probabilità di successo (29,33).

### **3.1.2 Reti di pathway cellulari e reti di interazione proteica**

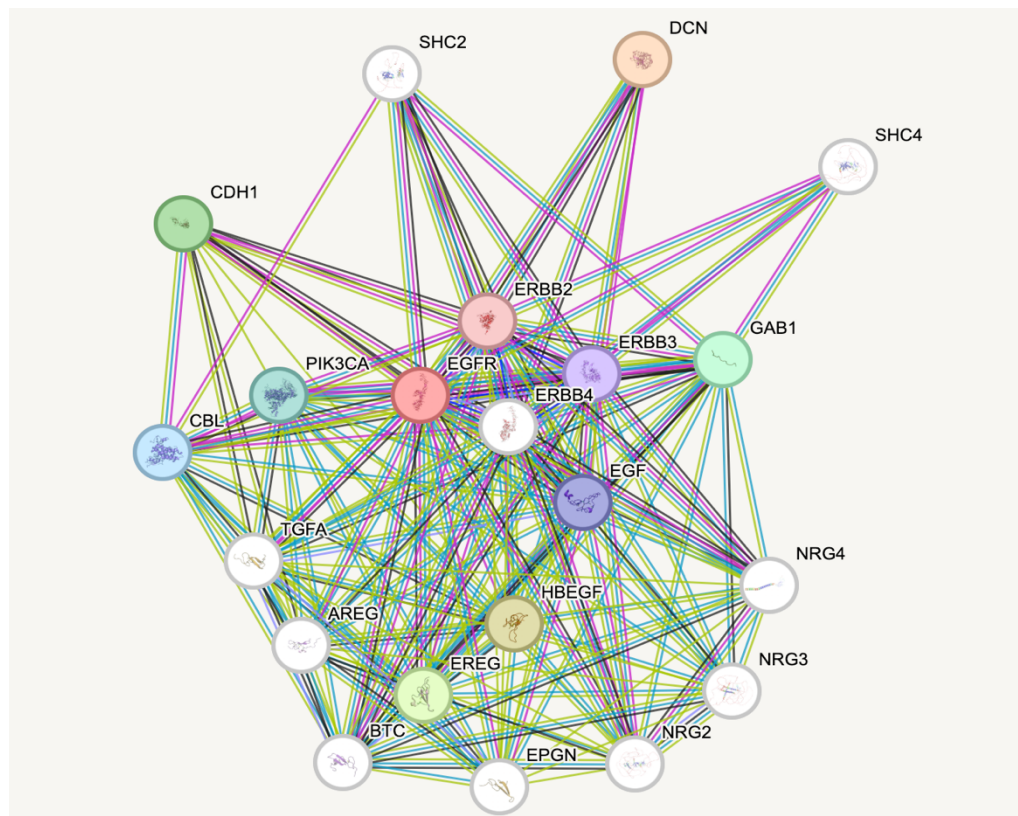
La comprensione dei meccanismi patologici alla base di una malattia non si esaurisce nell'identificazione di un singolo target. Le cellule funzionano attraverso reti complesse di segnalazione intracellulare, i cosiddetti *pathway cellulari*, in cui proteine, enzimi e recettori interagiscono tra loro in sequenze ordinate per trasmettere segnali e regolare processi fondamentali come la proliferazione, la sopravvivenza e la morte cellulare. Un'alterazione in uno qualsiasi dei nodi di questi pathway può propagarsi lungo tutta la rete, producendo effetti biologici difficilmente prevedibili analizzando le singole componenti in modo isolato (30).

In questo contesto, l'IA offre strumenti in grado di analizzare queste reti nella loro interezza, integrando informazioni provenienti da diverse fonti, come dati genomici, proteomici e trascrittomici, per costruire modelli computazionali della malattia che riflettano la reale complessità biologica del sistema. I pathway cellulari disfunzionali vengono rappresentati sotto forma di reti di interazione proteina-proteina, nelle quali i nodi corrispondono a geni o proteine e i collegamenti rappresentano le relazioni funzionali tra di essi. Attraverso l'analisi di questa struttura reticolare è possibile identificare i punti critici del sistema, ossia le proteine che occupano una posizione centrale nella rete e che, se alterate, influenzano il funzionamento di molte altre componenti (30).

Dal punto di vista farmacologico, questo approccio ha un'implicazione diretta: intervenire su una proteina centrale di un pathway disfunzionale può produrre un effetto terapeutico più ampio e duraturo rispetto al blocco di un singolo bersaglio “periferico”. Allo stesso tempo, conoscere la struttura della rete biologica consente di anticipare un problema frequente nella terapia farmacologica ossia la resistenza al farmaco durante il trattamento. Questo fenomeno si verifica quando le cellule, anche dopo che un farmaco ha bloccato una proteina bersaglio,

riescono a sopravvivere attivando vie di segnalazione alternative, aggirando il blocco farmacologico. Analizzando la rete nel suo insieme, l'IA può potenzialmente identificare questi pathway alternativi prima che la resistenza si manifesti, suggerendo combinazioni di target su cui agire simultaneamente (15,30).

Un esempio rilevante in questo senso è rappresentato dall'utilizzo di database specializzati come *Ingenuity Pathway Analysis* e *STRING*, che raccolgono informazioni sulle interazioni funzionali tra geni e proteine e consentono di ricostruire i pathway alterati in una determinata patologia. L'integrazione di questi database con algoritmi di *machine learning* permette di stratificare i pathway in base al loro grado di alterazione nei campioni biologici di pazienti rispetto ai tessuti sani, fornendo una mappa molecolare della malattia su cui basare le decisioni di *drug design* (15) (**Figura 7**).



**Figura 7.** Rete di interazione proteina-proteina centrata su EGFR (STRING database). I nodi rappresentano proteine, le connessioni le loro interazioni funzionali. I colori delle connessioni indicano il tipo di evidenza: giallo per co-espressione genica, azzurro per interazioni sperimentalmente validate, verde per associazioni da database biologici. Fonte: Szklarczyk et al., Nucleic Acids Research, 2023, string-db.org.

L'analisi delle reti di interazione proteica ha inoltre permesso di osservare che i farmaci approvati tendono a legarsi preferenzialmente a proteine altamente connesse nella rete

biologica, confermando che la centralità nella rete è un indicatore della rilevanza farmacologica di un target. Questo approccio, noto come *network medicine*, si basa sull'idea che le malattie non siano alterazioni di singole molecole isolate, ma disfunzioni di intere reti biologiche, e rappresenta oggi uno dei fondamenti teorici della scoperta computazionale di nuovi farmaci. L'IA ne ha quindi significativamente ampliato le possibilità applicative, consentendo di analizzare reti di dimensioni e complessità che sarebbero inaccessibili con i metodi tradizionali (15,30).

### **3.2 Progettazione razionale di farmaci assistita da IA (*AI-assisted Drug Design*)**

Una volta identificato e validato il target terapeutico, la fase successiva del processo consiste nella progettazione di molecole in grado di interagire con esso in modo selettivo ed efficace. Questo processo, noto come *progettazione razionale del farmaco*, si basa sulla conoscenza delle caratteristiche strutturali e funzionali del target per guidare la sintesi di nuovi composti candidati, riducendo il ricorso alla sperimentazione casuale (29,30).

In questo ambito si distinguono tradizionalmente due approcci principali: il primo, noto come *Structure-Based Drug Design* (SBDD), parte dalla struttura tridimensionale della proteina bersaglio per progettare molecole che si adattino alla sua conformazione e interagiscano con il suo sito di legame; il secondo, noto come *Ligand-Based Drug Design* (LBDD), viene utilizzato quando la struttura del target non è disponibile. In questo caso si parte dall'analisi di molecole già note per la loro efficacia biologica, con l'obiettivo di individuare quali caratteristiche chimiche le rendono attive e di utilizzare questa informazione per progettare nuovi composti con attività simile o superiore (15,30).

L'IA ha profondamente trasformato entrambi questi approcci, introducendo metodi in grado di accelerare la fase di progettazione, migliorare la precisione delle previsioni e gestire la complessità dei dati coinvolti (29,30).

Nei paragrafi seguenti vengono descritte le principali applicazioni dell'IA nell'ambito dello SBDD e del LBDD, con particolare attenzione agli strumenti che hanno avuto il maggiore impatto sul processo di *drug discovery*.

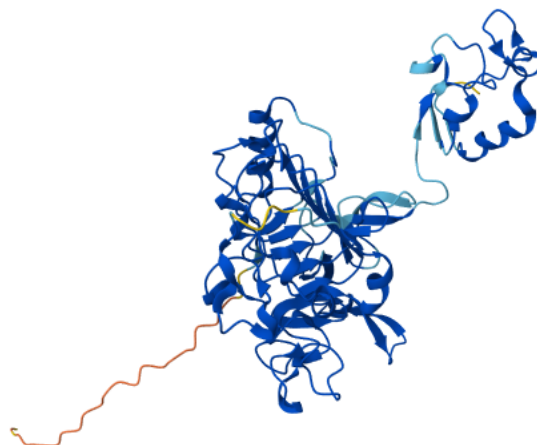
Accanto alla progettazione di piccole molecole, l'IA trova applicazione anche nei farmaci biologici, come anticorpi e vaccini, dove un passaggio cruciale è l'identificazione degli epitopi, ovvero le regioni di un antigene riconosciute dal sistema immunitario. Questi sono difficili da prevedere perché dipendono dalla struttura tridimensionale dell'antigene e non

dalla sola sequenza amminoacidica. Per affrontare questo problema, strumenti come *GraphBepi* integrano la struttura dell'antigene predetta da AlphaFold2 con la sua sequenza, analizzandole tramite reti neurali a grafo, così da prevedere con maggiore accuratezza quali regioni costituiscono l'epitopo; altri modelli, come DeepAIR, combinano analogamente informazioni di sequenza e di struttura per prevedere il legame tra i recettori immunitari e l'antigene, con applicazioni nella progettazione di anticorpi terapeutici (29).

### **3.2.1 Structure-Based Drug Design (SBDD): AlphaFold, Molecular Docking, Free Energy Perturbation**

La progettazione di farmaci basata sulla struttura del target, nota come *Structure-Based Drug Design* (SBDD), rappresenta uno degli approcci più consolidati nella farmacologia computazionale. Il principio fondamentale di questo metodo consiste nell'utilizzare la struttura tridimensionale della proteina bersaglio per progettare molecole in grado di legarsi al suo sito attivo in modo preciso e selettivo. La disponibilità di strutture proteiche ad alta risoluzione, ottenute sperimentalmente tramite cristallografia a raggi X o crio-elettromicroscopia, ha rappresentato per decenni il principale requisito per applicare questo approccio (29,33).

Un cambiamento radicale in questo scenario è stato introdotto da *AlphaFold*, un sistema di IA sviluppato da *DeepMind* (Google) che ha rivoluzionato la biologia strutturale rendendo possibile la predizione della struttura tridimensionale delle proteine direttamente dalla loro sequenza amminoacidica. Prima di *AlphaFold*, determinare la struttura di una proteina richiedeva anni di lavoro sperimentale e non era sempre possibile, per fattibilità, tempistiche e costi. *AlphaFold* ha reso disponibili predizioni di strutture per la quasi totalità delle proteine del proteoma umano, aprendo la strada all'applicazione dello SBDD su un numero di target prima inaccessibili (33). **(Figura 8)**



**Figura 8.** Struttura tridimensionale della proteina EGFR (recettore del fattore di crescita epidermico) predetta da AlphaFold. La colorazione riflette il livello di confidenza della predizione: blu scuro indica alta confidenza, celeste confidenza buona, giallo confidenza bassa e arancione/rosso confidenza molto bassa, tipicamente nelle regioni flessibili o disordinate. Fonte: *AlphaFold Protein Structure Database, Varadi et al., Nucleic Acids Research, 2022, alphafold.ebi.ac.uk.*

Tuttavia, l'utilizzo di strutture proteiche ottenute tramite predizione computazionale anziché determinate sperimentalmente introduce alcune limitazioni. Le strutture ottenute tramite predizione computazionale rappresentano la conformazione più probabile della proteina in condizioni statiche, ma non tengono conto del fatto che le proteine non sono strutture rigide, ma cambiano forma continuamente in risposta all'ambiente circostante ed ai loro ligandi. Questo è importante perché un farmaco non interagisce con una proteina statica, ma con una proteina dinamica, e la sua efficacia dipende anche da come i due si adattano reciprocamente durante il legame (33).

Una volta disponibile la struttura del target, il passo successivo nello SBDD è il *docking molecolare*, una tecnica computazionale che simula l'interazione fisica tra una molecola candidata e il sito di legame della proteina, calcolando la geometria e l'energia di tale interazione. Il docking permette di selezionare rapidamente, tra migliaia o milioni di molecole candidate, quelle con la maggiore probabilità di legarsi al target in modo efficace, riducendo il numero di esperimenti da condurre in laboratorio. L'IA ha migliorato significativamente la velocità e la precisione del docking, introducendo modelli in grado di gestire la flessibilità della proteina e di valutare con maggiore accuratezza l'affinità di legame (29,33).

Per una valutazione ancora più precisa dell'affinità di legame tra una molecola e il suo target, viene utilizzata una tecnica computazionale più avanzata nota come *Free Energy Perturbation* (FEP). Questa tecnica permette di confrontare due molecole molto simili tra

loro, per esempio due composti che differiscono solo per la presenza o assenza di un gruppo chimico, e di calcolare quale delle due si lega più fortemente al target. In pratica, simula cosa succederebbe dal punto di vista energetico se si apportasse una piccola modifica chimica a una molecola, consentendo di prevedere se quella modifica migliora o peggiora l'affinità di legame prima ancora di sintetizzare il composto in laboratorio. Questo guida le scelte di ottimizzazione chimica durante il processo di sviluppo del candidato farmacologico. Con l'IA, i calcoli FEP, tradizionalmente molto costosi in termini di tempo computazionale, possono oggi essere eseguiti in tempi significativamente ridotti, rendendoli applicabili su scala più ampia nel contesto industriale (33).

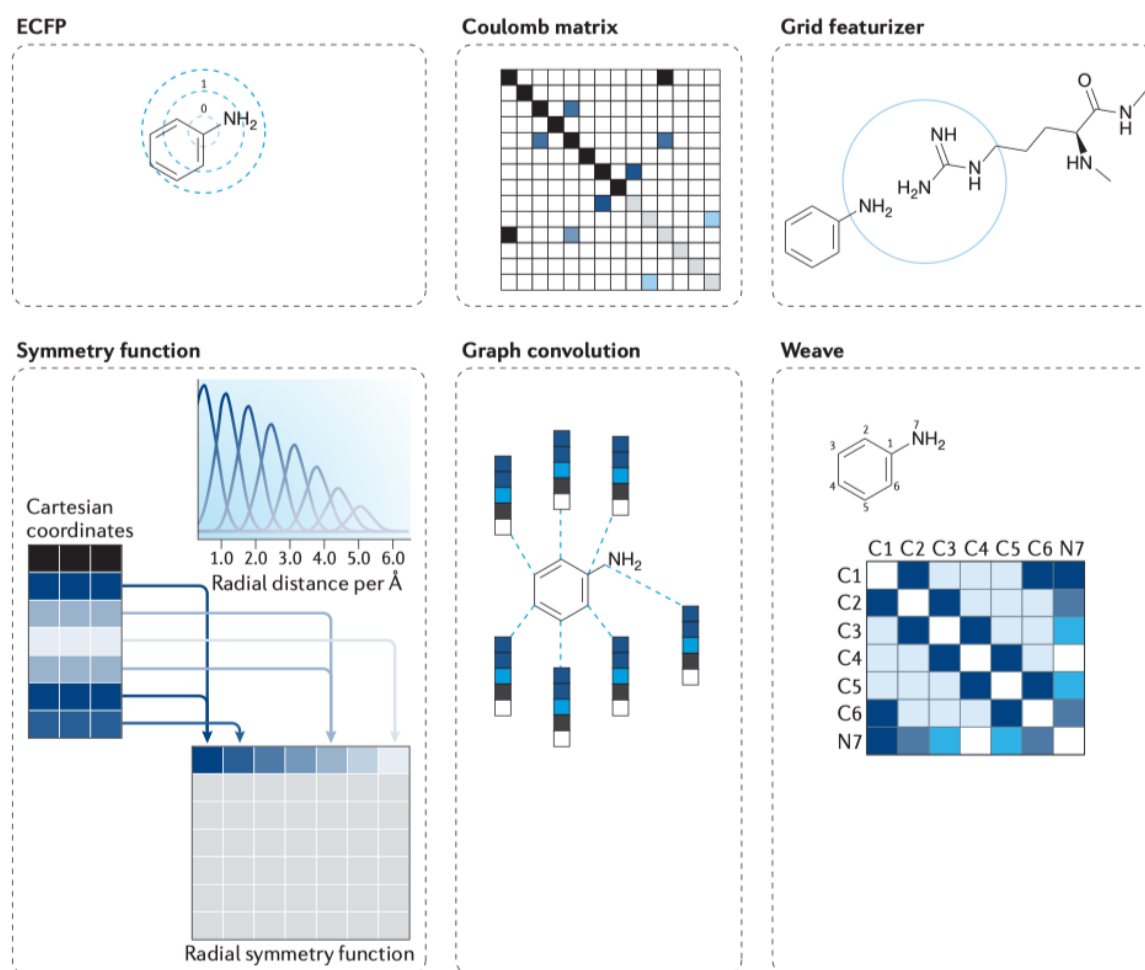
Un esempio concreto dell'impatto di questi strumenti è rappresentato dal loro utilizzo nella valutazione della selettività cardiaca dei farmaci. Il canale ionico hERG, una proteina coinvolta nella regolazione del ritmo cardiaco, è noto per essere un target aspecifico di numerosi farmaci, causando effetti collaterali cardiologici. L'utilizzo combinato di strutture proteiche predette da *AlphaFold* e calcoli FEP ha permesso di valutare il rischio di legame indesiderato con hERG già nelle fasi iniziali del *drug design*, contribuendo a ridurre il rischio di tossicità cardiaca (33).

### **3.2.2 Ligand-Based Drug Design (LBDD) e modelli QSAR (*Quantitative Structure-Activity Relationship*) potenziati dall'IA**

Quando la struttura tridimensionale del target non è disponibile o non è sufficientemente definita per applicare un approccio SBDD, la progettazione del farmaco può basarsi sulle caratteristiche chimiche delle molecole già note per la loro attività biologica. Questo approccio, noto come *Ligand-Based Drug Design* (LBDD), si fonda sull'idea che molecole con struttura chimica simile tendano ad avere effetti biologici simili, e utilizza questa relazione per identificare o progettare nuovi composti con proprietà farmacologiche desiderate (15,30).

Uno degli strumenti principali nell'ambito dell'LBDD è rappresentato dai modelli QSAR, acronimo di *Quantitative Structure-Activity Relationship*, ovvero relazione quantitativa tra struttura e attività. Questi modelli si basano sul concetto che l'attività biologica di una molecola, come ad esempio la sua capacità di inibire una proteina target, dipenda dalle sue caratteristiche chimiche e fisiche, e che questa relazione possa essere descritta matematicamente e utilizzata per fare previsioni su molecole nuove (15).

Per applicare questi modelli, la struttura chimica di ciascuna molecola viene tradotta in un insieme di valori numerici chiamati descrittori molecolari, cioè misure quantitative di proprietà chimiche e fisiche della molecola, come il peso molecolare, la solubilità, la distribuzione delle cariche elettriche o la forma tridimensionale. Questi descrittori vengono utilizzati come input per algoritmi di *machine learning* (**Figura 9**), che imparano a riconoscere quali combinazioni di caratteristiche chimiche sono associate a una maggiore attività biologica (15).



**Figura 9.** Diverse rappresentazioni molecolari utilizzate nei modelli di machine learning per il *drug discovery*. Ciascun metodo traduce la struttura chimica di una molecola in un formato matematico che il modello può elaborare, catturando caratteristiche diverse della molecola stessa (15).

L'introduzione del deep learning ha significativamente migliorato le prestazioni di questi modelli rispetto ai metodi statistici tradizionali. In particolare, le reti neurali multi-task, architetture in grado di apprendere simultaneamente più proprietà farmacologiche di una molecola, hanno dimostrato una capacità predittiva superiore rispetto ai modelli che valutano una sola proprietà alla volta, sfruttando le relazioni tra diversi endpoint biologici per migliorare la qualità complessiva delle previsioni (15).

Un ulteriore avanzamento in questo campo è rappresentato dall'introduzione di modelli basati su reti neurali a grafo, che rappresentano la molecola non come una lista di descrittori numerici ma come una struttura reticolare in cui gli atomi sono i nodi e i legami chimici sono le connessioni. Questo tipo di rappresentazione è più fedele alla natura tridimensionale delle molecole e consente di catturare informazioni strutturali che i descrittori tradizionali non riescono a rappresentare. In particolare, il modello GEM (*Geometry-Enhanced Molecular representation*) sviluppato da Fang et al. 2022 (34) ha introdotto la geometria tridimensionale della molecola come informazione aggiuntiva nel processo di apprendimento, migliorando significativamente la capacità di predire proprietà molecolari rilevanti per il *drug discovery*, come l'affinità di legame e la tossicità (34).

Nel complesso, l'applicazione del machine learning e del deep learning ai modelli LBDD e QSAR ha reso possibile una valutazione più rapida e accurata delle proprietà farmacologiche dei composti candidati, riducendo il numero di molecole da sintetizzare e testare sperimentalmente e accelerando il processo di identificazione dei lead compound (15,34).

### 3.2.3 Ottimizzazione del lead compound assistita da IA

Una volta identificata una molecola con attività biologica promettente, il cosiddetto *lead compound*, il processo entra nella fase di ottimizzazione, il cui obiettivo è migliorare progressivamente le caratteristiche farmacologiche del composto fino a renderlo un candidato farmacologico adatto alla sperimentazione clinica. Questa fase richiede di agire su più parametri contemporaneamente: aumentare l'affinità di legame con il target, migliorare la selettività per evitare effetti indesiderati e ottimizzare la potenza biologica del composto. L'IA ha reso questo processo più efficiente, riducendo il numero di cicli di sintesi e test necessari per raggiungere il profilo farmacologico desiderato (15,33).

L'affinità di legame misura quanto fortemente una molecola si lega al suo target e più è alta, minore è la quantità di farmaco necessaria per produrre un effetto terapeutico. Dal punto di vista energetico, questa affinità viene descritta attraverso l'energia libera di legame, un valore che quantifica la stabilità del complesso formato tra la molecola e la proteina bersaglio. Modelli computazionali basati su IA, come quelli che integrano calcoli FEP con reti neurali, consentono di prevedere questo valore con elevata precisione per un gran numero di composti candidati, guidando la scelta delle modifiche chimiche più promettenti da introdurre nella molecola (33).

La selettività è un aspetto altrettanto critico: un farmaco ideale dovrebbe legarsi esclusivamente al target di interesse, senza interagire con altre proteine dell'organismo che potrebbero causare effetti collaterali indesiderati. Raggiungere un buon livello di selettività è spesso difficile perché molte proteine condividono caratteristiche strutturali simili. L'IA contribuisce a questo problema analizzando simultaneamente le caratteristiche strutturali di famiglie di proteine correlate, identificando anche le minime differenze molecolari che possono essere sfruttate per progettare composti selettivi (15).

L'ottimizzazione della potenza riguarda invece la capacità di una molecola di produrre l'effetto biologico desiderato alla concentrazione più bassa possibile. In questo contesto, i modelli di machine learning vengono addestrati su dataset di composti con attività biologica nota per identificare quali modifiche chimiche, come l'aggiunta o la rimozione di gruppi funzionali, la modifica della geometria molecolare o la variazione della lipofilia, siano in grado di migliorare la potenza del composto senza comprometterne la selettività o le proprietà farmacocinetiche (15).

Un aspetto rilevante di questa fase è la natura iterativa del processo, che segue un ciclo noto come *design-make-test-analyze*, ossia si progetta computazionalmente una modifica alla molecola, la si sintetizza in laboratorio, la si testa biologicamente e si usano i risultati per affinare ulteriormente il modello predittivo. Ogni ciclo fornisce nuovi dati che migliorano la qualità delle previsioni successive. L'IA accelera questo processo riducendo il numero di composti da sintetizzare fisicamente, poiché è in grado di filtrare virtualmente un gran numero di candidati e selezionare solo quelli con il profilo farmacologico più promettente, rendendo l'ottimizzazione più rapida ed economica (15,33).

### **3.3 Generazione di nuove molecole (*De Novo Drug Design*)**

La progettazione de novo di farmaci, ovvero la generazione di molecole completamente nuove con proprietà farmacologiche desiderate, rappresenta uno degli ambiti in cui l'IA ha prodotto i progressi più significativi negli ultimi anni. A differenza degli approcci tradizionali, che si limitano a valutare molecole già esistenti o a modificare strutture note, i modelli generativi basati su deep learning sono in grado di esplorare regioni dello spazio chimico non ancora valutate sperimentalmente, concentrandosi su aree con maggiore probabilità di generare composti attivi e sicuri.

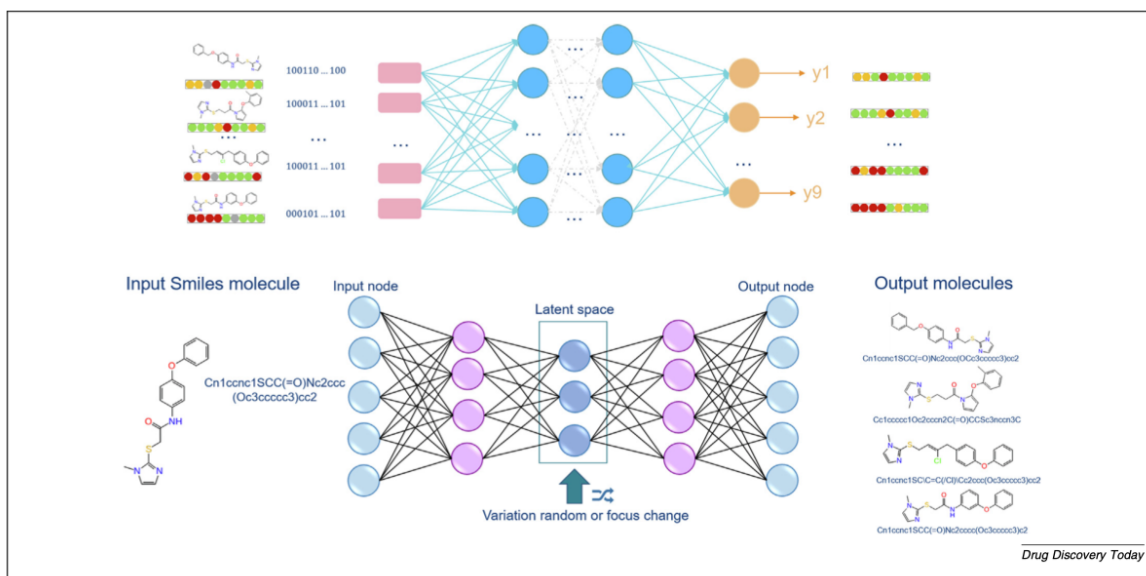
Il processo di “*de novo drug design*” si articola tipicamente in tre fasi principali: la generazione di nuove strutture molecolari, la valutazione delle loro proprietà farmacologiche e la successiva ottimizzazione (29,35).

Infatti, uno dei principali ostacoli del *drug discovery* tradizionale è rappresentato dalla vastità dello spazio chimico teoricamente esplorabile: è stato stimato che il numero di molecole *drug-like* potenzialmente sintetizzabili possa raggiungere circa  $10^{60}$  composti, rendendo impossibile una valutazione esaustiva con metodi sperimentali o computazionali convenzionali. I metodi classici di *de novo design* affrontano questo problema partendo da cataloghi di frammenti molecolari già noti, combinati tra loro seguendo regole chimiche, come le istruzioni per sintetizzare una molecola o i criteri fisico-chimici che una molecola deve soddisfare per essere assorbita dall'organismo. Tuttavia, questi approcci sono limitati dalla conoscenza incompleta delle interazioni molecolari e dalla difficoltà di descrivere l'intera complessità dello spazio chimico farmacologicamente rilevante (28,35).

I modelli di deep learning superano questi limiti: a differenza dei metodi tradizionali di machine learning, non richiedono una definizione manuale delle caratteristiche molecolari rilevanti, ma sono in grado di identificare autonomamente pattern complessi e relazioni non lineari tra struttura chimica e proprietà biologiche direttamente dai dati. Questo consente di orientare la ricerca verso le regioni più promettenti dello spazio chimico in modo molto più efficiente rispetto ai metodi precedenti (28,35).

Tra le architetture più utilizzate nel *de novo drug design* figurano le GAN, *Generative Adversarial Networks* e i modelli basati su *reinforcement learning*. Le GAN funzionano attraverso la competizione tra due reti neurali: una che genera nuove molecole e una che valuta se queste siano realistiche e farmacologicamente rilevanti, migliorando progressivamente la qualità delle strutture generate attraverso questo confronto continuo. Il *reinforcement learning* adotta invece un approccio diverso, in cui il modello impara a generare molecole con proprietà desiderate attraverso un sistema di premi e penalità, ottimizzando progressivamente le strutture generate verso profili farmacologici specifici, come valori ottimali di solubilità, proprietà farmacocinetiche e attività biologica.

Entrambi questi approcci hanno dimostrato risultati promettenti nell'espansione dello spazio chimico esplorabile, sebbene le molecole generate richiedano spesso ulteriori fasi di ottimizzazione sperimentale per quanto riguarda sintetizzabilità, stabilità e sicurezza (29,35) (**Figura 10**).



**Figura 10.** Due esempi di architetture di deep learning applicate alla progettazione molecolare. In alto: un modello che riceve in input le caratteristiche chimiche di una molecola e predice simultaneamente più proprietà farmacologiche. In basso: un modello generativo che, partendo da una molecola esistente, è in grado di generare nuove strutture molecolari con proprietà simili o migliorate (30).

Un'ulteriore evoluzione dei modelli generativi è rappresentata dai *diffusion models*, una classe di modelli di deep learning che imparano a generare nuove strutture molecolari partendo da dati esistenti. In pratica, questi modelli apprendono le caratteristiche che rendono una molecola biologicamente rilevante e usano questa conoscenza per generare nuove strutture con proprietà farmacologiche desiderate, andando oltre la semplice modifica di molecole già note. Questi modelli risultano particolarmente promettenti nella progettazione di librerie di composti innovativi per patologie caratterizzate da meccanismi molecolari complessi, come tumori, malattie neurodegenerative e malattie genetiche rare.

L'integrazione dei *diffusion models* con *RoseTTAFold* ha portato allo sviluppo di *RFdiffusion*, uno strumento per la progettazione de novo di proteine funzionali e anticorpi monoclonali con accuratezza a livello atomico, capace di progettare proteine dirette verso epitopi specifici e di accelerare lo sviluppo di nuovi farmaci biologici. Parallelamente, *AlphaFold 3*, anch'esso basato su modelli diffusivi, è in grado di predire la struttura di grandi complessi biomolecolari, comprese le interazioni tra proteine, acidi nucleici e piccole molecole, ampliando significativamente le possibilità del *de novo drug design* (29).

Un esempio concreto di applicazione di questi approcci è rappresentato da *iDRO - integrated deep learning-based mRNA optimization* - un algoritmo sviluppato da Gong et al. (36) per la progettazione di sequenze di RNA messaggero (mRNA) ottimizzate a partire dalla sequenza

amminoacidica della proteina target. iDRO genera sequenze ottimizzate per massimizzare l'efficienza di traduzione e l'espressione proteica, con applicazioni rilevanti nello sviluppo di vaccini a mRNA. Il modello è stato applicato all'ottimizzazione di sequenze codificanti per la proteina Spike di diverse varianti di SARS-CoV-2, evidenziando il suo potenziale nell'accelerare lo sviluppo di vaccini in risposta a nuove varianti virali e in scenari pandemici futuri (29).

Nel complesso, i modelli generativi basati su deep learning stanno ridefinendo il concetto stesso di progettazione molecolare, spostando il processo da una logica di selezione tra composti esistenti a una logica di creazione guidata di nuove entità chimiche con profili farmacologici ottimizzati. Nonostante i risultati promettenti, molte delle molecole generate richiedono ancora successive fasi di ottimizzazione sperimentale, soprattutto per quanto riguarda sintetizzabilità, stabilità e sicurezza farmacologica, un limite che i modelli attuali stanno progressivamente cercando di superare (28,29,35).

### **3.3.1 Ottimizzazione multiparametrica: ADMET, biodisponibilità, tossicità**

Una delle sfide più complesse nella fase di ottimizzazione del lead compound è la necessità di migliorare simultaneamente più proprietà farmacologiche, spesso in conflitto tra loro. Un composto può essere molto efficace nel legarsi al target, ma risultare tossico, poco solubile o metabolizzato troppo rapidamente dall'organismo. Per questo motivo, l'ottimizzazione del candidato farmacologico non può limitarsi alla sola attività biologica, ma deve tenere conto di un insieme di proprietà note con l'acronimo ADMET, che descrive cinque caratteristiche fondamentali di un farmaco nell'organismo: l'assorbimento (quanto il farmaco viene assorbito dopo la somministrazione), la distribuzione (come si distribuisce nei tessuti), il metabolismo (come viene trasformato dall'organismo), l'escrezione (come viene eliminato) e la tossicità (eventuali effetti dannosi) (15,30).

L'IA ha introdotto in questo contesto la possibilità di effettuare predizioni multitasking, ovvero di stimare simultaneamente più proprietà farmacologiche all'interno di un unico modello. In questo tipo di approccio, le diverse proprietà, come l'attività biologica, la solubilità, la permeabilità, la stabilità metabolica e il profilo di tossicità, vengono valutate contemporaneamente dalla stessa rete neurale. In pratica, la rete elabora le informazioni sulla molecola in modo condiviso per tutte le proprietà da prevedere infine produce una previsione separata per ciascuna proprietà. Perciò, le diverse proprietà si aiutano a vicenda durante

l'apprendimento, migliorando la qualità complessiva delle previsioni rispetto a modelli separati (15).

Questo approccio consente di orientare il processo di progettazione verso composti che presentino il miglior equilibrio possibile tra efficacia e sicurezza farmacologica.

La biodisponibilità, ovvero la frazione di farmaco che raggiunge effettivamente la circolazione sistemica dopo la somministrazione, è una delle proprietà più critiche da ottimizzare. Un composto con ottima attività in vitro può risultare clinicamente inutile se non viene assorbito adeguatamente dall'intestino o se non riesce a raggiungere il tessuto bersaglio in concentrazioni sufficienti. I modelli di machine learning addestrati su dataset di composti con dati farmacocinetici noti sono in grado di prevedere queste proprietà per nuove molecole, guidando le scelte di ottimizzazione chimica prima ancora che il composto venga sintetizzato (15,30).

Uno sviluppo interessante in questo ambito è rappresentato dagli approcci di *inverse QSAR*, che invertono la logica tradizionale dei modelli predittivi: invece di partire dalla struttura di una molecola per prevederne le proprietà, questi modelli partono dalle proprietà desiderate, come un certo livello di solubilità, una bassa tossicità e una buona permeabilità, e generano le strutture molecolari che le soddisfano. Questo approccio rappresenta un passo ulteriore verso una progettazione molecolare guidata dalle proprietà farmacologiche desiderate, piuttosto che dalla struttura chimica di partenza (15).

La predizione della tossicità è un altro ambito in cui l'IA ha dimostrato un impatto significativo nell'ottimizzazione del candidato farmacologico. I modelli predittivi possono identificare precocemente pattern strutturali associati a tossicità epatica, cardiaca o genotossicità, consentendo di escludere i composti a rischio già nelle fasi iniziali del *drug discovery*. Questo non solo riduce i costi e i tempi dello sviluppo, ma contribuisce anche a migliorare la sicurezza dei candidati farmacologici che arrivano alla sperimentazione clinica. Bisogna però sottolineare che la predizione della tossicità rimane uno degli ambiti più difficili per l'IA, poiché gli effetti tossici di un composto dipendono da molti fattori che vanno oltre la sola struttura chimica, come la dose, il profilo genetico del paziente e le eventuali interazioni con altri farmaci, rendendo difficile la raccolta di dati sufficientemente completi e affidabili per addestrare modelli predittivi di alta qualità (15,28).

Un esempio concreto di applicazione industriale di questi approcci è rappresentato dalla piattaforma ADMET sviluppata dall'azienda farmaceutica Bayer e descritta da Göller et al. (37), sviluppata nel corso di oltre due decenni con l'obiettivo di creare modelli predittivi per

le principali proprietà ADMET, tra cui il rischio di interazione con il canale hERG, uno dei bersagli di sicurezza più monitorati in fase preclinica per il suo ruolo nel rischio di aritmie cardiache. Questi modelli sono stati integrati direttamente nei software utilizzati da ricercatori durante la progettazione molecolare, in modo che per ogni nuova molecola proposta vengano calcolate automaticamente le previsioni ADMET, rendendo possibile una valutazione immediata del profilo di sicurezza senza dover condurre esperimenti separati. La piattaforma di Bayer rappresenta quindi un esempio concreto di come l'IA possa essere incorporata nei processi quotidiani di un'azienda farmaceutica, accelerando le decisioni di selezione dei candidati e riducendo il numero di esperimenti necessari nelle fasi precliniche. Un elemento particolarmente significativo di questa piattaforma è il sistema di aggiornamento automatico dei modelli: per diversi endpoint chiave, come la stabilità microsomiale e la stabilità negli epatociti, i modelli vengono riaddestrati automaticamente ogni settimana incorporando i nuovi dati sperimentali prodotti internamente, garantendo che le previsioni rimangano aggiornate rispetto alle molecole più recenti in sviluppo (37).

Nel complesso, l'ottimizzazione multiparametrica assistita da IA sta migliorando significativamente l'efficienza delle fasi precliniche del *drug discovery*, riducendo il numero di composti da sintetizzare e testare sperimentalmente e aumentando la probabilità che i candidati farmacologici selezionati abbiano successo nelle fasi successive dello sviluppo clinico (15,30).

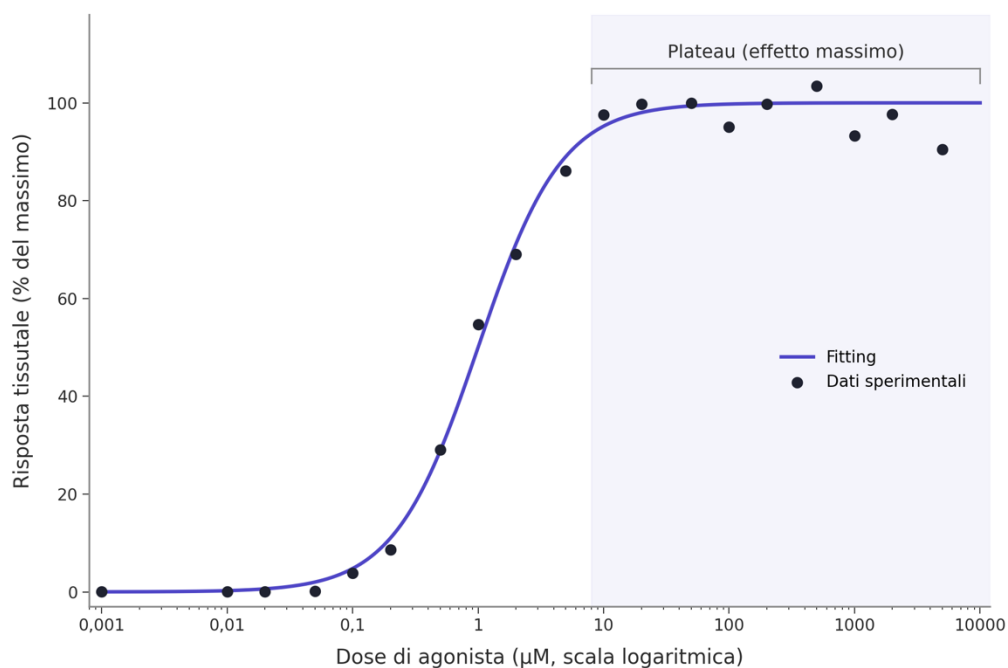
### **3.3.2 Modelli PK/PD integrati con IA e simulazioni in silico di efficacia terapeutica**

La farmacodinamica studia gli effetti di un farmaco sull'organismo e i meccanismi molecolari attraverso cui questi effetti si producono. I modelli PK/PD integrano la farmacodinamica con la farmacocinetica (discussa in precedenza) attraverso equazioni matematiche che mettono in relazione la concentrazione del farmaco nel sangue o nel tessuto bersaglio, descritta da parametri farmacocinetici come l'area sotto la curva, il volume di distribuzione e la clearance, con la risposta biologica osservata, consentendo di prevedere l'effetto terapeutico in funzione della dose e del tempo. Questo approccio è fondamentale nelle fasi di sviluppo farmacologico per ottimizzare il regime di dosaggio, identificare la finestra terapeutica e ridurre il rischio di tossicità o inefficacia clinica (30).

L'IA ha introdotto in questo contesto strumenti in grado di costruire modelli PK/PD più accurati e personalizzati rispetto a quelli tradizionali. I modelli classici si basano su

equazioni matematiche che descrivono il comportamento medio del farmaco in una popolazione, ma non tengono conto della variabilità individuale che comporta tra le altre cose differenze genetiche, età, peso corporeo, funzionalità renale ed epatica e che influenzano in modo significativo la risposta al trattamento. I modelli basati su machine learning possono invece integrare queste variabili individuali, costruendo previsioni personalizzate sulla risposta al farmaco per ciascun paziente (29,30).

La curva dose-risposta (**Figura 11**) descrive la relazione tra la dose di un farmaco somministrata e l'intensità dell'effetto biologico prodotto. Definire questa relazione è fondamentale per stabilire la dose terapeutica ottimale, identificare la soglia di tossicità e stimare l'indice terapeutico, ovvero il margine di sicurezza tra la dose efficace e quella tossica. I metodi tradizionali per definire queste curve richiedono numerosi esperimenti su modelli cellulari e animali, con costi e tempi significativi (29).



**Figura 11.** Curva dose-risposta sigmoideale che descrive la relazione tra la concentrazione di un farmaco e l'intensità dell'effetto biologico prodotto. All'aumentare della dose, la risposta cresce progressivamente fino a raggiungere un plateau che rappresenta l'effetto massimo raggiungibile.

I modelli di machine learning e deep learning consentono di predire le curve dose-risposta a partire dalle caratteristiche chimiche e biologiche del composto, riducendo il numero di esperimenti necessari. In particolare, modelli addestrati su dataset di composti con dati di attività biologica nota sono in grado di stimare come un nuovo composto si comporterà a

diverse concentrazioni, orientando la scelta delle dosi da testare sperimentalmente. Questo approccio è particolarmente utile nelle fasi precliniche, dove permette di concentrare le risorse sperimentali sulle finestre di dosaggio più promettenti (29,30).

Un aspetto particolarmente rilevante riguarda la capacità dell'IA di modellare le relazioni tra le proprietà chimiche di una molecola, la sua farmacocinetica e la farmacodinamica. Proprietà come la solubilità, la lipofilia, il legame alle proteine plasmatiche e la stabilità metabolica non influenzano il comportamento del farmaco in modo indipendente l'una dall'altra, ma interagiscono tra loro in modi complessi che i modelli matematici tradizionali faticano a rappresentare. I modelli di deep learning sono in grado di catturare queste interazioni in modo più completo, migliorando la capacità predittiva rispetto ai metodi convenzionali (30).

L'IA trova applicazione anche nella predizione degli effetti farmacodinamici, ovvero nella stima dell'intensità e della durata della risposta biologica al farmaco in funzione della sua concentrazione nel sito d'azione. Integrando dati farmacocinetici e farmacodinamici in un unico modello, l'IA consente di simulare scenari clinici complessi, come la risposta in popolazioni di pazienti con comorbidità o con polimorfismi genetici che alterano il metabolismo del farmaco, prima ancora che gli studi clinici vengano avviati. (29,30)

Un'applicazione concreta di questi modelli integrati si può vedere nella previsione delle interazioni tra farmaci, in particolare nelle situazioni in cui la somministrazione simultanea di più farmaci va ad alterare il profilo farmacocinetico di uno o più dei composti coinvolti, con potenziali conseguenze sulla sicurezza e sull'efficacia della terapia. I modelli di IA addestrati su dataset di interazioni note sono in grado di identificare precocemente potenziali interazioni problematiche, supportando le decisioni cliniche e riducendo il rischio di eventi avversi sia nelle fasi di sviluppo del farmaco che nella pratica clinica (29,30).

Le simulazioni *in silico*, ovvero le simulazioni condotte interamente al computer, senza ricorrere a esperimenti biologici, consentono di simulare virtualmente il comportamento di un farmaco in un sistema biologico complesso, integrando informazioni sulla struttura molecolare del composto, sul meccanismo d'azione del target e sulla fisiologia del paziente. Nelle fasi precliniche, queste simulazioni possono essere utilizzate per stimare l'efficacia di un composto in modelli di malattia virtuali, riducendo il numero necessario di esperimenti su animali. Nelle fasi cliniche, consentono di simulare la risposta di popolazioni virtuali di pazienti con caratteristiche demografiche e genetiche diverse, migliorando il disegno degli

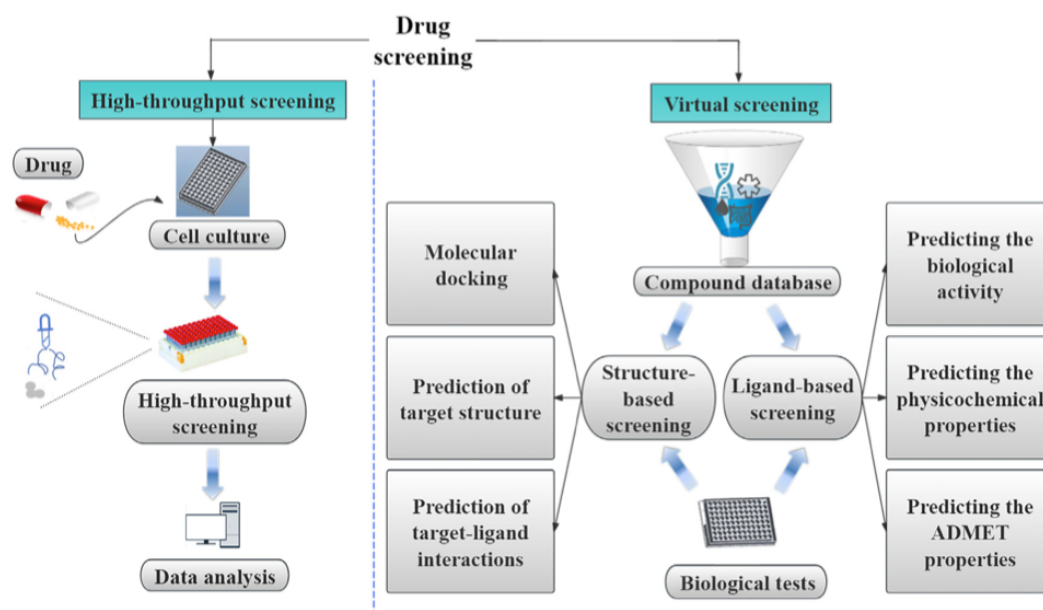
studi clinici e aumentando la probabilità di identificare il dosaggio ottimale per diverse sottopopolazioni di pazienti (29,30).

Un esempio rilevante in questo ambito è rappresentato dai modelli di simulazione della popolazione, noti come *population PK/PD models*, che utilizzano dati farmacocinetici e farmacodinamici raccolti su gruppi di pazienti per costruire modelli matematici della variabilità interindividuale nella risposta al farmaco. L'integrazione di questi modelli con algoritmi di machine learning consente di identificare i fattori che maggiormente influenzano questa variabilità, come l'età, il peso corporeo, la funzionalità renale o la presenza di specifici polimorfismi genetici, e di utilizzare queste informazioni per personalizzare il regime terapeutico per ciascun paziente (30).

Nel complesso, l'integrazione tra IA e modellistica PK/PD rappresenta uno degli strumenti più promettenti per rendere il processo di sviluppo farmacologico più efficiente, più sicuro e più orientato verso terapie personalizzate. Come per gli altri approcci computazionali descritti in questo capitolo, però, i risultati delle simulazioni *in silico* devono sempre essere verificati sperimentalmente, e la loro affidabilità dipende dalla qualità e dalla completezza dei dati su cui i modelli sono stati addestrati (29,30).

### **3.4 Ottimizzazione del processo di screening: virtual screening e supporto all'High-Throughput Screening (HTS)**

Lo screening di grandi librerie di composti chimici alla ricerca di molecole con attività biologica promettente rappresenta una delle fasi più costose e laboriose del *drug discovery*. L'*High-Throughput Screening* (HTS), ovvero la sperimentazione ad alta velocità di migliaia o milioni di composti in sistemi biologici automatizzati, ha rappresentato per decenni il principale strumento per questa fase, ma comporta costi elevati e tempi lunghi. Il *Virtual screening* affronta questo problema spostando una parte dello screening dal laboratorio in digitale: invece di testare fisicamente ogni composto, si utilizzano modelli computazionali per selezionare preventivamente i candidati più promettenti, riducendo il numero di esperimenti necessari e ampliando lo spazio chimico esplorabile oltre i limiti delle librerie fisicamente disponibili (15,29) (**Figura 12**).



**Figura 12.** Workflow del virtual screening assistito da IA, con le due categorie principali: il ligand-based virtual screening, che parte dalle caratteristiche chimiche di molecole già note, e lo structure-based virtual screening, che utilizza la struttura tridimensionale del target per identificare i composti candidati più promettenti (38).

I modelli di deep learning sono in grado di elaborare grandi dataset di composti con attività biologica nota, imparando a riconoscere i pattern chimici associati all'efficacia farmacologica e applicando questa conoscenza per valutare rapidamente nuovi composti sulla base delle loro proprietà chimico-fisiche, dell'attività biologica e dei profili farmacocinetici. In particolare, la predizione delle interazioni farmaco-target, nota con l'acronimo DTI, dall'inglese *Drug-Target Interaction*, è diventata uno degli ambiti centrali del Virtual screening assistito da IA, consentendo di stimare la probabilità che una molecola si leghi a un determinato bersaglio biologico prima ancora di sintetizzarla o testarla sperimentalmente (29).

Un aspetto rilevante di questi approcci è la loro capacità di integrare dati eterogenei provenienti da fonti diverse, informazioni chimiche sulla struttura delle molecole, dati biologici sui target, e dati clinici sugli effetti dei farmaci, per migliorare la qualità delle previsioni. Questo è particolarmente importante nel *drug discovery*, dove nessun tipo di dato è sufficiente da solo a descrivere la complessità delle interazioni biologiche. L'integrazione di più livelli informativi consente ai modelli di sfruttare simultaneamente differenti prospettive sulla stessa interazione farmaco-target, aumentando l'accuratezza predittiva e riducendo il rischio di falsi positivi nelle fasi di screening (29).

Parallelamente alla valutazione di molecole esistenti, i modelli di IA possono essere impiegati anche per la progettazione *de novo* di composti destinati a bersagli terapeutici specifici, contribuendo così non solo alla selezione ma anche alla generazione di nuovi candidati farmacologici con caratteristiche ottimizzate. Questo approccio integrato che combina *Virtual screening* e generazione molecolare, rappresenta uno dei paradigmi più promettenti del *drug discovery* computazionale moderno (15,29).

Un esempio concreto di questo approccio è rappresentato dal lavoro di Li et al. (39), che hanno sviluppato un modello generativo basato su una *rete neurale ricorrente condizionale*, un tipo di rete neurale progettata per generare nuove strutture molecolari non in modo casuale, ma orientandole verso un obiettivo preciso, in questo caso la capacità di legarsi e inibire il target RIPK1, una proteina coinvolta nei processi di necroptosi e infiammazione, considerata un target terapeutico promettente per diverse patologie infiammatorie e neurodegenerative, come la colite ulcerosa e il morbo di Crohn. Il modello è stato addestrato su un dataset di molecole note per apprenderne le caratteristiche chimiche, e ha generato 79,323 molecole potenzialmente attive nei confronti di RIPK1. Tra i composti individuati, un inibitore selettivo e particolarmente potente ha mostrato attività biologica nella riduzione dei processi di necroptosi e infiammazione sia in modelli cellulari in vitro che in modelli animali in vivo (29).

Un secondo esempio rilevante è rappresentato da *AI-Bind*, una pipeline sviluppata da Chatterjee et al. (40) per la predizione delle interazioni di legame tra proteine e ligandi. AI-Bind utilizza rappresentazioni apprese in modo non supervisionato su grandi librerie di composti chimici e sequenze proteiche, consentendo al modello di generalizzare le sue previsioni anche su proteine e molecole non presenti nel dataset di addestramento, una caratteristica particolarmente importante in contesti di emergenza, dove i target biologici possono essere nuovi e poco caratterizzati. Questo approccio è stato applicato allo screening di farmaci già esistenti con l'obiettivo di identificare molecole in grado di interagire con proteine del virus SARS-CoV-2 o con proteine umane rilevanti nel trattamento del COVID-19, con risultati validati attraverso simulazioni di docking molecolare ed esperimenti in vitro (29), rappresentando quindi una fase preliminare promettente che richiederebbe ulteriori studi clinici per confermare l'efficacia terapeutica dei composti identificati.

Un ulteriore strumento emergente nel supporto allo screening è il *morphological profiling*, un approccio che utilizza il deep learning per analizzare le alterazioni morfologiche delle cellule in risposta alla somministrazione di un composto. Invece di misurare un singolo parametro biochimico, questo metodo cattura il profilo morfologico completo della cellula,

forma, dimensione, distribuzione degli organelli, e utilizza queste informazioni per classificare il meccanismo d'azione del composto, identificare potenziali nuove indicazioni terapeutiche e supportare le decisioni di screening nelle fasi iniziali del *drug discovery*. Grazie all'imaging ad alta risoluzione e agli algoritmi di computer vision, è possibile analizzare automaticamente migliaia di immagini cellulari per estrarre informazioni farmacologicamente rilevanti che sarebbero difficili da ottenere con metodi biochimici tradizionali. Questo approccio si rivela particolarmente utile nelle fasi iniziali dello screening, dove permette di ottenere informazioni ricche sul comportamento biologico di un composto senza richiedere saggi biochimici specifici, e trova applicazione anche nel *drug repurposing*, dove i profili morfologici possono rivelare analogie funzionali tra composti con strutture chimiche diverse (41).

È importante sottolineare che, nonostante i progressi significativi, le predizioni computazionali ottenute tramite *Virtual screening* necessitano sempre di conferma sperimentale mediante saggi in vitro e in vivo (15,29).

### **3.5 Drug Repurposing: algoritmi per l'identificazione di nuove indicazioni terapeutiche per farmaci già approvati**

Il *drug repurposing*, ovvero il riposizionamento di farmaci già approvati verso nuove indicazioni terapeutiche, rappresenta uno degli approcci più promettenti per accelerare lo sviluppo di nuove terapie riducendo tempi e costi. Poiché i farmaci già approvati hanno già dimostrato di essere sicuri per l'uomo nelle dosi terapeutiche, il loro riposizionamento verso una nuova indicazione non richiede nuovamente tutti gli studi di sicurezza e tossicità, una delle fasi più lunghe e costose dello sviluppo farmacologico. È possibile perciò, in genere, direttamente dimostrare l'efficacia per la nuova indicazione (29,30).

L'IA ha ampliato significativamente le possibilità del *drug repurposing*, introducendo strumenti in grado di analizzare grandi reti di interazioni biologiche e farmaco-target per identificare associazioni non ancora note tra molecole esistenti e patologie diverse da quelle per cui sono state sviluppate. Un approccio particolarmente utilizzato si basa sull'analisi di reti biologiche complesse, come i *Knowledge Graph*, nelle quali farmaci, proteine e malattie vengono rappresentati come nodi interconnessi. Attraverso tecniche di machine learning che analizzano la struttura di queste reti, invece di esaminare le molecole singolarmente, è possibile identificare quali farmaci già esistenti sono strettamente collegati a target biologici

coinvolti in malattie diverse da quelle per cui sono stati sviluppati, suggerendo così nuove possibili indicazioni terapeutiche da verificare sperimentalmente (29,30).

I modelli di deep learning applicati a queste reti permettono di integrare informazioni eterogenee, dati chimici, genomici, fenotipici e clinici - migliorando la capacità predittiva dei modelli rispetto agli approcci tradizionali. In particolare, gli algoritmi di *network propagation* consentono di amplificare segnali deboli di associazione tra nodi della rete, facilitando l'identificazione di target potenzialmente rilevanti per nuove indicazioni terapeutiche anche quando le evidenze disponibili sono limitate. Parallelamente, i *Knowledge graph* integrano in un'unica struttura reticolare informazioni provenienti da fonti diverse, come articoli scientifici, database biologici e dati sperimentali, collegando tra loro entità come farmaci, geni, proteine e malattie attraverso relazioni etichettate che descrivono il tipo di connessione tra di esse. In questo modo è possibile interrogare la rete per scoprire percorsi indiretti tra un farmaco e una nuova malattia che non sarebbero emersi analizzando le singole fonti separatamente (30).

Un aspetto rilevante è che la capacità dei modelli di IA di elaborare rapidamente grandi quantità di molecole trova nel *drug repurposing* un'applicazione altrettanto importante. In particolare, l'integrazione tra modelli di deep learning e metodologie computazionali chimico-fisiche consente di effettuare screening virtuali su enormi librerie di farmaci già approvati, valutando rapidamente la loro potenziale attività su nuovi target. Questo approccio è particolarmente sviluppato nell'ambito dei farmaci a basso peso molecolare, per i quali esistono già grandi quantità di dati strutturali e farmacocinetici accumulati nel corso degli anni, informazioni su come sono fatti chimicamente, come vengono assorbiti e metabolizzati dall'organismo, che consentono di addestrare modelli predittivi più accurati e affidabili (29). Come sottolineato da Bender e Cortés-Ciriano (28), questo vantaggio è sostanziale: la qualità e la quantità dei dati disponibili è uno dei fattori determinanti per l'affidabilità dei modelli di IA nel *drug discovery*, e le piccole molecole beneficiano di decenni di ricerca che hanno prodotto database come ChEMBL e PubChem. I farmaci biologici, come anticorpi e proteine terapeutiche, rappresentano invece una categoria più recente e più complessa da caratterizzare, con meno dati storici disponibili, il che spiega perché i modelli di IA per questa categoria siano ancora meno maturi rispetto a quelli sviluppati per le piccole molecole (28).

Un contributo rilevante al *drug repurposing* viene anche dalla rivalutazione di target molecolari precedentemente abbandonati o comunque considerati poco promettenti. L'IA consente di riesaminare questi target integrando nuove evidenze biologiche accumulate nel

tempo, come nuovi dati genomici provenienti da studi su popolazioni più ampie, dati proteomici che descrivono meglio il comportamento della proteina in contesti patologici specifici, o dati clinici che rivelano associazioni tra quel target e determinate malattie. La capacità dell'IA di analizzare simultaneamente questi diversi livelli informativi, che corrisponde ad una quantità di dati che sarebbe impossibile da elaborare manualmente, permette di identificare pattern e connessioni che erano presenti nei dati ma che non erano stati ancora riconosciuti. In questo senso, l'IA amplifica la capacità dei ricercatori di estrarre informazioni utili da grandi volumi di dati, orientando la sperimentazione verso i candidati che mostrano il profilo più promettente e riducendo così il rischio di investire risorse su molecole destinate a fallire nelle fasi successive dello sviluppo (29,30).

Il *morphological profiling* trova nel *drug repurposing* una delle sue applicazioni più rilevanti. Il principio alla base di questo approccio è il seguente: se due composti diversi inducono alterazioni morfologiche simili nelle cellule, per esempio modifiche analoghe alla forma del nucleo, alla distribuzione dei mitocondri o alla struttura del citoscheletro, è probabile che agiscano attraverso meccanismi biologici simili, anche se la loro struttura chimica è completamente diversa. Se un farmaco già approvato produce nelle cellule le stesse alterazioni morfologiche, stessa forma del nucleo, stessa distribuzione degli organelli, stessa struttura del citoscheletro, di una molecola nota per essere attiva su un determinato target, si ipotizza che i due composti stiano interferendo con gli stessi processi biologici. Su questa base, il farmaco già approvato potrebbe essere attivo sullo stesso target, suggerendo una nuova possibile indicazione terapeutica da verificare sperimentalmente. In pratica, il deep learning analizza automaticamente migliaia di immagini cellulari, costruisce un profilo morfologico per ciascun composto e confronta questi profili tra loro, identificando somiglianze che un ricercatore non riuscirebbe mai a rilevare manualmente su una scala così ampia. Questo consente di scoprire connessioni tra farmaci e malattie che non emergerebbero né dalla sola struttura chimica, che può essere completamente diversa tra due composti con effetti biologici simili, né dai saggi biologici tradizionali, che misurano tipicamente un singolo parametro alla volta e non catturano il quadro complessivo degli effetti del composto sulla cellula (41).

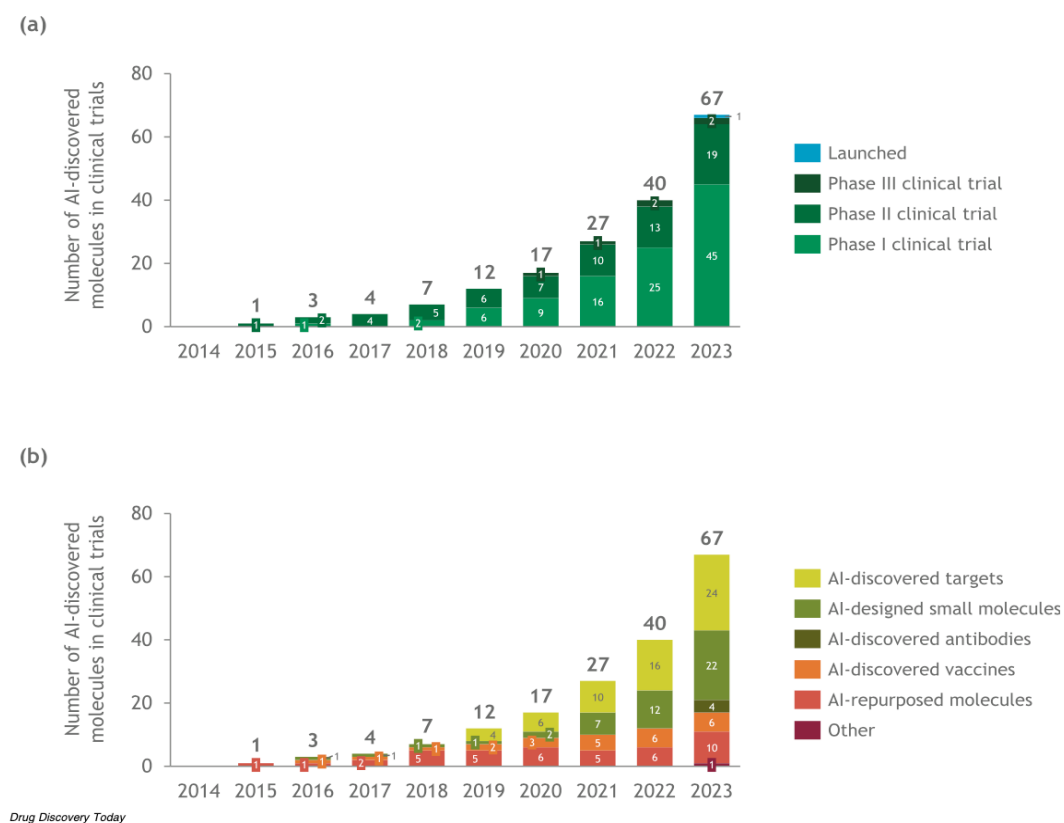
È tuttavia importante considerare i limiti di questi approcci. Come sottolineato da Bender e Cortés-Ciriano (28), i modelli predittivi applicati al *drug repurposing* si basano spesso su dati provenienti da fonti molto diverse, laboratori, ospedali e database clinici con protocolli sperimentali, strumenti di misura e criteri diagnostici differenti, che producono informazioni inconsistenti e difficilmente confrontabili. Questo rende difficile addestrare modelli

affidabili, perché l'IA utilizza i dati che le vengono forniti e non può compensare da sola le incongruenze presenti nelle fonti. Inoltre, gli effetti biologici di un farmaco dipendono da molti fattori contestuali, come la dose somministrata, il profilo genetico del paziente e le eventuali interazioni con altri farmaci, che raramente vengono registrati in modo sistematico nei dataset disponibili, perché la maggior parte di questi dati è stata raccolta per scopi clinici o sperimentali specifici e non con l'obiettivo di addestrare modelli di IA. I modelli computazionali non riescono quindi a tenere conto di questi fattori, semplicemente perché non hanno queste informazioni durante l'addestramento. Questo significa che le previsioni di *repurposing* generate dall'IA devono essere interpretate come ipotesi da verificare sperimentalmente (28).

Nel complesso, il *drug repurposing* assistito da IA rappresenta una strategia sempre più centrale nel *drug discovery* moderno, capace di valorizzare il patrimonio di farmaci già sviluppati, di rivalutare target molecolari già noti e di accelerare l'identificazione di nuove strategie terapeutiche, riducendo significativamente la dipendenza da processi sperimentali lunghi e costosi. Tuttavia, è importante non sopravvalutare questi risultati: le previsioni generate dai modelli di IA rappresentano ipotesi di lavoro che richiedono una validazione sperimentale rigorosa (28–30).

## 4. L'INTELLIGENZA ARTIFICIALE NEL PROCESSO DI *DRUG DEVELOPMENT*

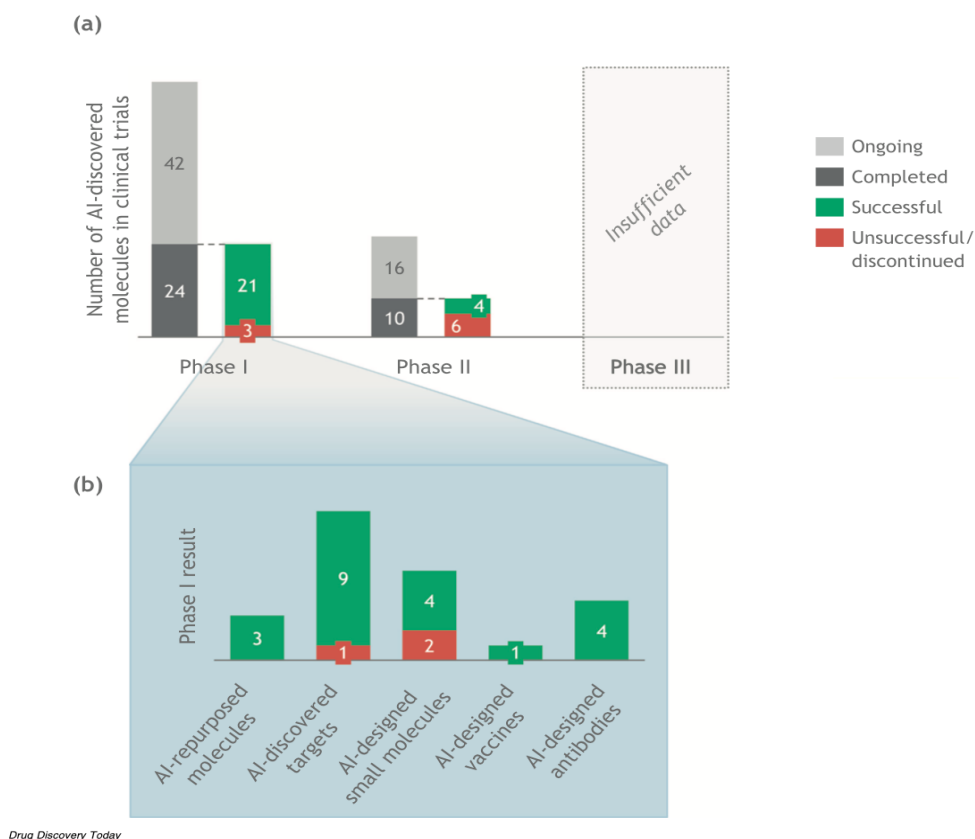
Il capitolo precedente ha illustrato come l'IA stia trasformando le fasi precoci della scoperta del farmaco, questo capitolo si concentra invece su una fase altrettanto critica e spesso ancora più costosa, ossia lo sviluppo clinico del farmaco. Infatti, una volta identificato un candidato farmacologico promettente, il percorso che lo porta dall'identificazione alla disponibilità sul mercato è lungo, complesso e caratterizzato da un tasso di fallimento molto elevato. Si stima che meno del 10% dei composti che entrano in fase clinica riesca effettivamente ad ottenere un'approvazione regolatoria, a fronte di costi che possono superare in alcuni casi il miliardo di dollari per ogni farmaco sviluppato (42) (**Figure 13 e 14**).



**Figura 13.** Numero di molecole scoperte dall'IA che sono entrate negli studi clinici. (a) molecole scoperte dall'IA in fase clinica. (b) molecole scoperte dall'IA in base alla categoria delle molecole (42).

In questo contesto, l'IA sta assumendo un ruolo sempre più rilevante, non solo nelle fasi precliniche dove supporta le fasi precoci della ricerca aiutando nella predizione delle proprietà ADMET e delle tossicità, ma anche nella progettazione e gestione degli studi

clinici, nell'identificazione di biomarcatori predittivi e nel monitoraggio della sicurezza post-marketing. La capacità di elaborare grandi volumi di dati eterogenei, di identificare pattern non evidenti con i metodi tradizionali e di costruire modelli predittivi personalizzati sta cambiando profondamente il modo in cui i farmaci vengono sviluppati, rendendolo un processo più efficiente, sicuro e orientato verso le esigenze specifiche del paziente, portando così ad una medicina sempre più personalizzata (29).



**Figura 14.** Successo delle molecole scoperte con l'IA negli studi clinici. (a) successi clinici delle molecole scoperte dall'IA in fase clinica. (b) molecole scoperte con l'IA che hanno completato la fase I degli studi clinici, in base alla categoria (42).

## 4.1 Studi preclinici

Gli studi preclinici rappresentano la prima fase del processo di *drug development* e hanno l'obiettivo di valutare la sicurezza e l'efficacia di un composto candidato per diventare farmaco, prima della fase clinica. In questa fase vengono raccolte informazioni fondamentali sul profilo farmacocinetico e farmacodinamico del composto, sulla sua tossicità e sul suo comportamento in sistemi biologici complessi. Tradizionalmente, questa fase si basa su una combinazione di esperimenti in vitro, condotti su cellule o tessuti in laboratorio, ed

esperimenti in vivo, condotti su modelli animali, con costi elevati, tempi lunghi e implicazioni etiche legate all'utilizzo degli animali. L'IA sta offrendo strumenti sempre più potenti per supportare e in alcuni casi ridurre la necessità di questi esperimenti, migliorando la qualità e la velocità degli studi preclinici (29,43).

La rilevanza di questi strumenti di IA è confermata dai dati, considerando che si stima che circa il 30% dei composti candidati preclinici fallisca a causa di problemi di tossicità, mentre circa il 40% fallisca per un profilo ADMET inadeguato (43). Le reazioni tossiche impreviste, non identificate nelle fasi precedenti, rappresentano inoltre una delle principali cause di ritiro dei farmaci dal mercato dopo l'approvazione (44); la quota di farmaci approvati che viene effettivamente ritirata per motivi di sicurezza è tuttavia contenuta e, secondo analisi recenti condotte su agenzie regolatorie diverse, si attesta intorno al 3%, con una tendenza alla diminuzione (45).

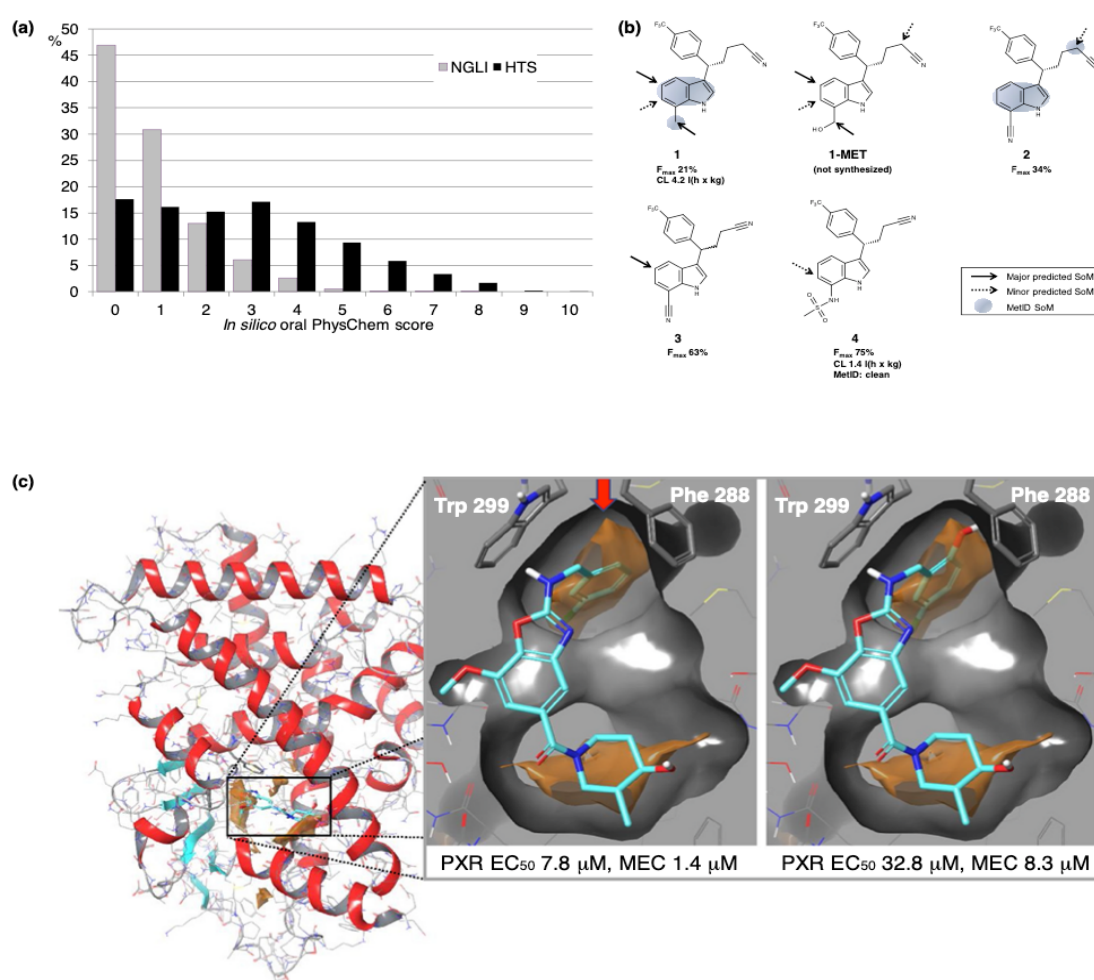
#### **4.1.1 Predizione delle proprietà ADMET: metabolismo epatico, interazioni farmacologiche e tossicità**

La valutazione delle proprietà ADMET, assorbimento, distribuzione, metabolismo, escrezione e tossicità, rappresenta uno dei passaggi più critici della fase preclinica. Un composto può dimostrarsi efficace in vitro su un determinato target, ma risultare inutilizzabile clinicamente se viene metabolizzato troppo rapidamente, se non raggiunge il tessuto bersaglio in concentrazioni sufficienti o se produce effetti tossici inaccettabili. Per questo motivo, la predizione precoce delle proprietà ADMET è fondamentale per ridurre il numero di composti da testare nelle fasi successive dello sviluppo, minimizzando i fallimenti (43).

Uno degli aspetti più rilevanti in questo contesto riguarda il metabolismo epatico mediato dal citocromo P450 (CYP450), una famiglia di enzimi presente principalmente a livello epatico che è responsabile del metabolismo della maggior parte dei farmaci. Le isoforme più rilevanti dal punto di vista farmacologico includono CYP3A4, CYP2D6, CYP2C9 e CYP2C19, ciascuna delle quali metabolizza un ampio numero di farmaci appartenenti a classi terapeutiche diverse. La capacità di un farmaco di essere substrato, inibitore o induttore di queste isoforme determina non solo la sua velocità di eliminazione, ma anche il rischio di interazioni farmacologiche clinicamente significative quando il paziente assume più farmaci contemporaneamente. I modelli di IA addestrati su grandi dataset di composti

per i quali è già stata misurata sperimentalmente la capacità di interagire con le isoforme del CYP sono in grado di predire con buona accuratezza il profilo metabolico di nuovi candidati farmacologici, consentendo di identificare precocemente i composti a rischio di interazioni farmacologiche clinicamente rilevanti (37,43).

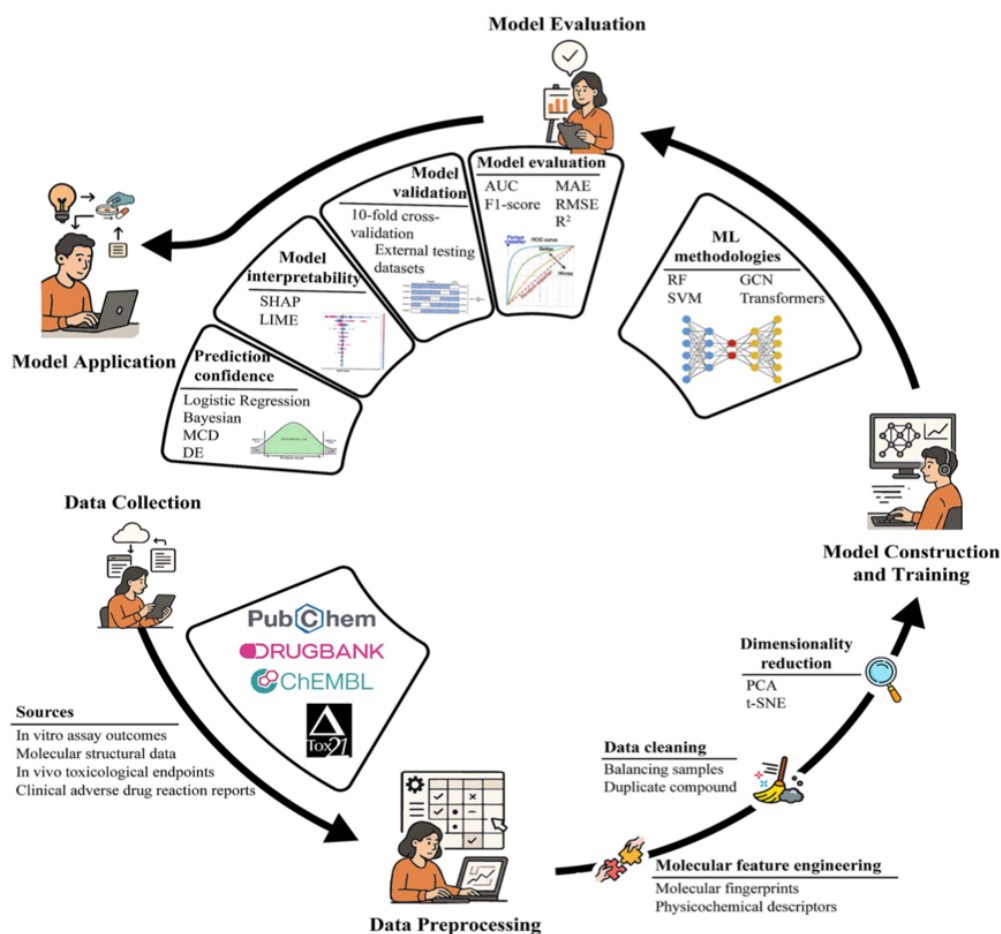
Un esempio concreto di applicazione industriale di questi approcci è rappresentato dalla piattaforma ADMET sviluppata dall'azienda farmaceutica Bayer, descritta da Göller et al. (37), che integra modelli di machine learning per la predizione automatica delle proprietà ADMET direttamente nei software di progettazione molecolare (**Figura 15**), riducendo il numero di esperimenti necessari nelle fasi precliniche, caso che è stato già approfondito nel *capitolo 2*.



Drug Discovery Today

**Figura 15.** Esempi dell'applicazione della piattaforma ADMET sviluppata da Bayer per l'ottimizzazione delle proprietà farmacocinetiche dei candidati farmaci (37).

Un ambito di particolare rilevanza clinica e farmacologica riguarda la predizione della tossicità epatica indotta da farmaci, nota con l'acronimo DILI (*Drug-Induced Liver Injury*), e della cardiotoxicità, in particolare il prolungamento dell'intervallo QT (**Figura 16**).



**Figura 16.** Workflow della predizione della tossicità guidata dall'IA (43).

Il DILI si verifica quando un farmaco o i suoi metaboliti danneggiano direttamente le cellule epatiche o innescano una risposta immunitaria che porta alla loro distruzione. Rappresenta una delle principali cause di ritiro di farmaci dal mercato e di fallimento nelle fasi avanzate dello sviluppo clinico, e la sua predizione è notoriamente difficile perché i meccanismi attraverso cui si manifesta sono molteplici e perché la suscettibilità individuale varia significativamente in base al profilo genetico, all'età e alla presenza di altre patologie. I modelli di deep learning addestrati su grandi dataset di composti per i quali è già stata documentata sperimentalmente la tossicità epatica sono in grado di riconoscere caratteristiche chimiche strutturali ricorrenti nei composti che causano DILI, consentendo di escludere precocemente i candidati a maggior rischio prima ancora di condurre esperimenti in vitro o in vivo (43).

Analogamente, il prolungamento dell'intervallo QT rappresenta un altro parametro di sicurezza critico nelle fasi precliniche. L'intervallo QT è una misura elettrocardiografica che descrive il tempo necessario al cuore per completare un ciclo di contrazione e rilassamento ventricolare. Quando questo intervallo si allunga oltre certi valori a causa del legame di un farmaco con il canale ionico hERG, il rischio di sviluppare aritmie ventricolari potenzialmente fatali aumenta significativamente. Per questo motivo, la valutazione del rischio di interazione con hERG è diventata uno degli endpoint di sicurezza obbligatori nelle fasi precliniche per qualsiasi nuovo candidato farmacologico. I modelli computazionali basati su IA consentono di predire questo rischio direttamente dalla struttura chimica della molecola, riducendo significativamente il numero di composti che arrivano in fase clinica con problemi di sicurezza cardiaca non identificati e che potrebbero causare danni gravi ai pazienti durante la sperimentazione (30,43).

#### **4.1.2 Sviluppo di modelli di tossicologia predittiva per ridurre i test su animali**

La sperimentazione su animali rappresenta uno dei pilastri della fase preclinica dello sviluppo farmacologico. Prima che un farmaco possa essere somministrato all'uomo, è necessario dimostrarne la sicurezza in modelli animali, valutando la tossicità acuta e cronica, la genotossicità, la teratogenicità, cioè il rischio di causare malformazioni nello sviluppo fetale, e altri parametri di sicurezza.

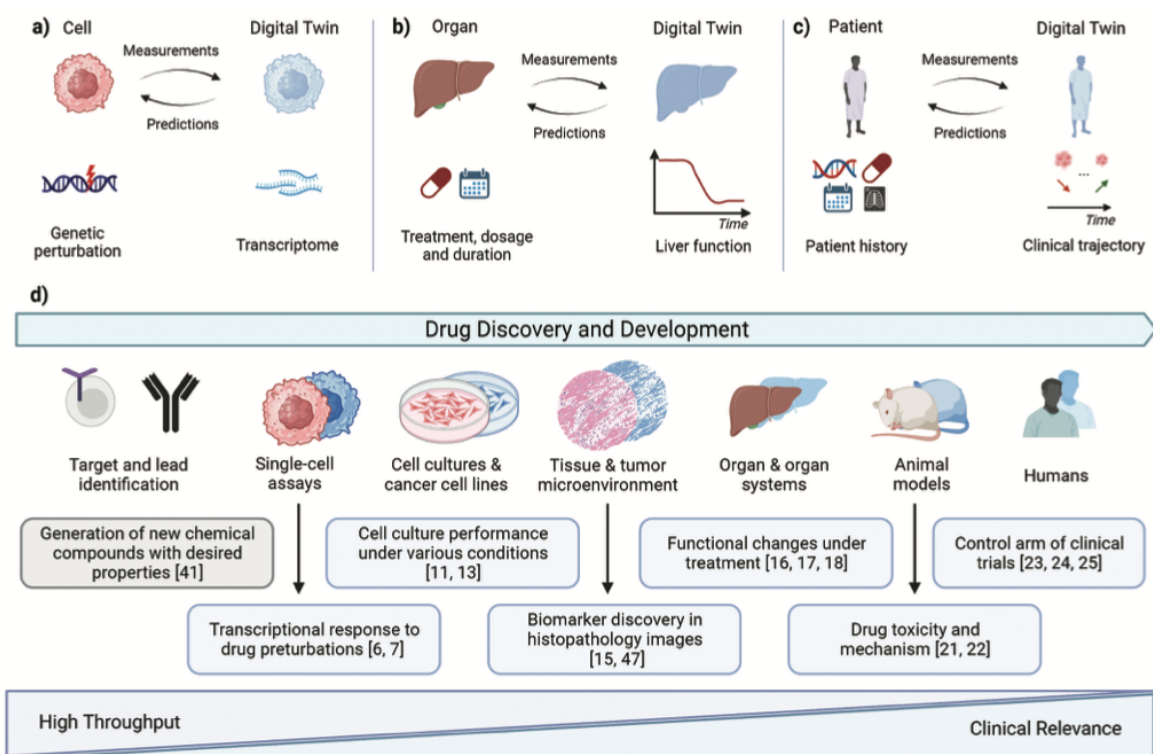
Tuttavia, i test su animali presentano varie limitazioni: hanno costi elevati, richiedono tempi lunghi (tipicamente compresi tra i 6 ed i 24 mesi per un singolo composto), sollevano questioni etiche legate al benessere degli animali e non sempre predicono in modo accurato la risposta nell'essere umano, a causa delle differenze biologiche tra le specie. L'IA si sta promuovendo come strumento complementare a questi test, con l'obiettivo di ridurre il numero di animali utilizzati senza compromettere l'efficacia e la qualità della valutazione della sicurezza, ottimizzando lo sviluppo preclinico (43).

Infatti, i modelli predittivi basati su machine learning e deep learning vengono addestrati su grandi dataset di composti per i quali sono già disponibili dati tossicologici sperimentali, imparando a riconoscere le caratteristiche chimiche strutturali associate a specifici profili di tossicità. Una volta addestrati, questi modelli sono in grado di valutare nuovi composti candidati direttamente dalla loro struttura chimica, senza necessità di esperimenti biologici. Questo approccio consente di filtrare precocemente i composti a maggior rischio tossicologico, concentrando la sperimentazione animale solo sui candidati che hanno

superato una valutazione computazionale preliminare favorevole e riducendo così il numero totale di animali utilizzati senza rinunciare alla sicurezza del processo (43).

Un esempio concreto di applicazione dell'IA a supporto della riduzione dei modelli animali è rappresentato dallo sviluppo di Digital Twins di ratto basati su reti neurali generative di tipo GAN (*Generative Adversarial Networks*).

Le GAN sono modelli composti da due reti neurali che lavorano in competizione tra loro: una rete generatrice che produce dati sintetici e una rete discriminatrice che valuta se tali dati siano realistici o meno. In questo contesto, le GAN vengono addestrate su dati tossicologici reali provenienti da esperimenti che permettono di riprodurre fedelmente le risposte biologiche tipiche dell'animale all'esposizione a farmaci. Una volta addestrate, queste reti sono in grado di simulare come un ratto virtuale risponderebbe a un nuovo composto, generando previsioni sull'evoluzione nel tempo di parametri clinici indicativi di epatotossicità, come i valori degli enzimi epatici ALT e AST. Questi modelli sono inoltre in grado di generare profili trascrittomici epatici virtuali, ovvero simulazioni di come cambia l'espressione dei geni nel tessuto epatico del ratto in risposta a un trattamento: in pratica, prevedono quali geni vengono attivati o silenziati nel fegato dopo la somministrazione del composto, fornendo informazioni sul meccanismo molecolare attraverso cui il farmaco potrebbe causare tossicità (**Figura 17**). Il concetto di modello animale virtuale dà quindi la possibilità di testare un composto su una replica digitale del ratto costruita a partire da dati biologici reali, ottenendo previsioni biologicamente plausibili sulla sua tossicità. Questo approccio consente di ridurre il numero di animali necessari in linea con il principio delle 3R: *Replacement, Reduction and Refinement*, che guida l'etica della sperimentazione animale (14).



**Figura 17.** I Digital Twins sono rappresentazioni virtuali di entità biologiche. Qui sono rappresentati come Digital Twins di (a) linee cellulari, (b) organi completi e (c) interi pazienti umani. Tra l'entità biologica reale e il suo gemello digitale esiste un flusso bidirezionale di informazioni: il Digital Twin viene inizializzato mediante misurazioni ottenute dall'entità biologica reale che, a sua volta, restituisce previsioni o risultati di simulazioni di esperimenti virtuali che supportano il processo decisionale relativo all'entità biologica stessa. (d) I Digital Twins influenzano tutte le fasi della scoperta e dello sviluppo di farmaci, inclusa la conduzione degli studi clinici (14).

Un aspetto particolarmente rilevante riguarda la capacità di questi modelli di predire la tossicità oltre che epatica anche su altri organi come reni, cuore e sistema nervoso. Questo è possibile perché i modelli vengono addestrati su dataset che contengono non solo informazioni sulla struttura chimica dei composti, ma anche dati sperimentali su quali organi sono stati danneggiati e con quale intensità. In questo modo il modello impara ad associare specifiche caratteristiche chimiche della molecola, come la presenza di determinati gruppi funzionali, la lipofilia o la capacità di attraversare le membrane cellulari, con il rischio di tossicità su un determinato organo. Questi modelli sono inoltre in grado di identificare meccanismi di tossicità che i metodi tradizionali faticano a rilevare nelle fasi precoci, perché analizzano simultaneamente un numero molto elevato di caratteristiche chimiche e biologiche che un ricercatore non riuscirebbe a valutare manualmente su larga scala, riconoscendo pattern complessi che emergono solo dall'analisi di grandi quantità di dati. Questi modelli sono in grado di fornire informazioni molto più dettagliate rispetto a una semplice classificazione binaria in 'tossico' o 'non tossico': sono infatti in grado di indicare

quale organo è a rischio, a quale dose la tossicità si manifesta, attraverso quale meccanismo biologico si produce il danno e in quali categorie di pazienti, come ad esempio i pazienti con alterazioni genetiche nel metabolismo del farmaco che potrebbero essere più vulnerabili (43).

È importante sottolineare che questi modelli predittivi non sostituiscono completamente la sperimentazione biologica: i composti che superano il filtro computazionale devono comunque essere validati attraverso esperimenti in vitro e in vivo prima di procedere alle fasi successive dello sviluppo. Il vantaggio non è infatti l'eliminazione completa di questi test, ma nel ridurre drasticamente il numero di composti su cui condurli; quindi, invece di testare sperimentalmente tutti i candidati, si portano ai test biologici solo quelli che hanno già superato una valutazione computazionale favorevole riducendo così i tempi e i costi complessivi della fase preclinica. La qualità dei modelli predittivi dipende inoltre in modo critico dalla qualità e dalla quantità dei dati tossicologici su cui sono stati addestrati, un aspetto che rimane una delle principali sfide del settore (28,43).

## **4.2 Innovazione negli studi clinici**

Gli studi clinici rappresentano la fase più critica e costosa dell'intero processo di sviluppo di un farmaco. Prima che un nuovo trattamento possa essere approvato dagli enti regolatori, deve dimostrare la propria sicurezza ed efficacia attraverso un percorso che si articola generalmente in tre fasi sequenziali: I, II e III, ciascuna delle quali coinvolge un numero crescente di pazienti e richiede anni di raccolta e analisi dei dati. Nonostante i progressi scientifici degli ultimi decenni, il tasso di successo complessivo di un farmaco attraverso tutte le fasi cliniche rimane compreso tra il 5% e il 10%, con costi per singolo farmaco approvato che spesso superano il miliardo di dollari (42).

Una delle ragioni di questa bassa efficienza risiede nel fatto che il processo di sviluppo clinico, sebbene si sia evoluto sostanzialmente nei metodi e nella qualità, è rimasto invariato nelle sue caratteristiche fondanti negli ultimi trent'anni, basandosi su approcci tradizionali che valutano l'effetto di un farmaco su grandi popolazioni di pazienti considerati come un gruppo uniforme, senza considerare le differenze interindividuali all'interno di una popolazione anche selezionata. Infatti, i trial arruolano popolazioni eterogenee senza

distinguere i sottogruppi molecolari, misurano un effetto medio sull'intera popolazione che può mascherare benefici importanti in sottogruppi specifici, se questi sottogruppi non sono previsti a priori. Personalizzare il processo di reclutamento e riconoscere sottogruppi di risposta/tossicità è un passo fondamentale del prossimo futuro (46).

È in questo contesto che l'IA si sta affermando come strumento capace di trasformare come gli studi clinici vengono progettati, condotti ed infine analizzati. Un segnale concreto viene dall'analisi condotta sulle pipeline cliniche delle aziende biotecnologiche specializzate in IA che mostra come i farmaci scoperti con IA stiano crescendo a un ritmo superiore al 60% di anno in anno e, in fase I, questi farmaci mostrano un tasso di successo dell'80-90%, significativamente superiore alla media storica del settore compresa tra il 40% e il 65%, suggerendo che l'IA sia particolarmente efficace nell'ottimizzare le proprietà ADMET e la sicurezza dei candidati. In fase II, tuttavia, il tasso di successo si attesta intorno al 40%, in linea con le medie storiche, indicando che il problema dell'efficacia clinica determinato dalla complessità biologica delle malattie rimane una sfida che l'IA non ha ancora risolto (42).

Le applicazioni dell'IA agli studi clinici sono molteplici e possono essere ricondotte a tre grandi aree, che verranno approfondite nei paragrafi seguenti.

La prima riguarda l'ottimizzazione del disegno dello studio e la selezione dei pazienti. Analizzando grandi quantità di dati clinici, l'IA può individuare più rapidamente i partecipanti idonei ad un trial e suddividerli in sottogruppi omogenei in base alle loro caratteristiche cliniche, molecolari e genetiche, aprendo la strada ad una sperimentazione più mirata, capace di tener conto delle differenze individuali nella risposta ai farmaci fino alla farmacogenomica, come viene affrontato nel paragrafo 4.2.1. La seconda riguarda l'identificazione di nuovi biomarcatori predittivi di risposta terapeutica, cioè degli indicatori biologici che predicono la risposta ad un trattamento: oltre a quelli molecolari tradizionali, l'IA permette di ricavarne di nuovi, i cosiddetti biomarcatori digitali, analizzando immagini, segnali fisiologici e dati prodotti da dispositivi, come discusso nel paragrafo 4.2.2. La terza, approfondita nel paragrafo 4.2.3, riguarda il monitoraggio remoto e continuo dei partecipanti: grazie ai dispositivi indossabili e all'analisi dei dati in tempo reale, definiti come Real World Data (RWD), l'IA consente di poter seguire i partecipanti in modo costante e a distanza, anche al di fuori delle visite programmate, migliorando la qualità dei dati raccolti e l'accessibilità agli studi (47).

#### **4.2.1 Ottimizzazione del disegno dello studio, stratificazione dei pazienti, medicina personalizzata e farmacogenomica**

Uno degli ambiti in cui l'IA sta dimostrando il maggiore potenziale negli studi clinici riguarda l'ottimizzazione del disegno dello studio e la selezione dei pazienti. Il reclutamento dei partecipanti rappresenta storicamente uno dei principali limiti del processo clinico: si stima che circa l'80% di tutti i trial clinici subisca ritardi significativi proprio a causa delle difficoltà nel raggiungere il numero di pazienti necessario nei tempi previsti. Questo problema è particolarmente evidente negli studi per malattie rare, nelle patologie con criteri di eleggibilità complessi e nei trial che richiedono popolazioni di pazienti con caratteristiche molecolari specifiche (14,38).

I modelli di IA possono affrontare questo problema analizzando in modo automatico le cartelle cliniche elettroniche di centinaia di migliaia di pazienti, identificando rapidamente quelli che soddisfano i criteri di inclusione ed esclusione di un determinato studio. Questo è possibile grazie al *natural language processing*, una branca dell'IA che consente ai computer di comprendere e analizzare il linguaggio scritto e che permette di estrarre informazioni rilevanti anche da testo libero, come le note cliniche dei medici, i referti radiologici e le lettere di dimissione, che tradizionalmente non sono analizzabili in modo automatico. Strumenti come Deep 6 AI e Mendel.AI sono in grado di analizzare milioni di cartelle cliniche identificando pazienti potenzialmente eleggibili che una persona umana avrebbe difficoltà a individuare manualmente su larga scala, promuovendo al tempo stesso una maggiore diversità e rappresentatività dei campioni di studio (38).

Un esempio particolarmente rilevante di applicazione dell'IA nell'ottimizzazione del reclutamento di partecipanti ai trial è rappresentato dal sistema RECTIFIER, sviluppato nell'ambito del trial clinico COPILOT-HF per pazienti con insufficienza cardiaca sintomatica. RECTIFIER utilizza GPT-4 che riesce a consultare in tempo reale le cartelle cliniche dei pazienti per valutare criteri di inclusione e criteri di esclusione. Il sistema ha dimostrato una concordanza con la valutazione di clinici esperti nel 98-100% dei casi, con un costo di soli 0.10 dollari per paziente valutato, rendendo questo approccio non solo accurato ma anche economicamente sostenibile su larga scala (47).

Un'altra innovazione significativa riguarda l'utilizzo dei dati raccolti da cartelle cliniche elettroniche, registri di malattia e database assicurativi per simulare i risultati che si otterrebbero applicando diversi criteri di inclusione senza dover condurre effettivamente il trial. In oncologia è stato sviluppato lo strumento Trial per emulare i risultati di trial clinici utilizzando dati reali di cartelle cliniche e tecniche statistiche che correggono le differenze tra i pazienti presenti nei dati della pratica clinica, rendendo i gruppi confrontabili come se fossero stati assegnati in modo casuale, come avviene in un trial randomizzato. La sua applicazione con diversi set di criteri di inclusione ha suggerito che alcuni trial oncologici potrebbero ottenere la stessa efficacia del farmaco anche allargando i criteri di inclusione, per esempio ammettendo pazienti più anziani o con comorbidità che normalmente verrebbero esclusi, consentendo così di arruolare più pazienti rapidamente e di ottenere risultati più applicabili alla popolazione reale. Analogamente, i RWD possono essere utilizzati per costruire bracci di controllo virtuali, simulando i pazienti che avrebbero ricevuto il placebo a partire da dati reali già disponibili, un approccio già accettato dalla FDA come strumento per supportare le decisioni nella progettazione degli studi. Questo consente in alcuni scenari di ridurre il numero di pazienti reali assegnati al placebo, il che ha implicazioni etiche rilevanti: meno pazienti ricevono un trattamento inattivo quando potrebbero invece beneficiare del farmaco sperimentale (46,47).

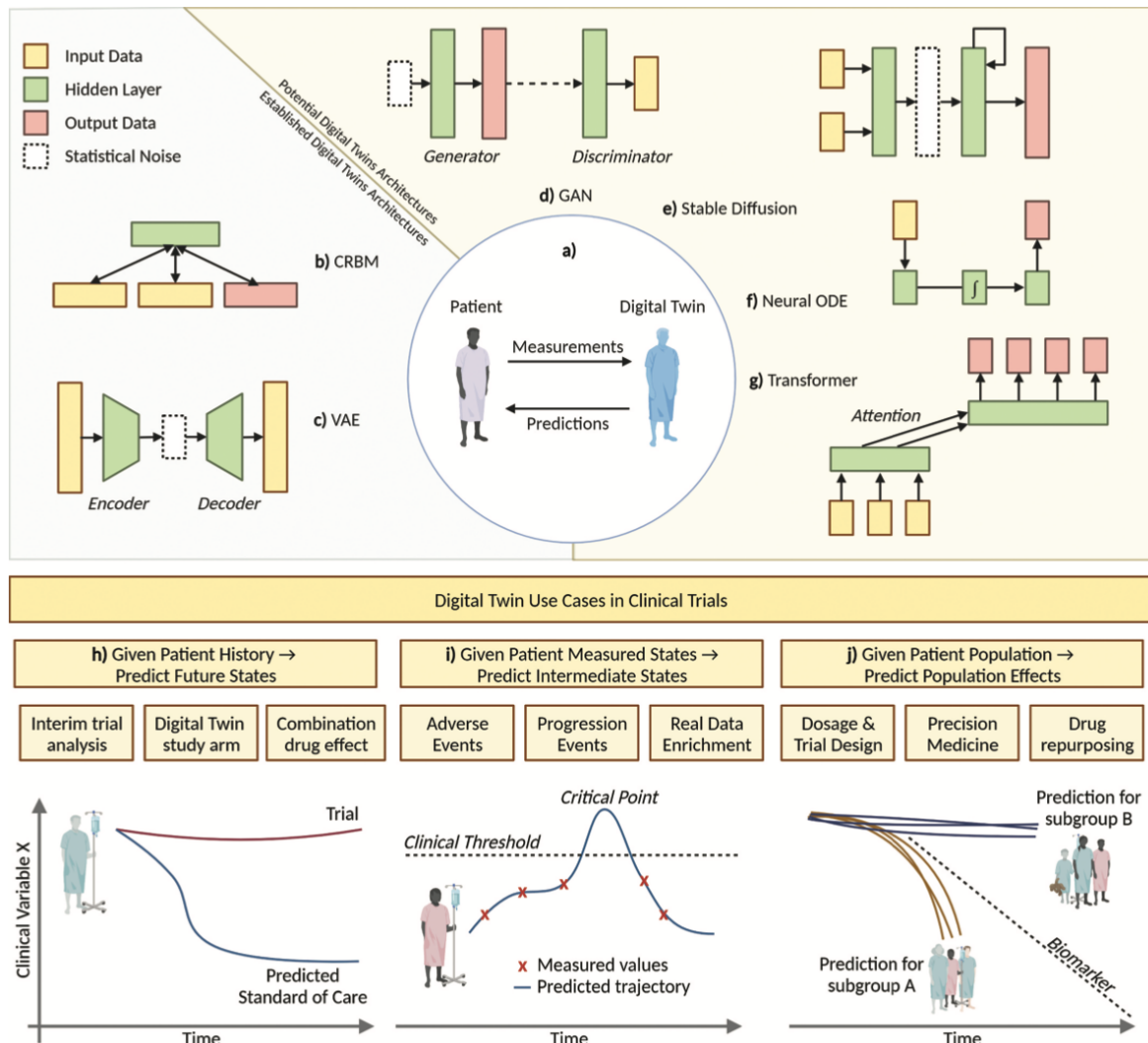
Un altro contributo fondamentale dell'IA riguarda la stratificazione dei pazienti, cioè la capacità di suddividere la popolazione arruolata in sottogruppi omogenei sulla base di caratteristiche molecolari, genetiche o cliniche, con l'obiettivo di identificare quali pazienti beneficerebbero maggiormente di un trattamento. L'approccio tradizionale consiste nel decidere in anticipo, prima di avviare il trial, di analizzare solo alcuni gruppi di pazienti specifici scelti sulla base di ipotesi biologiche già note. Il problema è che il numero di pazienti arruolati in un trial di solito non è adeguato da permettere di rilevare differenze statisticamente affidabili tra questi sottogruppi: quando si divide la popolazione in gruppi più piccoli si hanno meno dati e i risultati diventano meno robusti. I modelli di IA possono invece analizzare contemporaneamente tutte le caratteristiche dei pazienti misurate all'inizio del trial, come età, sesso, esami di laboratorio, diagnosi precedenti e profilo genetico costruendo una rappresentazione multidimensionale dell'intera popolazione, identificando gruppi di pazienti con caratteristiche simili e prevedendo come risponderà ciascun paziente al trattamento sulla base di come hanno risposto pazienti con caratteristiche simili. Questo approccio ha dimostrato accuratezza quando testato su gruppi di pazienti completamente

diversi da quelli su cui il modello è stato costruito, provenienti da altri ospedali o altri studi, a conferma che questi strumenti funzionano anche in contesti nuovi. Alcune simulazioni suggeriscono inoltre che modificare i criteri di partecipazione al trial durante il suo svolgimento sulla base di questi modelli, e quindi includendo progressivamente i pazienti con maggiore probabilità di rispondere, potrebbe ridurre il numero totale di pazienti necessari del 15-20%, con un impatto significativo sui tempi e sui costi degli studi (47).

Un altro aspetto particolarmente rilevante riguarda la farmacogenomica, ovvero lo studio di come le variazioni genetiche individuali influenzino la risposta ai farmaci. I polimorfismi del CYP450 assumono un'importanza ancora maggiore in fase clinica, dove le differenze nel metabolismo individuale possono determinare efficacia o tossicità di un farmaco per uno specifico paziente. Un esempio classico è il CYP2D6, un enzima epatico che metabolizza circa il 25% dei farmaci comunemente prescritti, inclusi molti antidepressivi, antipsicotici e analgesici. I pazienti possono essere classificati come metabolizzatori lenti, cioè con attività enzimatica molto ridotta, metabolizzatori intermedi o rapidi/ultrarapidi, ciascuno dei quali può rispondere in modo completamente diverso alla stessa dose di un farmaco. Ad esempio, un metabolizzatore ultrarapido elimina il farmaco così velocemente da non raggiungere mai concentrazioni terapeutiche efficaci, mentre un metabolizzatore lento può accumulare il farmaco fino a livelli tossici. L'IA, integrando dati genetici con dati clinici e molecolari, consente di stratificare i pazienti nei trial in base al loro profilo farmacogenomico, selezionando le popolazioni più appropriate per un determinato trattamento e riducendo il rischio di eventi avversi legati a differenze nel metabolismo (15,26).

Un'ulteriore innovazione nel disegno degli studi clinici è rappresentata dai *digital twins*, ovvero repliche virtuali di pazienti reali costruite a partire dai loro dati clinici, biologici e molecolari (**Figura 18**). Il concetto, originariamente sviluppato in ingegneria per testare virtualmente sistemi complessi prima della loro realizzazione fisica, è stato adattato al contesto clinico con l'obiettivo di simulare come un paziente specifico potrebbe rispondere a un trattamento senza dover effettivamente somministrarlo. I *digital twins* vengono inizializzati con le caratteristiche reali del paziente a un determinato momento temporale, come i valori di laboratorio, i parametri vitali, eventuali diagnosi precedenti e il profilo genetico, e successivamente simulano l'evoluzione del suo stato clinico partendo da quel

punto, generando traiettorie cliniche virtuali che possono essere confrontate con l'evoluzione osservata nel gruppo di trattamento reale.



**Figura 18.** Architetture dei gemelli digitali (Digital Twins, DT) e relativi casi d'uso negli studi clinici. **(a)** Il flusso di informazioni tra i DT e i pazienti reali è bidirezionale, poiché i DT costituiscono rappresentazioni virtuali dei pazienti e possono contribuire al miglioramento del trattamento clinico. **(h)** Casi d'uso dei DT basati sulla previsione degli stati futuri del paziente a partire dalla sua storia clinica. **(i)** Casi d'uso supportati dalle previsioni dei DT degli stati intermedi del paziente sulla base degli stati misurati. **(j)** Casi d'uso dei DT fondati sulla previsione degli effetti a livello di popolazione per una determinata popolazione e i relativi parametri, ad esempio per definire i criteri di inclusione in uno studio clinico (14).

Ad esempio, la società Unlearn.AI ha sviluppato e testato modelli di digital twins per trial clinici in Alzheimer e sclerosi multipla, sfruttando dati storici provenienti da bracci di controllo placebo di trial precedenti per costruire modelli generativi in grado di simulare le traiettorie cliniche individuali. I *digital twins* consentono di affiancare a ciascun paziente arruolato una traiettoria clinica virtuale, cioè una previsione di come si sarebbe evoluto in assenza del trattamento, generata dal modello a partire dai dati storici. Integrando queste

previsioni nell'analisi è possibile aumentare la potenza statistica dello studio e ridurre così la dimensione del braccio di controllo a parità di affidabilità dei risultati. La riduzione del numero di pazienti esposti al placebo è un obiettivo perseguito anche nei trial tradizionali, i *digital twins* offrono però un modo per ottenerla mantenendo la randomizzazione e sfruttando i dati già disponibili, con evidenti vantaggi etici. Si tratta tuttavia di un approccio ancora in fase di sviluppo e validazione: la verifica delle traiettorie generate dai digital twins resta limitata e manca ad oggi una regolamentazione specifica per il loro impiego negli studi clinici (8).

#### **4.2.2 Identificazione di biomarcatori digitali e molecolari**

Un biomarcatore è una caratteristica biologica misurabile, come il livello di una proteina nel sangue, la presenza di una mutazione genetica o un parametro fisiologico, che fornisce informazioni sullo stato di salute di un paziente, sulla progressione di una malattia o sulla risposta a un trattamento. Tradizionalmente, i biomarcatori utilizzati negli studi clinici sono di natura molecolare o genetica, e la loro identificazione richiede test di laboratorio spesso costosi, invasivi e non sempre accessibili in tutti i contesti clinici. L'IA sta aprendo una nuova categoria di biomarcatori, i cosiddetti biomarcatori digitali, ovvero indicatori predittivi ricavati dall'analisi computazionale di immagini, segnali fisiologici o dati prodotti da dispositivi indossabili, senza necessità di prelievi o test molecolari diretti (15,47).

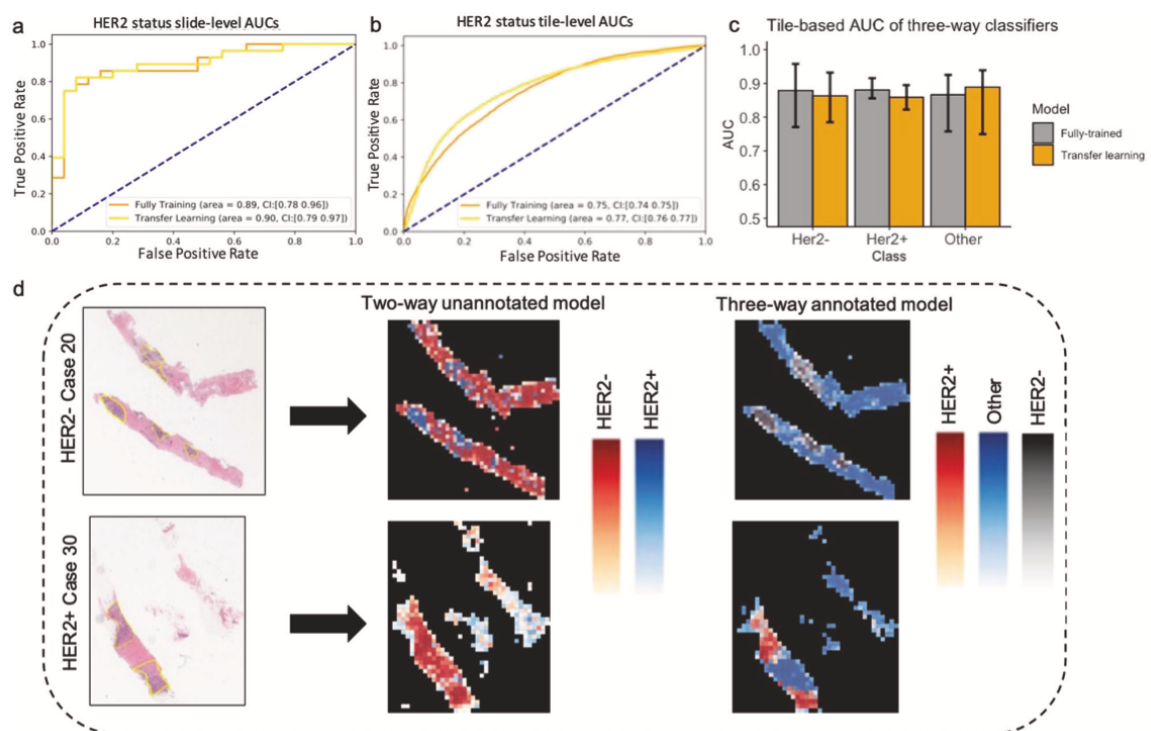
Uno degli ambiti più promettenti riguarda l'applicazione del deep learning all'analisi delle comuni immagini istologiche prodotte durante la valutazione patologica di un tumore con l'obiettivo di ricavare informazioni molecolari complesse senza dover eseguire test genetici aggiuntivi. Queste immagini, che il patologo osserva per valutare la morfologia del tumore, possono contenere più informazioni di quanto si può cogliere: variazioni sottili nella forma, nelle dimensioni e nell'organizzazione delle cellule possono essere riconosciuti da modelli di deep learning che vengono addestrati su migliaia di immagini di pazienti per i quali si conosce già il profilo molecolare. Questi modelli, perciò, vengono addestrati su caratteristiche visive come la forma dei nuclei, la densità cellulare o la presenza di cellule immunitarie nel tessuto tumorale, e applicano questa conoscenza per fare previsioni su nuovi pazienti. In questo modo, l'immagine istologica diventa essa stessa un biomarcatore, capace di fornire informazioni clinicamente rilevanti senza bisogno di ulteriori esami.

Un esempio particolarmente rilevante riguarda il carcinoma gastrico. Uno studio multicentrico retrospettivo condotto raccogliendo dati provenienti da dieci ospedali in sette paesi diversi, su più di 2800 pazienti con carcinoma gastrico ha sviluppato classificatori basati su deep learning per rilevare lo stato di instabilità dei microsatelliti, noto con l'acronimo MSI, e la positività al virus di Epstein-Barr direttamente da immagini istologiche di routine. L'instabilità dei microsatelliti è una condizione in cui si verificano errori nel meccanismo di riparazione del DNA cellulare, portando a un accumulo di mutazioni. Questo stato molecolare è clinicamente rilevante perché i tumori MSI-positivi rispondono meglio alle immunoterapie con inibitori dei checkpoint immunitari, cioè farmaci che riattivano il sistema immunitario del paziente contro le cellule tumorali, rispetto ai tumori MSI-negativi in cui il meccanismo di correzione funziona normalmente, rendendo la sua identificazione fondamentale per selezionare i pazienti che beneficeranno maggiormente di questo tipo di trattamento. Quando testato su pazienti provenienti da ospedali diversi da quelli utilizzati per addestrarlo, quindi su coorti esterne, il modello ha raggiunto un'accuratezza (misurata tramite curva ROC, cioè un indice che va da 0 a 1, dove quest'ultimo rappresenta una predizione completa) fino a 0.86, dimostrando una forte predizione. Un risultato particolarmente significativo è stato che la performance del modello migliorava sostanzialmente quando veniva addestrato su dati provenienti da più centri clinici di paesi diversi, confermando l'importanza di disporre di dati rappresentativi di popolazioni differenti. Questo approccio potrebbe consentire un pre-screening rapido ed economico dei pazienti tramite l'analisi dell'immagine istologica già disponibile che permette di identificare rapidamente i pazienti con alta probabilità di essere positivi, riducendo il numero di test molecolari necessari e rendendo accessibile la stratificazione molecolare anche in contesti con risorse sanitarie limitate (48).

Risultati analoghi sono stati ottenuti nel carcinoma coloretale, dove un sistema di deep learning è stato sviluppato per predire simultaneamente lo stato di sette caratteristiche molecolari rilevanti dal punto di vista terapeutico, tra cui MSI, ipermutazione, instabilità cromosomica, metilazione del DNA e mutazioni di geni chiave come BRAF, TP53 e KRAS direttamente da immagini istologiche di routine. Le mutazioni di BRAF e KRAS sono particolarmente rilevanti in farmacologia clinica perché determinano la risposta o la resistenza alle terapie con anticorpi diretti contro il recettore EGFR, come Cetuximab e panitumumab: i pazienti con mutazione di KRAS non beneficiano di questi trattamenti e la loro identificazione precoce consente di evitare terapie costose e inefficaci. Il modello ha

raggiunto un'accuratezza di 0.86 per la predizione dell'MSI e nella validazione su una coorte esterna indipendente ha raggiunto un'accuratezza di 0.98. Un'analisi delle caratteristiche cellulari delle regioni istopatologiche più predittive ha evidenziato che un'elevata presenza di linfociti nel tessuto tumorale era fortemente associata allo stato MSI, emergendo come biomcatore digitale di natura morfologica rilevabile automaticamente dall'IA nell'immagine istologica (23).

Nel carcinoma mammario, un approccio analogo ha dimostrato la capacità di predire lo stato di amplificazione del gene HER2, ovvero una proteina che quando prodotta in eccesso stimola la crescita incontrollata delle cellule tumorali e che guida la scelta di terapie mirate con anticorpi monoclonali, direttamente da immagini istologiche di routine, con un'accuratezza di 0.90 nella validazione interna e di 0.81 su un gruppo di pazienti provenienti da un'altra struttura (**Figura 19**).



**Figura 19.** Predizione dello stato HER2 da immagini istologiche (H&E) mediante rete neurale convoluzionale. I pannelli (a) e (b) mostrano le curve ROC per la classificazione HER2 (49).

Ancora più significativo è il fatto che lo stesso modello è stato in grado di predire la risposta all'anticorpo monoclonale con un'accuratezza di 0.80, superando la classificazione basata sul solo stato HER2 determinato con i metodi tradizionali e dimostrando così che l'immagine istologica contiene informazioni sul comportamento biologico del tumore, come

l'organizzazione del tessuto, la forma delle cellule e la presenza di cellule immunitarie, che vanno oltre quello che il test molecolare tradizionale è in grado di rilevare, e che l'IA riesce a sfruttare per fare previsioni più accurate sulla risposta al farmaco (49).

Parallelamente ai biomarcatori digitali derivati da immagini, l'IA sta potenziando l'identificazione di biomarcatori molecolari attraverso l'integrazione di dati omici, cioè dati provenienti dall'analisi del genoma, del proteoma, del trascrittoma e del metaboloma di un paziente, che descrivono rispettivamente il patrimonio genetico, le proteine prodotte, i geni attivi e le molecole metaboliche presenti in un determinato momento. L'integrazione di questi diversi livelli di informazione biologica, combinata con dati di imaging clinico e cartelle cliniche elettroniche, consente di costruire profili molecolari del paziente molto più completi di quanto possibile con un singolo tipo di analisi.

In oncologia, un esempio concreto di questa integrazione è rappresentato dalla biopsia liquida, cioè l'analisi di frammenti di DNA tumorale circolante che il tumore rilascia nel sangue del paziente, ottenibile con un semplice prelievo venoso e senza biopsie tissutali invasive. Questa tecnica produce grandi quantità di dati molecolari, ma presenta una difficoltà di fondo: il DNA tumorale costituisce solo una piccola frazione del DNA libero presente nel sangue, gran parte del quale proviene da cellule sane.

In questa fase possono intervenire i modelli di IA, utili per distinguere il segnale tumorale dal rumore di fondo per estrarre dai dati le informazioni clinicamente rilevanti, come la caratterizzazione molecolare del tumore e la sua evoluzione nel tempo, integrandole con i dati clinici del paziente per guidare la stratificazione e le decisioni terapeutiche (50).

Un'ulteriore applicazione dell'IA riguarda i punteggi di rischio poligenico, cioè i modelli computazionali analizzano migliaia di varianti genetiche distribuite nel genoma di un individuo per stimare la sua predisposizione a sviluppare una determinata malattia conferendogli un punteggio di rischio. Nel contesto degli studi clinici, questi strumenti consentono di selezionare preferenzialmente i pazienti geneticamente più a rischio di sviluppare l'evento che il farmaco intende prevenire, come ad esempio un infarto o una recidiva tumorale, rendendo lo studio più efficiente e riducendo il numero di pazienti necessari per dimostrare l'efficacia del trattamento (26).

### 4.2.3 Monitoraggio remoto ed analisi dei dati in tempo reale

Uno dei limiti più significativi dei trial clinici tradizionali riguarda la modalità di raccolta dei dati sui partecipanti, che avvengono solitamente durante le visite programmate e solo raramente sono previsti dispositivi di raccolta dati in tempo reale e continuativi. L'IA, combinata con l'uso sempre più diffuso ed accessibile dei dispositivi digitali indossabili e dei sensori ambientali, sta rendendo possibile un monitoraggio continuo e remoto dei partecipanti agli studi clinici, con potenziali implicazioni sia per la qualità dei dati raccolti sia per l'accessibilità degli studi stessi (26,47).

I dispositivi indossabili come gli smartwatch, le fasce cardiache e i sensori di attività fisica sono oggi in grado di raccogliere in modo continuo una vasta gamma di parametri fisiologici, tra cui frequenza cardiaca, saturazione dell'ossigeno nel sangue, variabilità della frequenza cardiaca, temperatura cutanea, qualità del sonno e livelli di attività fisica. I modelli di IA sono in grado di analizzare questi flussi continui di dati, identificando variazioni anche minime nei parametri fisiologici consentendo interventi tempestivi se necessari (26,47).

Un esempio concreto di biomarcatore digitale derivato da dispositivi indossabili è rappresentato dallo strumento *HearO*, sviluppato dall'azienda Cordio, che analizza registrazioni vocali dei pazienti con insufficienza cardiaca per identificare segni precoci di congestione polmonare, ovvero l'accumulo di liquido nei polmoni che è uno dei principali indicatori di peggioramento della malattia. Il sistema analizza variazioni sottili nella qualità della voce che precedono la comparsa dei sintomi clinici evidenti, consentendo di identificare i pazienti a rischio di scompenso prima che la situazione diventi critica. Analogamente, i sistemi di posizionamento GPS integrati negli smartphone possono aiutare i ricercatori a identificare automaticamente quando un paziente si reca in ospedale o al pronto soccorso in modo non programmato, un evento che potrebbe indicare un peggioramento della malattia o un effetto avverso del farmaco. Il sistema segnala l'accesso in tempo reale, consentendo ai ricercatori di richiedere tempestivamente la cartella clinica di quella visita per ottenere i dettagli clinici rilevanti per lo studio (47).

La decentralizzazione dei trial clinici, ovvero lo spostamento del monitoraggio dei pazienti dall'ospedale al domicilio, rappresenta una delle trasformazioni più significative che l'IA sta implementando. Questo approccio riduce il carico logistico e organizzativo per i partecipanti, che possono recarsi con meno frequenza presso i centri clinici, e amplia

geograficamente il bacino di reclutamento. Un minore carico per i partecipanti può inoltre tradursi in una maggiore diversità demografica degli studi, includendo gruppi di pazienti tradizionalmente sottorappresentati come anziani, persone con difficoltà di mobilità o residenti in aree rurali (38,47).

Oltre al monitoraggio dei parametri fisiologici, l'IA interviene anche nella valutazione automatica degli esiti del trial, che può essere distinta in due ambiti.

Un primo ambito riguarda l'identificazione e l'aggiudicazione degli endpoint clinici. In un trial, gli endpoint sono esiti prestabiliti dal protocollo, utilizzati per misurare l'efficacia o la sicurezza di un intervento. L'aggiudicazione consiste nel determinare se uno specifico evento clinico, ad esempio un ricovero, un decesso o una complicanza cardiovascolare, si sia effettivamente verificato e se soddisfatti i criteri definiti dal protocollo. Da questa valutazione dipende la classificazione finale degli eventi e, di conseguenza, una parte rilevante dei risultati di efficacia e sicurezza dello studio. (47)

Tradizionalmente, questa attività è affidata a commissioni indipendenti di esperti, spesso indicate come clinical events committees, che esaminano manualmente cartelle cliniche, referti e documentazione di follow-up. Si tratta di un processo accurato, ma costoso, lento e difficilmente scalabile nei trial di grandi dimensioni. I modelli di *natural language processing* possono supportare questa attività analizzando automaticamente documenti clinici non strutturati e proponendo una prima classificazione degli eventi. Nel trial INVESTED, uno studio multicentrico condotto negli Stati Uniti e in Canada sul vaccino antinfluenzale in pazienti con malattie cardiovascolari, un modello NLP è stato utilizzato per identificare i ricoveri per insufficienza cardiaca a partire dalla documentazione clinica. Il modello ha mostrato una buona concordanza con il clinical events committee e, dopo ulteriore messa a punto, ha raggiunto prestazioni prossime alla riproducibilità osservata tra revisori umani (47).

Più recentemente, i *large language models* hanno ampliato questa possibilità anche ai dati conversazionali. In una *secondary analysis* del trial randomizzato China CT-FFR Study 3, le interviste telefoniche di follow-up sono state trasformate in testo e analizzate da un modello linguistico dedicato, sviluppato per classificare automaticamente eventi clinici riferiti dai partecipanti, tra cui decessi, ricoveri, procedure invasive e uso di farmaci. Il modello ha raggiunto un accordo complessivo del 93.7% rispetto allo standard di riferimento, con sensibilità e specificità elevate, superando sia modelli linguistici generalisti sia adjudicator umani meno esperti (19). Tuttavia, gli autori sottolineano che questo

approccio deve essere interpretato come uno strumento di preaggiudicazione: alcune categorie di eventi, in particolare i ricoveri descritti in modo ambiguo, rimangono difficili da classificare e gli esiti critici richiedono comunque supervisione umana. (51)

Un secondo ambito, distinto ma complementare, riguarda il riconoscimento e la classificazione degli eventi avversi. In questo caso l'obiettivo non è stabilire se un endpoint pre-specificato sia stato raggiunto, ma individuare possibili segnali di sicurezza, inclusi eventi indesiderati descritti in testo libero nelle note cliniche, nei diari dei pazienti o nella documentazione di follow-up. La distinzione è importante: l'aggiudicazione degli endpoint applica criteri prestabiliti a eventi attesi dal protocollo, mentre la farmacovigilanza mira anche a intercettare eventi inattesi, pattern ricorrenti o segnali di rischio che possono emergere durante lo studio. In questo contesto, modelli NLP e LLM possono facilitare l'estrazione, la codifica e la classificazione automatica degli eventi avversi dalle narrazioni cliniche. (47) Tuttavia, le evidenze disponibili derivano in larga parte da analisi retrospettive o da dataset annotati, e l'impiego prospettico di questi strumenti nei trial clinici richiede ancora validazione, integrazione nei flussi regolatori e chiara definizione della supervisione umana (52).

Infine, l'IA può contribuire al monitoraggio continuo della qualità e della sicurezza dei dati durante lo svolgimento del trial. Sistemi di analisi automatizzata possono identificare più rapidamente dati mancanti, incoerenze, eventi potenzialmente rilevanti o segnali di sicurezza che richiedono revisione. Questo non sostituisce le procedure formali di interim analysis, né il ruolo dei comitati indipendenti di monitoraggio, ma può fornire un livello aggiuntivo di sorveglianza operativa. Se integrati in processi pre-specificati e sottoposti a governance clinica e statistica, questi strumenti potrebbero rendere più tempestiva l'identificazione di segnali di efficacia, futilità o danno, migliorando l'efficienza del trial e la tutela dei partecipanti (47).

### **4.3 Farmacovigilanza e sicurezza post-marketing: utilizzo dell'IA per l'analisi di *Real-World Data* e la rapida identificazione di eventi avversi**

La farmacovigilanza è definita dall'Organizzazione Mondiale della Sanità (OMS) come la scienza e l'insieme delle attività relative alla detezione, alla valutazione, alla comprensione e alla prevenzione degli effetti avversi dei farmaci o di qualsiasi altro problema correlato al loro utilizzo. Il suo obiettivo principale è garantire la sicurezza dei farmaci dopo la loro immissione in commercio, monitorando continuamente le reazioni avverse che emergono nella pratica clinica reale. Le reazioni avverse ai farmaci, note anche con l'acronimo ADR, dall'inglese *Adverse Drug Reactions*, rappresentano un problema di salute pubblica significativo: si stima che costituiscano tra il 5% e il 7% di tutte le visite al pronto soccorso e che causino un prolungamento della durata media del ricovero ospedaliero sia per i pazienti ambulatoriali che per quelli ricoverati (52).

Quando un farmaco viene approvato e reso disponibile ai pazienti, la sperimentazione clinica che ne ha dimostrato l'efficacia e la sicurezza è stata condotta su un numero relativamente limitato di pazienti, selezionati secondo criteri precisi e seguiti per un periodo di tempo definito. Nella pratica clinica, invece, il farmaco viene somministrato a popolazioni molto più ampie e diversificate, come pazienti anziani, pazienti con multiple patologie concomitanti, pazienti in terapia con altri farmaci, e per periodi di tempo molto più lunghi. In questo contesto possono emergere reazioni avverse rare, interazioni farmacologiche non anticipate o problemi di sicurezza che non erano stati rilevati durante la sperimentazione clinica.

#### **4.3.1 L'analisi dei Real-World Data**

Gli RWD sono dati raccolti al di fuori dell'ambiente controllato degli studi clinici. Questi dati rappresentano la principale fonte di informazioni per la farmacovigilanza post-marketing, in quanto riflettono la pratica clinica. Sia la FDA che l'EMA hanno istituito infrastrutture specifiche per sfruttare i RWD per la sicurezza dei farmaci, che sono rispettivamente la *Sentinel Initiative* e la *Data Analysis and Real World Interrogation Network*, nota come DARWIN, riconoscendo ufficialmente il loro valore nella sorveglianza post-marketing (52).

Tuttavia, l'utilizzo dei RWD per la farmacovigilanza presenta sfide significative: questi dati provengono da fonti eterogenee, contengono spesso una proporzione elevata di valori

mancanti ed errori di registrazione, e la loro raccolta è complicata da problemi legali ed etici legati alla privacy dei pazienti che ne limitano l'accesso.

Nonostante queste difficoltà, l'IA sta offrendo strumenti sempre più potenti per analizzare questi enormi volumi di dati in modo sistematico. I modelli di machine learning sono infatti in grado di identificare associazioni complesse tra farmaci ed eventi avversi che i metodi statistici tradizionali faticano a rilevare, soprattutto quando le relazioni tra le variabili non sono lineari e quando i dati provengono da popolazioni molto diverse tra loro.

Una revisione sistematica della letteratura scientifica pubblicata tra il 2010 e il 2024 ha identificato 36 studi che applicano l'IA a RWD strutturati per scopi di farmacovigilanza, con un aumento significativo delle pubblicazioni dopo il 2019, confermando la crescente attenzione verso questo campo. La fonte di dati più utilizzata in questi studi è la cartella clinica elettronica, impiegata nel 78% dei casi (52).

#### **4.3.2 Identificazione degli eventi avversi e sistemi di segnalazione automatizzata**

L'identificazione delle reazioni avverse rappresenta il cuore della farmacovigilanza post-marketing. Il sistema tradizionale si basa principalmente sulla segnalazione spontanea da parte di medici, farmacisti e pazienti a database dedicati come il FAERS, il sistema americano di segnalazione degli eventi avversi gestito dalla FDA, che raccoglie milioni di segnalazioni provenienti da tutto il mondo. Tuttavia, questo sistema presenta un limite ben documentato ossia la sotto segnalazione. Infatti, solo una piccola parte delle reazioni avverse che si verificano nella pratica clinica viene effettivamente segnalata. Inoltre, le segnalazioni spontanee spesso mancano di informazioni importanti sulla storia clinica del paziente che potrebbero fare la differenza nell'interpretazione di una potenziale reazione avversa e la sua correlazione con il farmaco (52).

I modelli di IA possono affrontare questi limiti. Il 64% degli studi identificati nella revisione sistematica condotta da Dimitsaki et al. si concentra sulla identificazione automatica delle reazioni avverse, utilizzando algoritmi di machine learning addestrati su cartelle cliniche elettroniche per identificare pattern associati a specifiche reazioni avverse senza dover attendere una segnalazione esplicita da parte di un medico o un farmacista. È importante sottolineare che questi modelli non forniscono certezze ma probabilità: la valutazione finale di un potenziale segnale di sicurezza rimane sempre a carico del professionista sanitario. L'algoritmo più comunemente utilizzato in questo contesto è il *random forest*, un modello

che costruisce molti alberi decisionali e combina i loro risultati per ottenere una previsione più affidabile, utilizzato nel 47% degli studi analizzati. Questi modelli sono in grado di analizzare simultaneamente un gran numero di variabili cliniche, come diagnosi, farmaci prescritti, valori di laboratorio e parametri vitali, identificando combinazioni di fattori associati ad un aumentato rischio di reazione avversa che sarebbero difficili da rilevare manualmente (52).

Un aspetto particolarmente innovativo riguarda l'applicazione dell'IA causale, ovvero approcci che cercano non solo di identificare associazioni statistiche tra farmaci ed eventi avversi, ma di stabilire relazioni di causa-effetto tra variabili, distinguendo se una reazione avversa sia stata davvero causata dal farmaco o sia invece attribuibile ad altre caratteristiche del paziente, come l'età, le patologie preesistenti o l'assunzione di altri farmaci contemporaneamente. Questo approccio è particolarmente rilevante in farmacovigilanza perché i RWD sono per natura osservazionali e non sperimentali, rendendo difficile distinguere se un evento avverso sia stato davvero causato dal farmaco o sia invece attribuibile ad altre caratteristiche del paziente. Nonostante il suo potenziale, la revisione ha evidenziato che solo l'8% degli studi ha applicato approcci di causal machine learning, indicando che si tratta di un campo ancora in fase iniziale di sviluppo (52).

Un limite rilevante dei modelli attualmente sviluppati è la scarsa trasparenza delle loro decisioni, il cosiddetto problema della *black box*. Solo l'11% degli studi ha utilizzato tecniche di IA spiegabile, ossia metodi che permettono di comprendere quali elementi abbiano contribuito a una determinata previsione. Questo aspetto rappresenta una criticità rilevante, poiché i professionisti sanitari sottolineano la necessità di poter interpretare le motivazioni alla base delle previsioni generate dagli algoritmi. Un ulteriore limite emerso è che solo il 14% dei modelli sviluppati è stato applicato in contesti clinici reali: la maggior parte degli studi ha sviluppato e testato i modelli su dati storici già disponibili, senza mai verificare se funzionassero effettivamente quando integrati nei sistemi di gestione di un ospedale o di una farmacia, suggerendo che molti di questi strumenti non sono ancora pronti per un utilizzo diffuso nella pratica clinica (52).

#### **4.3.3 Implicazioni per il farmacista clinico**

Il farmacista clinico rappresenta una figura professionale centrale nel sistema di farmacovigilanza, in quanto è spesso il primo punto di contatto tra il paziente e il sistema

sanitario per questioni legate alla sicurezza dei farmaci. L'integrazione dell'IA nella farmacovigilanza apre nuove prospettive per il ruolo del farmacista clinico, sia in termini di strumenti disponibili sia in termini di responsabilità professionali. In Italia, il farmacista non è autorizzato a modificare autonomamente una prescrizione medica: il suo ruolo nella farmacovigilanza è principalmente quello di segnalare le reazioni avverse sospette attraverso la rete nazionale di farmacovigilanza, di informare il paziente sull'uso corretto dei farmaci e di comunicare eventuali problemi di sicurezza al medico prescrittore.

L'integrazione dell'IA nella farmacovigilanza apre nuove prospettive per questo ruolo. I sistemi di IA potrebbero supportare il farmacista clinico nell'identificazione di potenziali reazioni avverse, analizzando automaticamente il profilo terapeutico completo del paziente e segnalando combinazioni di farmaci associate a un rischio aumentato. In questo scenario, il farmacista non modifica la terapia autonomamente, ma utilizza le indicazioni del sistema di IA come supporto per comunicare tempestivamente il rischio identificato al medico prescrittore, contribuendo così a una farmacovigilanza più proattiva e sistematica (52).

La farmacogenomica assume un ruolo di particolare rilievo anche in questo contesto: un paziente metabolizzatore povero del CYP2D6 che assume un antidepressivo metabolizzato da questo enzima potrebbe accumulare il farmaco fino a livelli tossici. Un sistema di IA integrato nella cartella clinica elettronica potrebbe identificare automaticamente questa situazione di rischio sulla base del profilo genetico del paziente e segnalarla al farmacista, che potrebbe a sua volta informare tempestivamente il medico prescrittore prima che si verifichi una reazione avversa (46,52).

Al tempo stesso, l'utilizzo di questi strumenti richiede che il farmacista clinico sviluppi nuove competenze: la capacità di interpretare l'output di algoritmi di IA, di valutarne criticamente l'affidabilità e i limiti, e di integrare le raccomandazioni automatizzate con il proprio giudizio professionale e la conoscenza del singolo paziente. L'IA non sostituisce il ragionamento farmacologico del professionista, ma lo potenzia fornendo strumenti di supporto decisionale più potenti e tempestivi, rafforzando così il ruolo del farmacista clinico come figura chiave nella tutela della sicurezza del paziente (52).

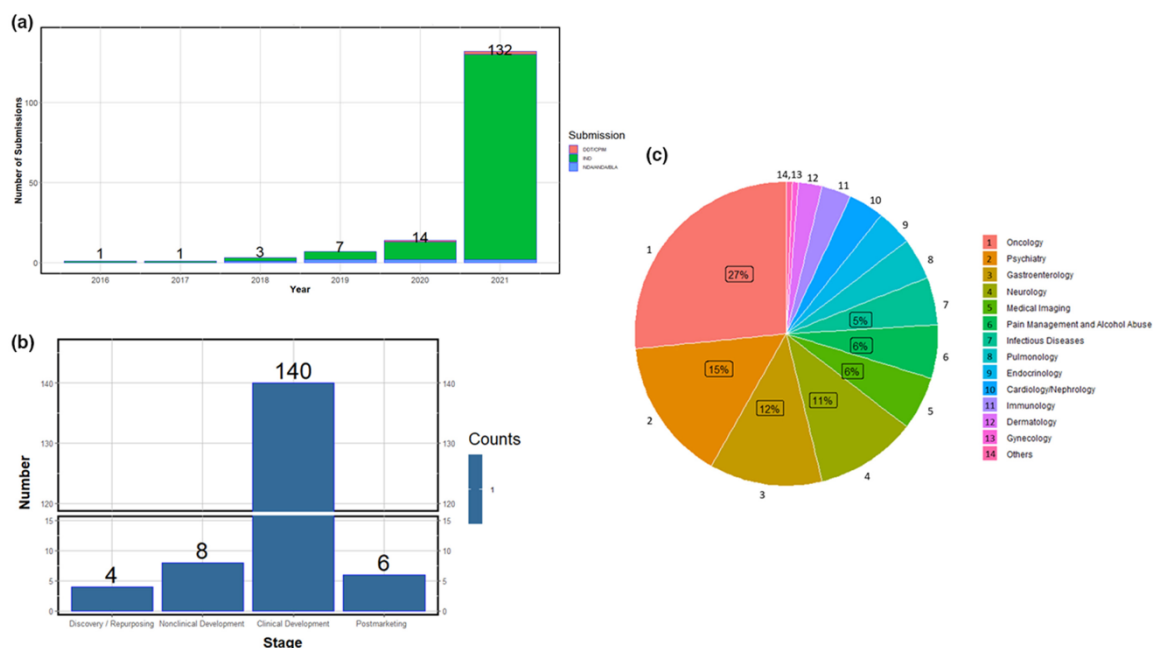
## 5. ASPETTI REGOLATORI E STANDARD DI QUALITÀ NELLO SVILUPPO DI FARMACI BASATI SULL' IA

L'introduzione dell'IA nel ciclo di vita del farmaco solleva una questione che trascende il piano puramente tecnico-scientifico: quella della affidabilità regolatoria. Un modello algoritmico, per quanto sofisticato, non può essere accolto come supporto alle decisioni in materia di sicurezza, efficacia o qualità di un medicinale se non è possibile dimostrarne, in modo documentato e verificabile, la fiducia che merita per lo specifico impiego previsto (53). Le principali autorità regolatorie, la statunitense *Food and Drug Administration* (FDA) e la *European Medicines Agency* (EMA), hanno affrontato questo tema intensamente a partire dal 2023 (54), concludendo, nel gennaio 2026, con un documento congiunto di principi comuni.

Il presente capitolo ricostruisce il quadro regolatorio di riferimento, ne analizza i dieci principi guida per una “buona pratica dell'IA” e discute criticamente le sfide ancora aperte in tema di validazione, trasparenza, tutela dei dati e cybersecurity.

### 5.1 Il panorama regolatorio: le posizioni di FDA ed EMA

Il punto di partenza comune alle due agenzie è una constatazione empirica: negli ultimi anni l'impiego di IA e ML nelle sottomissioni regolatorie, cioè nei dossier che le aziende presentano alle autorità regolatorie, è cresciuto in modo marcato (55). Un'analisi delle pratiche presentate al *Center for Drug Evaluation and Research* (CDER) tra il 2016 e il 2021 ha infatti documentato un numero crescente di proposte che includevano componenti di IA/ML, distribuite su diverse aree terapeutiche, come l'oncologia, la psichiatria, la gastroenterologia e la neurologia (55), (**Figura 20**). Di fronte a questa rapida diffusione, entrambe le agenzie hanno scelto un approccio “basato sui principi” e proporzionato al rischio, preferendo quindi fissare principi generali e flessibili, calibrando i controlli in base a quanto sarebbero gravi le conseguenze di un eventuale errore del modello (53,54), evitando, almeno in questa fase, regole prescrittive e rigide che rischierebbero di diventare rapidamente obsolete. È possibile ritenere che questa scelta risponda a un'esigenza precisa: l'IA evolve molto più rapidamente di quanto si aggiornino le norme, e una regola troppo legata a una specifica tecnica rischierebbe di risultare superata poco dopo la sua adozione, finendo per ostacolare l'innovazione anziché guidarla.



**Figura 20.** Distribuzione delle sottomissioni regolatorie contenenti componenti di intelligenza artificiale/machine learning (IA/ML) presentate al *Center for Drug Evaluation and Research* (CDER) della FDA nel periodo 2016–2021: (a) per anno; (b) per area terapeutica; (c) per fase del ciclo di sviluppo del farmaco (55).

### 5.1.1 L'approccio della FDA: il framework di valutazione della credibilità

Dopo la pubblicazione di un *discussion paper* nel maggio 2023 (rivisto nel febbraio 2025) (56) e la raccolta di oltre 800 commenti (57), all'inizio di gennaio 2025 la FDA ha emanato la propria prima bozza di linea guida dedicata ai farmaci, intitolata *Considerations for the Use of Artificial Intelligence to Support Regulatory Decision-Making for Drug and Biological Products* (12). Il documento, preparato dal CDER in collaborazione con CBER (*Center for Biologics Evaluation and Research*), CDRH (*Center for Devices and Radiological Health*), CVM (*Center for Veterinary Medicine*), Oncology Center of Excellence e Office of Combination Products, rappresenta, una volta finalizzato, il primo quadro organico, cioè il primo documento che organizza in modo completo le regole per la progettazione, lo sviluppo, la documentazione e il controllo continuativo dei modelli di IA usati per supportare le decisioni regolatorie.

Un elemento cruciale, spesso trascurato, riguarda l'ambito di applicazione. La linea guida si applica all'IA che produce informazioni o dati a sostegno di decisioni su sicurezza, efficacia o qualità nelle fasi non cliniche, cliniche, post-marketing e di manufacturing; resta invece esplicitamente esclusa l'IA impiegata nella scoperta del farmaco (*drug discovery*) e nelle

attività di mera efficienza operativa interna, quindi per gli usi amministrativi che non incidono sulla sicurezza del paziente, sulla qualità del prodotto o sull'attendibilità degli studi (12). Si tratta di una scelta significativa: l'autorità riconosce che l'IA possa comprimere i tempi della ricerca esplorativa, ma mantiene un atteggiamento prudente sul suo ruolo nelle attività regolatorie critiche, in particolare la FDA accetta che l'IA acceleri la fase iniziale di ricerca, ma resta prudente quando l'IA incide su decisioni importanti per la sicurezza o per l'approvazione.

L'elemento centrale del documento è il *risk-based credibility assessment framework*, un processo articolato in sette fasi (12). Per *credibilità* (credibility) si intende la fiducia, supportata da evidenze, nelle prestazioni di un modello per uno specifico *contesto d'uso* (Context of Use, COU), ossia il ruolo e l'ambito che il modello assume nel rispondere a una determinata domanda (12). La credibilità non è dunque una proprietà assoluta del modello, ma è sempre relativa all'impiego.

Il processo si articola in sette fasi:

1. definizione della domanda di interesse;
2. definizione del contesto d'uso;
3. valutazione del rischio del modello;
4. elaborazione di un piano per stabilirne la credibilità;
5. esecuzione del piano;
6. documentazione dei risultati ed eventuali scostamenti;
7. determinazione dell'adeguatezza del modello per il contesto d'uso (12).

Il rischio del modello (*model risk*) è definito come combinazione di due fattori indipendenti: l'influenza del modello (*model influence*), ossia il peso dei risultati prodotti dall'IA rispetto alle altre evidenze disponibili per rispondere alla domanda, e la conseguenza della decisione (*decision consequence*), la gravità di un esito avverso derivante da una decisione errata (12). Tanto più alto è il rischio così determinato, tanto più stringenti dovranno essere le attività di validazione, i criteri di accettazione delle prestazioni e il livello di documentazione richiesto (12). La FDA chiarisce con due casi esempio il concetto: un modello che stratifica i pazienti per il monitoraggio ospedaliero dopo la somministrazione (rischio alto, perché è l'unico criterio su cui si fonda una decisione che può avere conseguenza gravissime per il paziente) e un sistema di visione artificiale per il controllo del volume di riempimento dei flaconi (rischio medio, perché affiancato da test di controllo di qualità indipendenti per il rilascio sul mercato) (12). È interessante notare come i concetti chiave del framework derivino, da come dichiara la FDA, dallo standard tecnico preesistente ASME V&V40, originariamente sviluppato per i modelli computazionali dei dispositivi medici (12).

La linea guida insiste inoltre su due aspetti destinati a ricorrere in questo capitolo: la qualità dei dati, che devono essere *fit for use* (rilevanti, rappresentativi e affidabili in termini di accuratezza, completezza e tracciabilità), e il *mantenimento lungo il ciclo di vita* (life cycle maintenance), reso necessario dal fenomeno del *data drift*, cioè il calo delle prestazioni del modello quando i dati reali che riceve nel tempo diventano diversi da quelli su cui era stato addestrato (12). Coerentemente con tale impostazione, l'agenzia incoraggia un coinvolgimento precoce (early engagement) tra sponsor, cioè l'azienda che sviluppa il farmaco, e autorità (12). Sul piano operativo, l'evoluzione è già visibile: nel dicembre 2025 la FDA ha riconosciuto ufficialmente come valido il primo strumento basato su IA per l'uso in studi clinici di sviluppo, una piattaforma di supporto alla valutazione di biopsie epatiche negli studi clinici sulla steatoepatite metabolica (NASH/MASH) (58), pur trattandosi di uno strumento di valutazione istologica a supporto degli endpoint degli studi clinici e non di progettazione del farmaco. Tuttavia, sul fronte dei farmaci veri e propri, a fine 2025 i candidati scoperti tramite IA più avanzati si trovavano ancora nelle fasi cliniche e nessuno aveva ottenuto l'approvazione (59).

### **5.1.2 L'approccio dell'EMA: il *reflection paper* e la strategia di rete**

In Europa il percorso regolatorio sull'IA nei medicinali è stato avviato con il *Reflection paper on the use of Artificial Intelligence (AI) in the medicinal product lifecycle* (54), la cui bozza è stata posta in consultazione pubblica nel luglio 2023 e la cui versione definitiva è stata adottata da CHMP e CVMP nel settembre 2024 (54). Rispetto alla FDA, che esclude dal proprio ambito la fase di scoperta del farmaco e si concentra più sull'IA a supporto delle decisioni regolatorie, l'EMA adotta una prospettiva più ampia: il documento copre l'intero ciclo di vita del medicinale, quindi dalla scoperta, allo sviluppo non clinico e clinico, alla valutazione regolatoria, fino al manufacturing e alla fase post-autorizzativa, e si estende anche ai medicinali veterinari e all'IA integrata nei dispositivi medici impiegati negli studi clinici (54).

Due sono i pilastri dell'impostazione europea: un approccio basato sul rischio e un approccio antropocentrico (*human-centred*). L'EMA invita le aziende produttrici a riflettere preliminarmente sul potenziale del sistema di IA di generare rischi per il paziente e, in caso affermativo, a richiedere una consulenza scientifica precoce. Emerge inoltre una chiara preferenza a priori per modelli trasparenti, più semplici e interpretabili, e l'aspettativa che

l'azienda sia pronta a condividere ogni aspetto rilevante dello sviluppo, dell'addestramento (incluse le fonti dei dati e la gestione dei bias) e del monitoraggio in uso del modello (54).

Il *reflection paper* non è un atto isolato, ma si inserisce in una strategia istituzionale più vasta. La *European medicines agencies network strategy to 2028* (EMANS 2028) individua tra i propri sei assi tematici quello del “*leveraging data, digitalisation and artificial intelligence*”, con l'obiettivo dichiarato di sfruttare il potenziale dell'IA lungo tutto il ciclo di vita del medicinale, garantendo al contempo interoperabilità, qualità dei dati, gestione dei bias e conformità alle norme europee in materia di protezione dei dati, riservatezza e cybersecurity (60). A ciò si aggiunge un apposito piano congiunto HMA-EMA in materia di dati e IA (60), mentre sul piano legislativo questo quadro sarà completato da tre interventi: l'EU AI Act, la proposta di Biotech Act e la riforma della legislazione farmaceutica europea. Quest'ultima, il cui testo di compromesso finale è stato pubblicato a marzo 2026, ma non è ancora stato approvato in via definitiva, introduce i *regulatory sandbox*, cioè spazi sperimentali protetti in cui, sotto la supervisione dell'autorità e con regole temporanee adattate al caso, è possibile testare prodotti o metodi innovativi che non rientrerebbero nei requisiti regolatori standard, e apre a un uso più ampio dei dati e dell'IA come supporto alle decisioni delle autorità regolatorie (61).

Sul versante nazionale, anche l'Agenzia Italiana del Farmaco (AIFA) si sta progressivamente dotando di strumenti predittivi basati su IA, non solo per la valutazione dei dossier regolatori ma anche a supporto dell'Health Technology Assessment (HTA) (62), ossia l'analisi integrata dei benefici clinici, dei costi e dell'impatto sistemico delle nuove tecnologie sanitarie (62). Ne emerge un sistema regolatorio su più livelli, in cui le regole decise a livello europeo vengono via via adottate e applicate dalle autorità nazionali.

### **5.1.3 Convergenze e differenze tra i due approcci**

Pur partendo da sistemi regolatori differenti, FDA ed EMA convergono su alcuni capisaldi: la proporzionalità al rischio, la centralità della supervisione umana, l'importanza della qualità dei dati e della trasparenza metodologica, e la necessità di un monitoraggio continuo dei modelli (54,63). Le differenze riguardano soprattutto il campo di applicazione: la FDA circoscrive il proprio quadro alle attività che incidono sulle decisioni regolatorie, escludendo formalmente la *discovery* (63); l'EMA adotta invece una visione di ciclo di vita completo (54). Questa parziale divergenza del campo di applicazione spiega la rilevanza dell'iniziativa congiunta del 2026, oggetto del paragrafo seguente, con lo scopo di ridurre le differenze tra

le regole statunitensi ed europee e di rendere più uniformi le norme a livello internazionale (64).

## 5.2 Principi guida per le “Good AI Practices”: i dieci principi comuni

### EMA-FDA

Il 14 gennaio 2026 EMA e FDA hanno pubblicato congiuntamente il documento *Guiding principles of good AI practice in drug development*, individuando dieci principi comuni per una buona pratica dell'IA nel ciclo di vita del medicinale (64). L'iniziativa è il primo esito concreto di una rinnovata cooperazione UE-USA e si allinea sia alla EMANS 2028 sia al piano HMA-EMA su dati e IA (64). I principi sono rivolti agli sviluppatori di medicinali e ai titolari (o richiedenti) di autorizzazione all'immissione in commercio, e sono concepiti come fondamento per le future linee guida specifiche nelle rispettive giurisdizioni (64).

È opportuno precisare che i dieci principi EMA-FDA del 2026, dedicati ai *farmaci e prodotti biologici*, non vanno confusi con i dieci principi della *Good Machine Learning Practice* (GMLP) pubblicati nell'ottobre 2021 congiuntamente da FDA, Health Canada e MHRA del Regno Unito, che riguardano invece lo sviluppo dei *dispositivi medici* basati su ML (65). Si tratta di due insiemi distinti, sebbene ispirati a una filosofia comune, cioè l'approccio basato sul rischio, trasparenza, gestione del ciclo di vita; il presente paragrafo si riferisce ai principi del 2026, pertinenti al ciclo di vita del farmaco (64).

Nella definizione adottata dal documento, l'IA è intesa come l'insieme delle tecnologie “a livello di sistema” impiegate per generare o analizzare evidenze attraverso le fasi non clinica, clinica, post-marketing e di manufacturing. I principi ribadiscono che i farmaci continuano a essere autorizzati sulla base di qualità, efficacia e sicurezza dimostrate, e che l'IA deve rafforzare, e non indebolire, tali requisiti, contribuendo anche a ridurre il ricorso alla sperimentazione animale attraverso una migliore predizione di tossicità ed efficacia nell'uomo. I dieci principi sono così sintetizzabili (64):

1. **Antropocentricità per progettazione (human-centric by design).** Lo sviluppo e l'uso delle tecnologie di IA devono allinearsi a valori etici e centrati sull'essere umano.
2. **Approccio basato sul rischio (risk-based approach).** Validazione, mitigazione del rischio e supervisione devono essere proporzionate al contesto d'uso e al rischio del modello, secondo la stessa logica del framework di credibilità della FDA.

3. **Aderenza agli standard (adherence to standards).** Le tecnologie devono rispettare gli standard legali, etici, tecnici, scientifici, di cybersecurity e regolatori applicabili, comprese le Good Practices.
4. **Contesto d'uso chiaro (clear context of use).** Ogni tecnologia di IA deve avere un contesto d'uso ben definito, ossia un ruolo e un ambito espliciti che ne giustificano l'impiego.
5. **Competenze multidisciplinari (multidisciplinary expertise).** Lungo l'intero ciclo di vita devono essere integrate competenze relative sia alla tecnologia di IA sia al suo contesto d'uso (clinico, regolatorio, di data science, ingegneristico).
6. **Governance e documentazione dei dati (data governance and documentation).** Provenienza delle fonti, fasi di elaborazione e scelte analitiche vanno documentate in modo dettagliato, tracciabile e verificabile, con un'adeguata governance che includa la tutela dei dati sensibili lungo tutto il ciclo di vita.
7. **Buone pratiche di progettazione e sviluppo del modello (model design and development).** Lo sviluppo deve seguire le migliori pratiche di ingegneria del software e impiegare dati idonei all'uso, bilanciando interpretabilità, spiegabilità e prestazioni predittive, così da promuovere trasparenza, affidabilità, generalizzabilità e robustezza a tutela della sicurezza del paziente.
8. **Valutazione delle prestazioni basata sul rischio (risk-based performance assessment).** La valutazione deve riguardare l'intero sistema, comprese le interazioni uomo-IA, impiegando dati e metriche appropriati al contesto d'uso e una validazione delle prestazioni predittive condotta con metodi di test ed evaluation adeguatamente disegnati.
9. **Gestione del ciclo di vita (life cycle management).** Vanno implementati sistemi di gestione della qualità basati sul rischio, con monitoraggio programmato e rivalutazione periodica del modello per affrontare fenomeni come il data drift.
10. **Informazione chiara ed essenziale (clear, essential information).** Occorre fornire, in linguaggio semplice e accessibile, informazioni pertinenti al pubblico di riferimento, utenti e pazienti inclusi, su contesto d'uso, prestazioni, limiti, dati sottostanti, aggiornamenti e interpretabilità della tecnologia.

Nel loro insieme, i dieci principi consolidano e armonizzano gli orientamenti già espressi separatamente dalle due agenzie: il framework di credibilità basato sul rischio e sul contesto d'uso della FDA si ritrova nei principi 2, 4 e 8; l'enfasi europea sulla trasparenza, sull'interpretabilità e sull'approccio antropocentrico nei principi 1, 7 e 10; l'attenzione comune alla governance dei dati e alla manutenzione nel tempo nei principi 6 e 9. Si tratta, dichiaratamente, di principi di base e non di regole rigide: sono pensati per cambiare insieme alla tecnologia e per essere completati, in seguito, da linee guida più specifiche. (64).

### **5.3 Le sfide regolatorie: validazione, trasparenza, data privacy e cybersecurity**

Il quadro di principi delineato nei paragrafi precedenti, per quanto solido, lascia aperte alcune sfide importanti che condizionano l'effettiva integrazione dell'IA nei processi regolatori. Esse possono essere ricondotte a quattro nodi principali: la validazione degli algoritmi, la trasparenza (con il connesso problema della "black box"), la tutela dei dati personali e la sicurezza informatica.

#### **5.3.1 La validazione degli algoritmi**

A differenza dei modelli statistici tradizionali, un modello di ML può comportarsi in modo imprevedibile su dati diversi da quelli con cui è stato addestrato e le sue prestazioni dipendono in misura determinante dalla qualità e dalla rappresentatività dei dati, cioè dal fatto che i dati rispecchiano davvero la popolazione su cui il modello verrà poi usato. Per questo entrambe le agenzie insistono sul requisito dei dati *fit for use*, cioè adeguati allo scopo (pertinenti e affidabili), e sulla necessità di una validazione condotta su un gruppo di dati separato, che il modello non ha mai usato per imparare, così da verificarne le reali capacità su casi nuovi (63). Le prestazioni vanno misurate con indicatori adatti al contesto d'uso come la sensibilità, cioè la capacità di individuare correttamente i casi positivi, e la specificità, cioè la capacità di individuare i negativi, accompagnati sempre dal margine di incertezza, e verificando che i risultati restino stabili se la prova viene ripetuta (63).

Una difficoltà di fondo è che l'accuratezza tecnica di un modello non garantisce di per sé la sua validità clinica: un bersaglio individuato con precisione non è necessariamente "*druggable*" e la capacità di accelerare le fasi iniziali della ricerca non si traduce automaticamente in un miglioramento dei tassi di successo clinico, che restano storicamente molto bassi, infatti circa l'86-90% dei programmi di sviluppo non arriva all'approvazione

(66). La validazione sperimentale resta dunque insostituibile, e questo spiega la prudenza con cui le autorità accettano i risultati prodotti dall'IA solo se confermati da prove esterne e indipendenti.

### **5.3.2 La trasparenza e il problema della “black box”**

Molti modelli di deep learning, pur efficaci, sono di fatto "scatole nere": il percorso che porta dal dato di partenza al risultato non è facilmente ispezionabile né comprensibile dall'essere umano. Questa opacità genera un problema regolatorio di primo piano, perché rende difficile verificare la correttezza del ragionamento del modello, individuarne i bias e attribuire le responsabilità in caso di errore. La FDA affronta il tema chiedendo una trasparenza dei metodi: nelle pratiche presentate vanno descritti in dettaglio l'architettura del modello, le variabili usate, i parametri e le ragioni delle scelte fatte (63). L'EMA, dal canto suo, esprime una preferenza per modelli più semplici e interpretabili a parità di prestazioni, soprattutto per gli usi in cui un errore avrebbe conseguenze gravi per il paziente (54); lo stesso orientamento è ripreso dal principio 7 del documento congiunto, che invita a bilanciare interpretabilità, spiegabilità e prestazioni del modello (64).

Legato alla trasparenza c'è anche il modo in cui gli studi sull'IA vengono descritti e pubblicati. Per renderli chiari, confrontabili e ripetibili, la comunità scientifica ha messo a punto apposite linee guida internazionali su come documentarli che fissano regole condivise di trasparenza. Adottarle è uno dei modi concreti con cui i principi delle autorità si trasformano in pratiche verificabili (62).

All'origine del problema della trasparenza c'è poi la questione dei bias del modello. Se i dati usati per addestrarlo non rappresentano bene tutta la popolazione, il modello rischia di funzionare peggio proprio sui gruppi meno presenti nei dati, con conseguenze sull'equità delle cure. Per questo gestire e documentare con cura i dati di addestramento non è solo un'esigenza tecnica, ma anche un modo per tutelare l'equità (54,63).

### **5.3.3 Il data drift e il mantenimento lungo il ciclo di vita**

Un modello validato in un certo momento non resta necessariamente affidabile nel tempo. Il fenomeno del *data drift*, ossia la progressiva "deriva" dei dati in ingresso, che col tempo diventano diversi da quelli su cui il modello era stato addestrato, può peggiorarne gradualmente le prestazioni, soprattutto quando il modello opera in contesti diversi da quello

in cui è stato sviluppato. Per questo la FDA dedica una specifica attenzione al *life cycle maintenance*, cioè ad un insieme pianificato di attività di monitoraggio e di verifica periodica per mantenere il modello affidabile lungo tutta la vita del prodotto (63); allo stesso modo, il principio 9 del documento congiunto richiede procedure di controllo della qualità proporzionate al rischio e verifiche periodiche programmate del modello (64). Nel campo dei dispositivi medici, questa stessa logica ha portato a sviluppare i *Predetermined Change Control Plans* (PCCP), piani che definiscono in anticipo quali modifiche al modello sono ammesse e con quali controlli (67): un approccio che potrebbe ispirare soluzioni simili anche per i modelli usati nel ciclo di vita del farmaco.

### **5.3.4 Tutela dei dati personali e cybersecurity**

L'addestramento e l'uso dei modelli richiedono l'accesso a grandi quantità di dati clinici, spesso molto sensibili. In Europa la gestione di questi dati deve rispettare il Regolamento generale sulla protezione dei dati (GDPR) (68) e le norme su riservatezza e sicurezza informatica, mentre lo spazio europeo dei dati sanitari (European Health Data Space, EHDS) renderà più facile produrre dati e prove utili alle autorità per decidere sui farmaci (69). Il principio 6 del documento congiunto richiede esplicitamente una governance adeguata dei dati, che comprenda la tutela della privacy e dei dati sensibili lungo tutto il ciclo di vita del modello (64).

La sicurezza informatica, pur richiamata tra gli standard di riferimento (principio 3), ha una collocazione particolare: la bozza FDA del 2025 precisa che essa non fa parte del proprio metodo di valutazione della credibilità, pur raccomandando di tenerne conto quando il modello viene messo in funzione (63). Questo non ne riduce l'importanza: l'integrità e la riservatezza dei dati e dei modelli sono presupposti della loro affidabilità e la loro tutela è destinata a essere disciplinata in modo più dettagliato dalle leggi di portata generale in materia di IA e sicurezza informatica, valide per più settori (come l'*EU AI Act*).

### **5.3.5 La questione della responsabilità**

Sullo sfondo di tutte queste sfide tecniche resta una domanda di natura giuridica ed etica: se l'algoritmo sbaglia, di chi è la responsabilità? Del medico, di chi ha sviluppato il software, dell'azienda farmaceutica? Per ora le autorità rispondono in modo indiretto ma coerente: insistendo sulla supervisione umana (principio 1) e sulla valutazione dell'intero "sistema uomo-IA" (principio 8), ribadiscono che la decisione finale, e quindi la responsabilità, deve

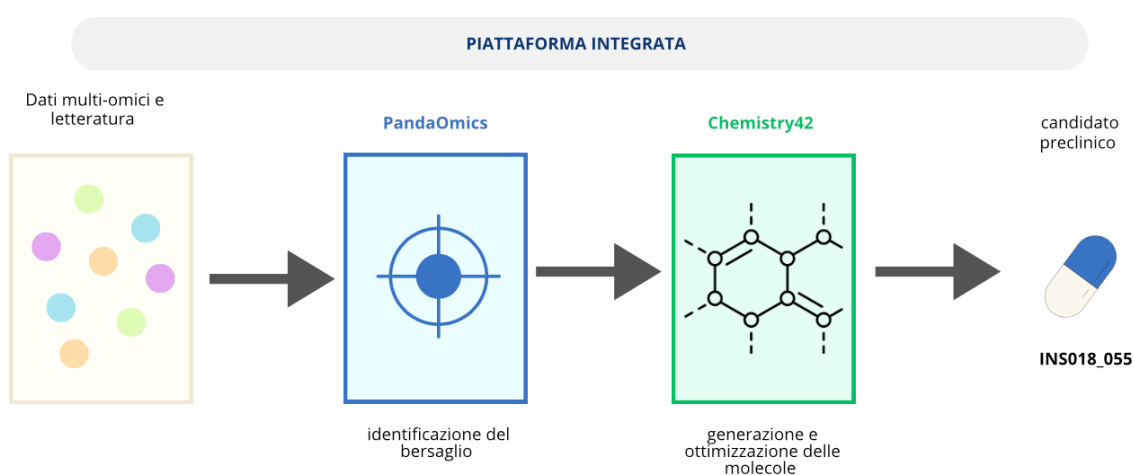
restare in capo a una persona competente (64). L'IA è dunque pensata come uno strumento di supporto e non come un sostituto del giudizio professionale: un principio particolarmente importante per il farmacista clinico, che nella pratica è spesso chiamato a interpretare e verificare i risultati prodotti da questi strumenti.

## 6. CASI DI STUDIO E APPLICAZIONI PRATICHE

Le potenzialità dell'IA descritte nei capitoli precedenti trovano oggi riscontro in programmi concreti di scoperta e sviluppo dei farmaci. Negli ultimi anni un numero crescente di molecole individuate o progettate con il supporto dell'IA è entrato nelle fasi cliniche, rendendo possibile una prima valutazione empirica dei vantaggi attesi (70). Questo capitolo analizza il modello di alcune aziende biotecnologiche nate attorno a queste tecnologie, esamina come l'IA intervenga nell'identificazione dei bersagli terapeutici, presenta molecole arrivate alla sperimentazione clinica e approfondisce due casi studio rappresentativi.

### 6.1 Aziende biotech "AI-native": modelli di business e pipeline

Una parte rilevante dell'innovazione nel settore è guidata da imprese che integrano l'IA lungo l'intero processo di ricerca e sviluppo, anziché applicarla a singole fasi isolate. Tra le pioniere nell'impiego dell'IA generativa per la scoperta di piccole molecole figura *Insilico Medicine* (71). La sua piattaforma articola il processo in due moduli distinti: *PandaOmics* per l'identificazione del bersaglio terapeutico a partire da dati multi-omici e dalla letteratura, e *Chemistry42* per la generazione e l'ottimizzazione delle strutture molecolari (72). Questa organizzazione modulare ha permesso, nel caso dell'inibitore di TNIK per la fibrosi polmonare idiopatica, di completare il percorso dalla scelta del bersaglio alla nomina del candidato preclinico in circa diciotto mesi (72) (**Figura 21**).



**Figura 21.** Architettura modulare della piattaforma di *Insilico Medicine*. Il modulo *PandaOmics* identifica il bersaglio terapeutico a partire da dati multi-omici e dalla letteratura; il modulo *Chemistry42* genera e ottimizza le strutture molecolari fino al candidato preclinico.

L'approccio computazionale alla base di questi risultati è consolidato da tempo. Già nel 2016 era stato dimostrato come reti neurali profonde addestrate su profili trascrizionali potessero classificare i farmaci per categoria terapeutica, aprendo la strada al riposizionamento computazionale dei composti (73). Lo stesso approccio si è successivamente esteso alla previsione dell'esito dei trial: la piattaforma *inClinico*, basata su *architetture transformer*, stima la probabilità di successo degli studi di fase II in funzione del bersaglio scelto e del disegno sperimentale (74).

Le aziende AI-native raramente dispongono del capitale e delle infrastrutture necessari allo sviluppo clinico completo. La sostenibilità di questi modelli si fonda in larga parte su accordi di licenza con grandi gruppi farmaceutici. *Insilico Medicine* ha concesso in licenza a *Exelixis* ISM3091, un inibitore della deubiquitinasi USP1 sviluppato per il trattamento dei tumori con mutazione BRCA, pronto per la fase regolatoria (75) e ha siglato con *Stemline Therapeutics* un accordo per un inibitore di KAT6A destinato al carcinoma mammario ormono-sensibile (76). Una logica analoga caratterizza *Exscientia*, con accordi in oncologia, immunologia e neuroinfiammazione (71).

Questi accordi riflettono un interesse crescente dell'industria verso le piattaforme di IA, testimoniato dall'aumento del valore complessivo delle partnership e dei finanziamenti raccolti nel settore (71). Resta tuttavia aperta la problematica principale, cioè se le molecole così individuate dimostreranno, nelle fasi avanzate, tassi di successo superiori a quelli storici della scoperta tradizionale (70).

## **6.2 L'IA nell'identificazione dei target: il caso dell'oncologia**

Tra le aree in cui le piattaforme di IA descritte hanno espresso il potenziale maggiore, l'oncologia occupa una posizione di primo piano. Essa rappresenta infatti l'area terapeutica in cui l'applicazione dell'IA per l'identificazione e la validazione dei target farmacologici ha registrato i progressi più significativi. Questa predominanza non è casuale: i tumori sono patologie caratterizzate da un'elevata complessità molecolare, in cui alterazioni genetiche, epigenetiche e proteomiche si combinano in modi difficilmente prevedibili con i metodi tradizionali. L'analisi di grandi volumi di dati genomici, trascrittomici e proteomici provenienti da coorti di pazienti oncologici ha reso l'IA uno strumento indispensabile per gestire questa complessità e identificare bersagli terapeutici rilevanti (15,29).

Nella pratica del *drug discovery* oncologico, l'identificazione del target si basa sulla dimostrazione che una specifica molecola, tipicamente una proteina, svolge un ruolo diretto nella progressione tumorale e che la sua modulazione farmacologica produce un effetto clinicamente rilevante. L'IA contribuisce a questo processo attraverso l'analisi di reti biologiche complesse, l'integrazione di dati multi-omici e l'applicazione di modelli predittivi in grado di classificare le proteine candidate in base alla loro probabilità di essere target farmacologici efficaci (15).

In questo ambito, Jeon et al. (77) hanno proposto un classificatore basato su *support vector machine*, un algoritmo di ML in grado di distinguere automaticamente due categorie di dati sulla base delle loro caratteristiche, addestrato su dati genomici tra cui essenzialità genica, espressione dell'mRNA, numero di copie del DNA e topologia delle reti di interazione proteica, per distinguere proteine target da proteine non-target in diversi tipi di tumore, tra cui carcinoma mammario, pancreatico ed ovarico (77). Questo approccio ha permesso di identificare 122 target oncologici globali e centinaia di target specifici per singoli tipi tumorali, alcuni dei quali validati sperimentalmente tramite inibitori peptidici che hanno mostrato effetti antiproliferativi significativi in modelli cellulari (77).

Più in generale, l'IA si è dimostrata particolarmente efficace nell'analisi di dati ad alta dimensionalità tipici dell'oncologia, come i profili di espressione genica, i dati di metilazione del DNA e le varianti somatiche. L'integrazione di questi livelli informativi in modelli di ML consente di identificare firme molecolari associate alla progressione tumorale, alla resistenza ai farmaci e alla risposta alle terapie target (15). In questo contesto, i modelli multi-task basati su *deep learning* hanno mostrato prestazioni superiori rispetto ai metodi tradizionali, potendo valutare simultaneamente l'influenza di molteplici variabili genomiche sulla sensibilità ai farmaci antitumorali (15).

Dal punto di vista farmacologico, l'identificazione dei target oncologici tramite IA si inserisce in un quadro più ampio che include la prioritizzazione dei candidati in base alla loro *druggability*, ovvero alla probabilità che il bersaglio possa essere modulato efficacemente da un farmaco. Questo processo tiene conto di caratteristiche strutturali della proteina, della sua accessibilità farmacologica e del rischio di effetti *off-target*, cioè di interazioni indesiderate del farmaco con proteine diverse dal bersaglio terapeutico di interesse, aspetti che i modelli computazionali sono in grado di valutare simultaneamente sull'intero genoma anziché uno alla volta come nei metodi tradizionali (15).

Tra i target oncologici di maggiore rilevanza clinica e farmacologica oggi al centro delle pipeline assistite da IA figurano i recettori tirosin-chinasici, tra cui il recettore del fattore di crescita epidermico (EGFR), e le proteine del checkpoint immunitario, tra cui PD-1 e il suo ligando PD-L1. Questi target sono stati individuati come prioritari da numerose aziende farmaceutiche che impiegano piattaforme IA, a conferma della convergenza del settore verso un insieme ristretto di molecole ad alto valore terapeutico nell'immuno-oncologia (29). Le piattaforme consentono oggi non solo di identificare tali target, ma anche di analizzare le alterazioni molecolari che ne determinano l'attivazione o la sovraespressione nei diversi sottotipi tumorali, orientando la progettazione di farmaci selettivi e il loro abbinamento ai profili genomici dei pazienti (15,29).

Un esempio concreto di questa progressione, dal bersaglio identificato in silico alla molecola in sviluppo, è dato dai programmi oncologici già citati: l'inibitore di USP1 e quello di KAT6A, entrambi individuati a partire da bersagli prioritizzati con metodi computazionali e successivamente concessi in licenza da *Insilico Medicine* a partner industriali per lo sviluppo clinico (71).

### **6.3 Farmaci scoperti con l'IA**

Accanto all'identificazione del bersaglio, l'IA è stata impiegata direttamente per individuare nuove entità terapeutiche. Un esempio emblematico riguarda la scoperta di nuovi antibatterici: un modello di *deep learning*, addestrato a prevedere l'attività antibatterica sulla base della sola struttura molecolare, ha consentito di identificare l'*halicina*, una molecola strutturalmente distante dagli antibiotici noti e attiva contro ceppi batterici multiresistenti (78). Il caso è rilevante perché dimostra la capacità di questi modelli di esplorare regioni dello spazio chimico difficilmente raggiungibili con lo screening convenzionale (79).

Questi risultati si inseriscono in una tendenza più ampia. Negli ultimi anni diverse molecole scoperte o progettate con il supporto dell'IA sono entrate nelle fasi cliniche, e una prima analisi dei loro esiti suggerisce tassi di successo nelle fasi precoci potenzialmente superiori alle medie storiche, pur con la cautela imposta dal numero ancora ridotto di molecole giunte in clinica e dal tempo insufficiente per osservarne gli esiti nelle fasi avanzate (70).

## 6.4 Collaborazioni tra industria farmaceutica e aziende tecnologiche

Lo sviluppo clinico richiede competenze, infrastrutture e capitali che le aziende IA-native raramente possiedono. Per questo il modello prevalente è quello della collaborazione strategica con grandi gruppi farmaceutici, che apportano esperienza regolatoria e capacità di sperimentazione in cambio dell'accesso alle piattaforme computazionali (71). Negli ultimi anni si è osservata una proliferazione di accordi, sia per la progettazione di piccole molecole sia per quella di farmaci biologici (71).

Oltre alla progettazione delle molecole, la collaborazione tra industria e aziende tecnologiche ha riguardato anche la generazione e l'analisi dei dati sperimentali. È stato dimostrato, ad esempio, come dati di imaging cellulare ad alta produttività possano essere riutilizzati per addestrare modelli capaci di prevedere l'attività biologica di nuovi composti, riducendo il ricorso a saggi sperimentali costosi e di difficile applicazione, in particolare nei sistemi cellulari fisiologicamente più rilevanti (80). In questo schema il dato industriale alimenta il modello e il modello orienta l'esperimento successivo, e ciò rappresenta una delle modalità più mature di integrazione tra le due realtà (80).

## 6.5 Caso studio 1: inibitore di TNIK per la fibrosi polmonare idiopatica

La fibrosi polmonare idiopatica (IPF) è una malattia interstiziale progressiva e a prognosi sfavorevole, con una sopravvivenza media di due-tre anni nei pazienti non trattati e opzioni terapeutiche mirate limitate a Nintedanib e Pirfenidone (72). In questo contesto di forte bisogno clinico, la piattaforma di *Insilico Medicine* ha individuato la chinasi TNIK (*TRAF2-and NCK-interacting kinase*) come bersaglio anti-fibrotico (72).

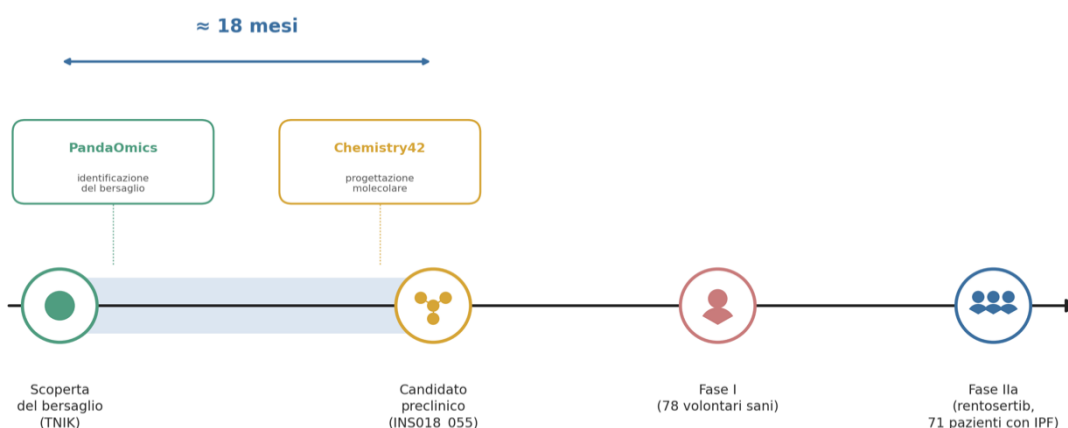
TNIK è stata classificata come bersaglio prioritario dal modulo di identificazione (*PandaOmics*), che ne ha valutato la rilevanza a partire da dati multi-omici e dalla letteratura, sulla base della sua connessione con vie note di progressione fibrotica, tra cui le vie di segnalazione WNT, TGF- $\beta$ , Hippo e NF- $\kappa$ B; pur essendo coinvolta in tali processi, non era mai stata studiata come bersaglio terapeutico nell'IPF (72).

A partire dalle strutture cristallografiche del dominio chinasi, il modulo di progettazione molecolare ha generato strutture in grado di legare il sito di legame dell'ATP, fino al composto INS018-055, un inibitore di TNIK a piccola molecola sviluppato per

somministrazione orale (72). La molecola, in seguito identificata con la denominazione comune internazionale rentosertib, è progredita nelle fasi cliniche successive (32). L'inibizione di TNIK ostacola le transizioni epitelio-mesenchimale e fibroblasto-miofibroblasto indotte dal TGF- $\beta$ , processi centrali nella fibrosi (72).

In un saggio di selettività condotto su un ampio pannello di chinasi, INS018-055 ha inibito TNIK con una concentrazione inibitoria semi-massimale (IC<sub>50</sub>) di 31 nM, valore che ne indica l'elevata potenza. Il composto ha mostrato buona selettività: le poche altre chinasi inibite a concentrazioni nanomolari risultavano anch'esse coinvolte in processi profibrotici, il che limita il rischio di effetti off-target indesiderati (72).

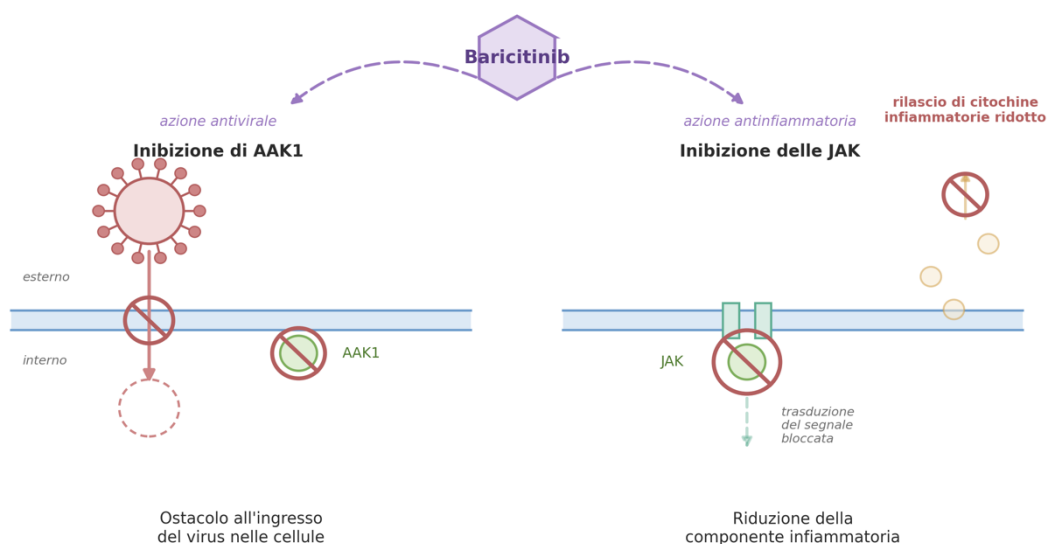
L'intero percorso, dalla scelta del bersaglio alla nomina del candidato preclinico, è stato completato in circa diciotto mesi (72). La sicurezza, la tollerabilità e la farmacocinetica del composto sono state successivamente valutate in uno studio clinico di fase I randomizzato, in doppio cieco e controllato con placebo, condotto su settantotto partecipanti sani (72). Il programma è poi proseguito nelle fasi successive: rentosertib ha completato uno studio di fase IIa, nel quale ha evidenziato un miglioramento dose-dipendente della funzione polmonare nei pazienti con fibrosi polmonare idiopatica (32). Il caso illustra come l'integrazione dei moduli di identificazione del bersaglio e di progettazione molecolare possa comprimere in modo sostanziale le tempistiche delle fasi iniziali dello sviluppo (32) (**Figura 22**).



**Figura 22.** Tempistiche del programma sull'inibitore di TNIK. Il percorso dalla scoperta del bersaglio alla nomina del candidato preclinico è stato completato in circa diciotto mesi; la molecola è stata successivamente valutata in uno studio di fase I su settantotto volontari sani. Elaborazione propria su dati di (32,72)

## 6.6 Caso studio 2: il riposizionamento del Baricitinib nella COVID-19

Un percorso complementare alla scoperta *de novo* è il riposizionamento (*drug repurposing*), ossia l'individuazione di nuove indicazioni per farmaci già approvati. L'esempio più noto di riposizionamento guidato dall'IA riguarda il baricitinib durante la pandemia di COVID-19. Nelle prime settimane del 2020, interrogando un *knowledge graph*, ovvero una rappresentazione strutturata delle relazioni tra geni, proteine, malattie e farmaci estratta dalla letteratura biomedica, il baricitinib, un inibitore delle Janus chinasi (JAK) già impiegato nell'artrite reumatoide, è stato identificato come candidato per il trattamento della patologia (81). L'ipotesi alla base era duplice: l'inibizione della chinasi AAK1, regolatrice dell'endocitosi, avrebbe potuto ostacolare l'ingresso del virus nelle cellule, mentre l'azione sulle JAK avrebbe ridotto la componente infiammatoria (81) (**Figura 23**).



**Figura 23.** Duplice meccanismo d'azione ipotizzato per il baricitinib nella COVID-19. L'inibizione di AAK1 ostacola l'ingresso del virus nelle cellule; l'inibizione delle Janus chinasi riduce la componente infiammatoria. Elaborazione propria su dati di (81).

La plausibilità del meccanismo è stata successivamente sostenuta da evidenze sperimentali, a conferma dell'ipotesi formulata per via computazionale (82). La validazione clinica è giunta con uno studio randomizzato, in doppio cieco e controllato con placebo, nel quale l'associazione di Baricitinib e Remdesivir si è dimostrata superiore al solo Remdesivir nel ridurre i tempi di recupero dei pazienti ospedalizzati, con una minore incidenza di eventi avversi gravi (83).

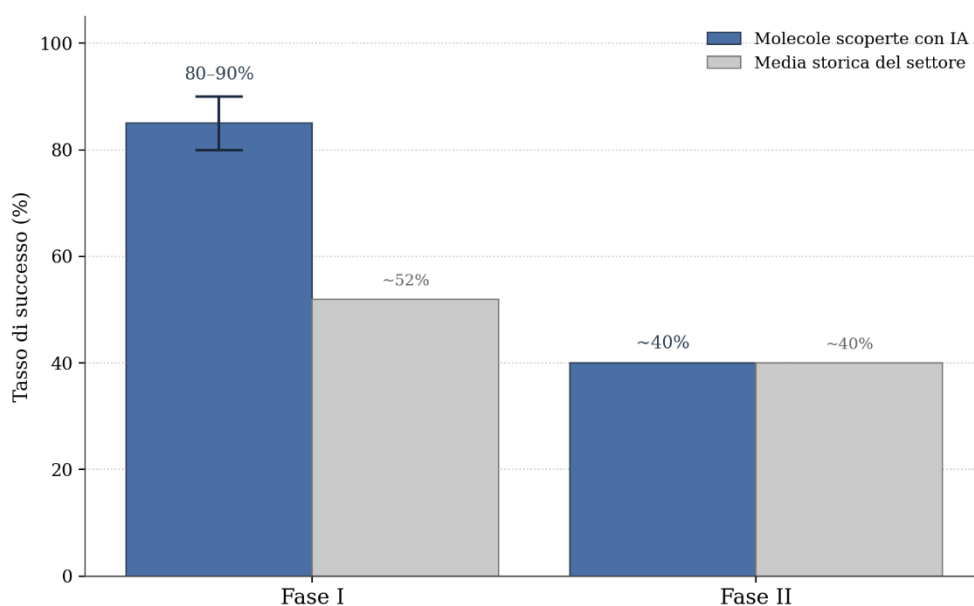
Il caso è particolarmente istruttivo perché mostra l'intero arco di un'applicazione riuscita: dall'ipotesi generata per via computazionale, alla conferma del meccanismo, fino alla validazione in uno studio clinico controllato e al successivo impiego terapeutico. Dal punto di vista farmaceutico, il riposizionamento offre un vantaggio rilevante: il profilo di sicurezza del farmaco è già stato caratterizzato nei modelli preclinici e nell'uomo, e gran parte della valutazione preclinica risulta già completata, con una conseguente riduzione dei tempi e dei rischi rispetto allo sviluppo di una nuova entità chimica (84).

## **6.7 L'IA a supporto delle decisioni terapeutiche: un esempio dalla diagnostica oncologica**

Sebbene la scoperta dei farmaci costituisca il nucleo di questo capitolo, vale la pena richiamare brevemente un ambito contiguo, in cui l'IA incide direttamente sull'impiego dei farmaci, ossia la selezione dei pazienti tramite biomarcatori. In oncologia, l'efficacia delle terapie a bersaglio molecolare e dell'immunoterapia dipende dall'identificazione dei pazienti che ne trarranno beneficio. I modelli di deep learning applicati alle immagini istologiche colorate con ematossilina-eosina hanno mostrato di poter rilevare alterazioni molecolari rilevanti per la scelta terapeutica, come l'instabilità dei microsatelliti nel carcinoma del colon-retto, un marcatore che orienta l'uso degli inibitori dei checkpoint immunitari (85). Analogamente, l'analisi computazionale dell'architettura spaziale tra cellule tumorali e linfociti infiltranti il tumore si è dimostrata in grado di predire la risposta agli inibitori di PD-1 e PD-L1 in pazienti con tumore polmonare e ginecologico (86).

## 7. SFIDE, PROSPETTIVE FUTURE E CONSIDERAZIONI ETICHE

Le evidenze attuali delineano un settore in rapida trasformazione, nel quale l'IA si configura come uno strumento di supporto in grado di accelerare i processi, ridurre i costi e migliorare la probabilità di successo di scoperta e sviluppo farmacologico (29,30). Permane tuttavia un divario significativo tra le potenzialità attribuite a tali tecnologie e il loro impatto clinico. Nonostante il numero crescente di farmaci progettati mediante IA giunti alla sperimentazione clinica, ad oggi nessuno ha ottenuto l'approvazione delle agenzie regolatorie (29). Il candidato più avanzato, rentosertib, ha completato uno studio di fase IIa con risultati positivi, rappresentando la prima validazione clinica di *proof-of-concept* di un farmaco con bersaglio identificato e composto progettato mediante IA (32). Gli stessi autori sottolineano tuttavia la numerosità campionaria limitata e la necessità di validare tali risultati in studi di fase successiva su popolazioni più ampie ed eterogenee, condizione che colloca il farmaco ancora a distanza dall'approvazione (32). Va inoltre osservato che la categoria stessa di "farmaco progettato dall'IA" non corrisponde a una classificazione regolatoria definita, poiché tali composti comportano un rilevante intervento umano nelle fasi di sviluppo, circostanza che rende complessa l'attribuzione (42). Un'analisi delle pipeline cliniche delle aziende AI-native ha inoltre evidenziato come le molecole scoperte mediante IA presentino in fase I un tasso di successo dell'80-90 %, superiore alle medie storiche, mentre in fase II tale tasso si riduca a circa il 40%, in linea con i valori tradizionali del settore (42) (**Figura 24**).



**Figura 24.** Tassi di successo delle molecole scoperte mediante IA in fase I e fase II, confrontati con le medie storiche del settore. In fase I le molecole AI mostrano un tasso dell'80-90%, superiore al riferimento storico; in fase II il tasso si riduce a circa il 40%, in linea con i valori tradizionali. Figura ottenuta dai dati presenti in Jayatunga et al. (42).

Questo andamento suggerisce che l'IA può eccellere nella generazione di molecole con proprietà farmacologiche adeguate, ma che il vantaggio si attenua nelle fasi in cui assume rilievo l'efficacia clinica. Questo divario riflette una serie di ostacoli che non sono prevalentemente di natura computazionale, ma riguardano la qualità dei dati, l'affidabilità dei modelli e le implicazioni etiche e professionali del loro impiego.

## 7.1 Sfide tecniche e scientifiche

Nonostante i risultati promettenti descritti nei capitoli precedenti, l'applicazione dell'IA nello sviluppo dei farmaci presenta limiti strutturali che derivano dalla natura stessa dei dati chimici e biologici. Bender e Cortés-Ciriano (28) hanno osservato che i progressi ottenuti dall'IA nel riconoscimento di immagini e del linguaggio non si sono tradotti in avanzamenti analoghi nel *drug discovery* (28). La ragione di questa differenza risiede verosimilmente nei dati. Nel dominio delle immagini un oggetto è rappresentato da pixel con una disposizione spaziale definita, e l'etichetta, ossia la categoria assegnata al dato per addestrare il modello, può essere attribuita in modo non ambiguo sulla base del solo contenuto visivo: un'immagine raffigura un gatto oppure non lo raffigura. Nel dominio chimico questa corrispondenza non esiste. Sono disponibili circa tremila descrittori molecolari, ciascuno dei quali quantifica un aspetto diverso della molecola, quali il peso molecolare, la lipofilia o la distribuzione delle

cariche. La difficoltà non risiede nel numero di descrittori, ma nel fatto che proprietà biologiche diverse dipendono da aspetti molecolari diversi, e non è noto a priori quale descrittore catturi l'informazione rilevante per una determinata proprietà. A differenza del dominio delle immagini, dove il pixel costituisce un'unità di rappresentazione valida per qualsiasi compito, nel dominio chimico non esiste una rappresentazione di riferimento, e la sua scelta avviene per tentativi (28). Gli autori sintetizzano questo limite osservando che, nella maggior parte dei casi, non si conosce quali caratteristiche di una struttura determinino un dato effetto biologico (28).

A ciò si aggiunge il limite delle rappresentazioni molecolari più diffuse. Le notazioni basate sulla connettività, come il formato SMILES o i grafi molecolari, descrivono i legami presenti tra gli atomi, ma non rappresentano in modo completo le proprietà rilevanti per la modellazione. Tali rappresentazioni non catturano la stereochimica complessa, la direzionalità dei legami a idrogeno né la natura dinamica e tridimensionale delle molecole (28). Il limite dell'IA nel *drug discovery* non è quindi prevalentemente computazionale, ma riguarda la rappresentazione e la quantità dei dati disponibili. Bender e Cortés-Ciriano (28) evidenziano come il volume di dati chimico-biologici sia ridotto rispetto a quello di altri domini: a fronte dei milioni di immagini etichettate disponibili per l'addestramento, i dati annotati con esiti rilevanti in vivo, come i farmaci associati al rischio di danno epatico, ammontano a poche centinaia o migliaia di voci. Questa scarsità deriva dall'elevato costo e dai lunghi tempi richiesti dalla sperimentazione biologica (28).

Una seconda criticità riguarda la qualità e la coerenza dei dati di addestramento. Guo et al.(17) hanno evidenziato come l'affidabilità dei modelli di ML e *deep learning* dipenda direttamente dalla qualità dei dataset utilizzati (17). Nella tossicologia predittiva uno stesso composto può ricevere etichette discordanti, o persino opposte, a seconda del dataset in cui compare. Guo et al. (17) illustrano questo fenomeno con il danno epatico indotto da farmaci, per il quale esistono dataset costruiti a partire da fonti diverse: i foglietti illustrativi approvati, le segnalazioni di reazioni avverse successive alla commercializzazione e i dati estratti dalla letteratura scientifica. Poiché queste fonti definiscono e misurano la tossicità con criteri differenti, un medesimo farmaco può essere classificato come epatotossico in un dataset e come non epatotossico in un altro (17). Tale discordanza compromette la riproducibilità dei modelli e la coerenza delle predizioni. Gli autori osservano inoltre che i dataset tossicologici sono tipicamente di dimensioni limitate, raramente superiori ad alcune migliaia di composti,

e che i modelli sviluppati su dataset più ampi non hanno mostrato prestazioni superiori rispetto a quelli costruiti su dataset più piccoli (17). Ne deriva che la qualità e la coerenza dei dati incidono sulle prestazioni più della loro numerosità.

L'ambiguità delle etichette costituisce un ulteriore ostacolo. Concetti fondamentali come tossicità, efficacia o meccanismo d'azione non rappresentano proprietà intrinseche e univoche di una molecola, ma dipendono da molteplici fattori, quali la dose, il sistema biologico considerato, le condizioni sperimentali e le caratteristiche individuali del paziente (28). Poiché l'apprendimento supervisionato richiede dati a cui sia stata assegnata un'etichetta corretta, l'ambiguità di tali etichette rappresenta un problema sostanziale. Lo stesso limite emerge nella relazione tra attività su un bersaglio molecolare ed effetto clinico. Bender e Cortés-Ciriano (28) hanno analizzato dati di farmacologia secondaria, ossia i dati relativi all'attività di un farmaco su bersagli diversi da quello terapeutico previsto, integrandoli con le segnalazioni di reazioni avverse contenute nel FDA *Adverse Event Reporting System* e nel database SIDER (28). L'analisi evidenzia l'assenza di una relazione univoca tra l'attività su un singolo bersaglio e la comparsa di una reazione avversa: per molti bersagli, anche quando la concentrazione plasmatica del farmaco è sufficiente ad agire sul bersaglio stesso, l'effetto avverso atteso si manifesta solo in una frazione dei casi (28). L'attività su un bersaglio aumenta dunque la probabilità di osservare un determinato effetto, ma non ne costituisce un criterio sufficiente, perché l'effetto complessivo dipende anche dal metabolismo del composto, dal suo profilo farmacocinetico, dalle interazioni con altri bersagli e dalla variabilità individuale (28).

Il caso della ketamina illustra questa complessità. Il suo meccanismo d'azione è stato storicamente attribuito all'antagonismo del recettore NMDA. Tuttavia, altri antagonisti dello stesso recettore, come memantina e lanicemina, non hanno mostrato la medesima efficacia antidepressiva negli studi clinici, circostanza che suggerisce un meccanismo d'azione differente e non ancora pienamente chiarito (28). Studi successivi hanno implicato anche il sistema oppioide e hanno individuato un metabolita della ketamina attivo nei modelli animali di depressione. A ciò si aggiunge la dipendenza dalla dose, poiché il farmaco agisce come anestetico ad alte dosi e manifesta un effetto antidepressivo a dosi più basse (28). Questo esempio dimostra come l'azione di un composto non possa essere ricondotta a una singola interazione molecolare né descritta attraverso un numero limitato di parametri, rendendo difficile l'assegnazione di etichette affidabili per l'addestramento dei modelli.

L'interpretabilità dei modelli rappresenta una sfida tecnica di rilievo. Le reti neurali profonde sono spesso descritte come modelli "*black box*", in quanto il loro processo decisionale è difficile da ricostruire: l'output deriva dalla combinazione di numerosi strati e di milioni di parametri, e risulta difficile identificare quali caratteristiche del dato in ingresso abbiano determinato una specifica predizione (17,43). Questa scarsa trasparenza ostacola l'applicazione in ambito regolatorio. Le agenzie devono infatti poter comprendere il razionale alla base di una predizione per poterla validare e approvare, e un modello la cui logica non è ricostruibile non soddisfa questo requisito (43). I modelli basati su alberi decisionali, come il *random forest*, risultano intrinsecamente più interpretabili, poiché la classificazione segue un percorso di regole esplicite che può essere ripercorso passo dopo passo; le reti neurali profonde, al contrario, non rendono trasparente tale percorso (17). Per attenuare questo limite sono state proposte tecniche dedicate, quali l'analisi dell'importanza delle variabili, che quantifica il contributo di ciascuna caratteristica alla predizione, i metodi di estrazione di regole e gli strumenti SHAP e LIME, che ricostruiscono a posteriori il peso dei singoli fattori in una specifica decisione del modello (43).

Una criticità ulteriore riguarda la validazione dei modelli e il controllo dei bias metodologici. La performance apparente di un algoritmo non è necessariamente indicativa della sua validità clinica, poiché un modello può ottenere risultati eccellenti sui dati con cui è stato sviluppato senza che ciò si traduca in un beneficio reale per il paziente (6). Briganti e Le Moine (87) hanno dimostrato che la maggior parte degli studi che confrontano l'efficacia dell'IA con quella dei clinici presenta un disegno poco affidabile e una carenza di validazione primaria, ossia di una verifica degli algoritmi su campioni provenienti da fonti diverse da quella utilizzata per l'addestramento (87). Gli autori segnalano in particolare il rischio di *overfitting*, condizione in cui un modello si adatta in modo eccessivo ai dati di addestramento, memorizzandone anche le componenti casuali e prive di valore generale. Un modello in questa condizione ottiene risultati eccellenti sui dati con cui è stato sviluppato, ma non li replica su dati nuovi, perché non ha appreso le relazioni generali bensì le particolarità del singolo dataset (87). Una meta-analisi sul confronto tra software di deep learning e radiologi ha rilevato che il 99% degli studi non presentava un disegno affidabile, e che solo una frazione minima validava i risultati su immagini provenienti da popolazioni di origine diversa, condizione necessaria per stimare la reale capacità di generalizzazione del modello. Tali osservazioni sostengono la necessità di sottoporre le tecnologie basate su IA ad una validazione estensiva mediante studi clinici rigorosi (87).

Un esempio concreto di bias metodologico è descritto da Bretthauer et al. (88), i quali hanno commentato uno studio nel quale un modello di deep learning era stato sviluppato per predire la sopravvivenza dei pazienti con carcinoma coloretale a partire da immagini istopatologiche (88). Gli autori hanno evidenziato come tale studio non avesse considerato il *lead-time bias*. Questo bias si verifica perché la sopravvivenza viene misurata a partire dal momento della diagnosi: quando un tumore viene individuato precocemente tramite screening, la diagnosi è anticipata rispetto alla comparsa dei sintomi, e il tempo di sopravvivenza misurato risulta maggiore anche in assenza di un reale prolungamento della vita del paziente o di un beneficio terapeutico (88). Per il carcinoma coloretale il lead time, ossia l'intervallo tra la diagnosi da screening e quella che sarebbe avvenuta su base clinica, è stato stimato in almeno due anni (88). Ad esso si associa l'*overdiagnosis bias*, che si verifica quando lo screening individua neoplasie a evoluzione così lenta da non manifestarsi clinicamente nel corso della vita del paziente, e che gli autori descrivono come una forma estrema di lead-time bias (88). L'utilizzo di dataset che non distinguono i tumori diagnosticati tramite screening da quelli identificati in seguito a sintomi clinici introduce quindi una distorsione sistematica nei modelli (88). Per ovviare a tale distorsione, Bretthauer et al. (88) raccomandano di stratificare le analisi in base alla modalità di diagnosi, separando i casi da screening da quelli sintomatici, e di includere tale variabile tra le informazioni riportate negli standard di rendicontazione degli studi predittivi (88). Il riconoscimento di questi bias è ritenuto cruciale per lo sviluppo di strumenti realmente affidabili.

## 7.2 Prospettive future

Le prospettive future dell'uso dell'IA nello sviluppo farmaceutico sono legate all'integrazione di modelli generativi avanzati, allo sviluppo di una medicina personalizzata su larga scala e all'automazione progressiva dei laboratori. Queste direzioni rispondono a esigenze distinte. I modelli generativi mirano ad ampliare lo spazio chimico esplorabile e ad accelerare le fasi iniziali della scoperta; la medicina personalizzata punta a adattare le terapie alle caratteristiche del singolo paziente; l'automazione dei laboratori intende ridurre il principale collo di bottiglia del processo, costituito dalla sintesi e dalla validazione sperimentale.

L'IA generativa rappresenta una delle direzioni più promettenti. A differenza dei modelli predittivi, che si limitano ad analizzare molecole esistenti, i modelli generativi progettano

nuove strutture molecolari ottimizzate rispetto a proprietà farmacologiche desiderate. Vanhaelen et al. (35) hanno descritto l'emergere della cosiddetta *generative chemistry*, un campo che utilizza il deep learning per generare entità chimiche nuove con proprietà specifiche (35). L'integrazione di reti generative avversarie, in cui due reti neurali competono per produrre molecole sempre più realistiche, e di *reinforcement learning*, che ottimizza iterativamente le molecole premiando le caratteristiche desiderate, consente di esplorare lo spazio chimico in modo guidato (35). Tale esplorazione è particolarmente rilevante perché lo spazio delle molecole sintetizzabili è stimato nell'ordine di  $10^{60}$ , una dimensione che rende impraticabile una ricerca esaustiva con i metodi tradizionali (35). Gli autori sottolineano tuttavia che molti studi rimangono confinati alla validazione *in silico*, ossia condotta unicamente per via computazionale. La sintesi e la validazione sperimentale restano il principale collo di bottiglia del processo, poiché il tempo necessario a testare in laboratorio le molecole generate dall'IA supera ampiamente quello richiesto per costruire e addestrare i modelli (35).

Un'ulteriore evoluzione è rappresentata dai modelli linguistici di grandi dimensioni e dai modelli di diffusione. Zhang et al. (29) hanno indicato l'adozione di queste tecnologie come strategia per superare il limite più evidente dell'attuale impiego dell'IA nel *drug discovery*, ossia l'assenza di farmaci progettati mediante tali metodi e approvati dalle agenzie regolatorie (29). I modelli linguistici di grandi dimensioni, noti come *Large Language Models*, sono modelli addestrati su enormi quantità di testo e capaci di integrare e individuare pattern in dati eterogenei provenienti da fonti diverse. Questa capacità li rende adatti a compiti quali la predizione dell'efficacia di un farmaco, l'identificazione di potenziali effetti collaterali e la scoperta di nuovi bersagli (29). I modelli di diffusione sono invece modelli generativi che imparano a costruire nuovi dati a partire da dati esistenti, e che hanno raggiunto lo stato dell'arte nella generazione di immagini. Applicati al *drug discovery*, trovano impiego nella progettazione di strutture proteiche e anticorpali: lo strumento *RFdiffusion* consente la progettazione *de novo* di proteine con precisione a livello atomico, mentre AlphaFold3 predice la struttura di complessi molecolari di grandi dimensioni, quali quelli tra proteine e acidi nucleici (29). Tali modelli generativi consentono di generare nuovi composti apprendendo dai dati chimici esistenti e producendo molecole con funzioni desiderate. Questa capacità risulta particolarmente rilevante per le malattie a meccanismo molecolare complesso, quali i tumori, le malattie neurodegenerative e le malattie rare, nelle

quali la difficoltà di individuare bersagli e di progettare farmaci efficaci rende prezioso un approccio in grado di esplorare in modo mirato strutture molecolari nuove (29).

La medicina personalizzata su larga scala costituisce una seconda direzione di sviluppo. Con tale espressione si intende la possibilità di estendere ad un'ampia popolazione di pazienti un approccio terapeutico adattato alle caratteristiche di ciascun individuo, superando la logica del trattamento uniforme. Moingeon et al. (30) hanno descritto l'emergere di una medicina di precisione computazionale, nella quale l'IA consente di progettare terapie e misure preventive calibrate sulla fisiologia, sulle caratteristiche della malattia e sull'esposizione a fattori di rischio ambientali del singolo paziente (30). Tale approccio si fonda sulla costruzione di modelli di malattia a partire da dati di profilazione molecolare, ottenuti confrontando i profili dei pazienti con quelli di soggetti sani, e sulla successiva stratificazione dei pazienti in sottogruppi omogenei, definiti non più sulla base dei soli sintomi clinici ma dei meccanismi molecolari sottostanti (30). Questa suddivisione in sottogruppi è il presupposto di una medicina di precisione, in quanto consente di indirizzare le terapie verso popolazioni di pazienti accuratamente definite (30).

In questo contesto, i *digital twins* rappresentano un'innovazione di rilievo. Bordukova et al. (14) hanno descritto tali repliche digitali come rappresentazioni virtuali di entità biologiche, che spaziano dalla singola cellula fino all'intero organismo, connesse alla loro controparte reale da un flusso bidirezionale di dati (14). Il *digital twin* viene inizializzato con le misurazioni di un'entità biologica reale e, applicando tecniche computazionali avanzate, simula i risultati di esperimenti reali; le sue predizioni vengono poi restituite all'entità di partenza per supportare le decisioni cliniche (14). Negli studi clinici, un *digital twin* è una rappresentazione virtuale del paziente che ne riproduce le caratteristiche longitudinali, consentendo di generare traiettorie cliniche realistiche. In questo modo è possibile simulare bracci di controllo, ridurre la necessità di arruolamento di pazienti reali e aumentare la potenza statistica attraverso un maggior numero di dati simulati, con un impatto significativo su tempi, costi ed etica della sperimentazione (14). Gli autori prospettano come evoluzione futura lo sviluppo di modelli multimodali, capaci di apprendere simultaneamente da fonti di dati di natura diversa, quali dati genetici, di imaging e di caratterizzazione dei farmaci, e di un *foundation model*, ossia un unico modello di base in grado di scalare dal livello cellulare a quello umano e di essere adattato a usi specifici (14).

L'automazione dei laboratori rappresenta la terza direzione di sviluppo, e risponde direttamente al collo di bottiglia rappresentato dalla sintesi e dalla validazione sperimentale. Vanhaelen et al. (35) prospettano l'integrazione tra IA e sistemi robotici automatizzati, capaci di assemblare molecole complesse su richiesta, affiancata dalla pianificazione della sintesi chimica basata su deep learning, che individua i percorsi di reazione necessari a produrre una molecola (35). Tale integrazione si combina con il progresso delle tecniche di simulazione computazionale, ossia di riproduzione al computer del comportamento delle molecole e delle loro interazioni, ed è considerata dagli autori un passo verso un *drug discovery* progressivamente automatizzato (35). Nel complesso, le evidenze convergono nel descrivere un settore in evoluzione verso modelli sempre più *data-driven*, automatizzati e generativi (14,29,35). Il superamento dei limiti attuali richiederà tuttavia ulteriore innovazione metodologica e una maggiore integrazione tra i diversi componenti, poiché, come osservato dagli stessi autori, le tecnologie generative non hanno ancora prodotto farmaci approvati e la loro affidabilità dipende dalla disponibilità di dati condivisi e di qualità (29,35).

### 7.3 Considerazioni etiche

L'applicazione dell'IA nello sviluppo farmaceutico e alla pratica clinica solleva considerazioni etiche rilevanti, riguardanti l'impatto sulla forza lavoro, l'equità nell'accesso alle cure e la responsabilità in caso di errore dell'algoritmo. Il timore di una sostituzione dei professionisti sanitari rappresenta uno degli ostacoli all'adozione delle tecnologie "intelligenti", ossia degli strumenti basati su IA impiegati a supporto della pratica clinica. Briganti e Le Moine (87) hanno osservato che l'opinione prevalente in letteratura ritiene che l'IA completerà l'attività dei clinici, anziché sostituirla (87). Gli autori segnalano tuttavia come la digitalizzazione precoce dei processi sanitari abbia comportato un incremento del carico amministrativo, riconosciuto come una delle principali componenti del burnout dei medici (87). L'introduzione di tali tecnologie pone quindi la necessità di un aggiornamento della formazione medica e dello sviluppo di competenze digitali, al punto che alcune scuole di medicina hanno iniziato ad affiancare al curriculum clinico percorsi di formazione ingegneristica e di alfabetizzazione digitale (87).

L'equità nell'accesso alle cure costituisce una seconda area di criticità, poiché l'impiego diffuso delle tecnologie intelligenti rischia di ampliare le disuguaglianze sanitarie anziché

ridurle. Briganti e Le Moine (87) hanno evidenziato come il monitoraggio continuo tramite dispositivi medici e le possibili violazioni della privacy abbiano il potenziale di aumentare la stigmatizzazione dei cittadini con patologie croniche o socialmente svantaggiati (87). Gli autori osservano inoltre che tale dinamica può penalizzare coloro che non sono in grado di adottare i nuovi standard di stile di vita, riducendone l'accesso all'assicurazione sanitaria e alle cure: in sistemi sanitari orientati alla responsabilità individuale del cittadino verso la propria salute, l'incapacità di conformarsi ai parametri rilevati dai dispositivi può tradursi in uno svantaggio assicurativo (87). A ciò si aggiunge la questione della proprietà e della protezione dei dati, sulla quale gli autori rilevano un orientamento crescente verso la titolarità del paziente, ossia il riconoscimento al paziente del controllo sui propri dati sanitari, soluzione che produce effetti positivi sul coinvolgimento del paziente stesso (87).

La questione dell'equità si estende ai dati impiegati per addestrare i modelli. Zhang et al. (29) hanno osservato che donne, minoranze ed anziani sono stati storicamente sottorappresentati in numerosi studi clinici (29). I modelli addestrati su tali dati possono non generalizzare correttamente queste popolazioni, poiché apprendono relazioni valide per i gruppi prevalenti nel dataset ma non per quelli sottorappresentati, con il rischio di raccomandazioni posologiche inaccurate o di effetti avversi non rilevati nelle popolazioni escluse (29). La condivisione e l'integrazione di dati provenienti da fonti diverse possono contribuire a identificare e correggere tali distorsioni, pur sollevando ulteriori preoccupazioni in termini di privacy e di gestione di informazioni sanitarie sensibili (29). Sul piano regolatorio, gli autori sottolineano la necessità di stabilire linee guida e normative per l'uso dei dati dei pazienti, garantendo che i dati siano adeguatamente anonimizzati e che ne sia prevenuto l'uso improprio (29).

Il problema della discriminazione e del *selection bias* è stato approfondito anche nel contesto del reclutamento negli studi clinici. Il *selection bias* indica una distorsione che si genera quando i soggetti inclusi non sono rappresentativi della popolazione di riferimento, mentre la discriminazione algoritmica indica un trattamento iniquo di determinati gruppi. Lu et al (89) hanno rilevato che i modelli di ML possono perpetuare i bias presenti nei dati di addestramento, riproducendo nei risultati le disparità già contenute nei dati e generando così esiti discriminatori (89). Privacy e sicurezza dei dati sono risultati i principali problemi etici associati all'uso dell'IA nel reclutamento, poiché le piattaforme richiedono spesso l'accesso a dati sensibili dei pazienti, quali le immagini mediche, che devono essere protetti e resi

accessibili unicamente al personale autorizzato (89). Il consenso informato rappresenta un ulteriore aspetto critico, in quanto i pazienti devono essere adeguatamente informati sull'utilizzo dei propri dati e devono poter scegliere di non parteciparvi (89). Gli autori sottolineano che la trasparenza nell'impiego dell'IA, con una comunicazione chiara sul funzionamento della tecnologia e sul ruolo che essa svolge nel processo decisionale, è considerata essenziale per garantire un utilizzo equo e responsabile (89).

La responsabilità in caso di errore dell'algoritmo costituisce una questione ancora irrisolta. Briganti e Le Moine (87) hanno segnalato l'assenza, a livello globale, di un quadro giuridico che definisca il concetto di responsabilità nel caso di adozione o rifiuto delle raccomandazioni di un algoritmo, condizione che espone il medico a potenziali conseguenze legali (87). Sul piano regolatorio, Bordukova et al. (14) hanno osservato che l'*FDA Modernization Act 2.0* consente l'uso di metodi computazionali come alternativa alla sperimentazione animale, senza tuttavia menzionare esplicitamente i *digital twins*, e che non esistono al momento regolamentazioni specifiche per il loro impiego negli studi clinici (14). Tale lacuna normativa deriva dal fatto che la tecnologia si è sviluppata più rapidamente del quadro regolatorio, che non ne ha ancora codificato l'utilizzo. Gli autori sottolineano la necessità di una collaborazione tra ricercatori e autorità regolatorie per definire requisiti sicuri, tecnicamente realizzabili e impattanti, e richiamano la rilevanza delle questioni etiche connesse alla fiducia che ricercatori, clinici e pazienti devono poter riporre nelle decisioni assunte sulla base di tali modelli (14). La natura "*black box*" dei modelli e la loro limitata trasparenza rendono particolarmente complesso il soddisfacimento dei requisiti di interpretabilità e affidabilità richiesti dalla valutazione della sicurezza dei farmaci (43).

#### **7.4 Il ruolo del farmacista clinico nell'era dell'IA**

L'evoluzione dell'IA ridefinisce il ruolo dei professionisti sanitari, incluso quello del farmacista clinico. Le evidenze disponibili concordano nel descrivere l'IA come uno strumento di supporto decisionale, non come un sostituto del giudizio umano (30,87). Moingeon et al. (30) hanno affermato che la responsabilità ultima delle decisioni diagnostiche o terapeutiche informate dagli algoritmi spetta ai professionisti sanitari, e che l'esperienza umana resta necessaria per fornire trasparenza sul razionale delle classificazioni e per interpretare gli output dei modelli predittivi, descritti come "*black box*" (30). In questo quadro si colloca il farmacista clinico, chiamato a un ruolo nell'interpretazione critica dei

risultati prodotti dagli algoritmi, dei quali valuta la coerenza con il contesto clinico e con le caratteristiche del singolo paziente.

La farmacovigilanza assistita da IA rappresenta un primo ambito di applicazione. Zhang et al. (43) hanno descritto come i modelli linguistici di grandi dimensioni possano supportare il monitoraggio della sicurezza dei farmaci mediante l'analisi in tempo reale delle cartelle cliniche elettroniche e dei report di eventi avversi, facilitando segnalazioni precoci e valutazioni dinamiche del rischio (43). Gli autori riportano l'esempio della piattaforma UniTox, che impiega un modello GPT-4 per estrarre e categorizzare automaticamente le informazioni di tossicità dai documenti di etichettatura di farmaci approvati dalla FDA, operazione che consente di costruire un database di tossicità standardizzato a partire da fonti testuali non strutturate (43). Tali strumenti migliorano l'efficienza dell'estrazione di dati strutturati, pur presentando limiti legati alla qualità dei dati, alla dipendenza da pattern statistici anziché da inferenza causale e alla natura poco trasparente del processo decisionale (43). La dipendenza da pattern statistici implica che il modello individua correlazioni ricorrenti nei dati senza coglierne i nessi causali, circostanza che ne limita l'affidabilità nell'interpretazione dei meccanismi di tossicità (43). In questo contesto il farmacista clinico è chiamato a integrare tali strumenti nella propria attività mantenendo un ruolo di supervisione e di validazione delle segnalazioni generate.

Il counseling personalizzato, ossia l'attività di consulenza al paziente adattata alle sue specifiche caratteristiche e necessità terapeutiche, costituisce un secondo ambito di rilievo. La medicina di precisione computazionale descritta da Moingeon et al. (30) prospetta terapie e misure preventive adattate alle necessità di ciascun paziente, con un coinvolgimento crescente dei pazienti nella gestione della propria salute (30). In questo scenario, il farmacista clinico contribuisce alla traduzione delle informazioni generate dai modelli predittivi in indicazioni comprensibili per il paziente: a partire, ad esempio, da una stima di rischio o da una raccomandazione posologica prodotta da un algoritmo, il farmacista ne spiega il significato, ne verifica la pertinenza rispetto alla terapia in corso e supporta l'aderenza al trattamento. Lu et al. (89) hanno sottolineato l'importanza di informare adeguatamente i pazienti sull'utilizzo dei loro dati e di garantire una comunicazione chiara sul ruolo che la tecnologia svolge nel processo decisionale, principio che rafforza la funzione di mediazione del farmacista nel rapporto tra strumento algoritmico e paziente (89).

Il farmacista, in virtù della propria competenza farmacologica e del rapporto diretto con il paziente, si colloca in una posizione privilegiata per accompagnare la transizione verso una pratica clinica integrata dall'IA, cioè con l'impiego delle nuove tecnologie a complemento dell'attività clinica, secondo il concetto di *augmented medicine* descritto da Briganti e Le Moine (87), nel quale gli strumenti digitali ampliano le capacità del professionista senza sostituirne il giudizio (87). In tale prospettiva, il contributo del farmacista resta orientato a mantenere al centro la sicurezza e il benessere del paziente.

## 8. CONCLUSIONI

Il presente lavoro ha analizzato, da una prospettiva farmacologica, il ruolo dell'IA nella ricerca e nello sviluppo dei farmaci, seguendone l'applicazione lungo l'intero percorso che conduce dall'identificazione di un bersaglio terapeutico fino alla sperimentazione clinica. Il punto di partenza è stato la constatazione di una difficoltà strutturale del modello tradizionale di ricerca e sviluppo. Tale modello è caratterizzato da tempi che superano il decennio, da costi che le diverse stime collocano tra il miliardo e mezzo e i due miliardi e mezzo di dollari per ciascun farmaco approvato, e, soprattutto, da un tasso di fallimento molto elevato, che porta a non raggiungere il mercato la maggioranza dei composti che entrano nella sperimentazione clinica. La progressiva diminuzione del rendimento della ricerca farmaceutica ha reso evidente la necessità di strumenti capaci di rendere le fasi iniziali del processo più razionali, rapide e predittive. È in risposta a questa esigenza che l'IA si è imposta negli ultimi anni come una tecnologia potenzialmente trasformativa.

L'analisi condotta nei capitoli precedenti ha mostrato come l'IA non costituisca un singolo strumento, ma un insieme di metodologie computazionali che intervengono in fasi distinte del processo. Il passaggio dai metodi tradizionali della chemo-informatica, basati su regole definite a priori tradizionalmente, alle tecniche di apprendimento automatico e, più recentemente, all'apprendimento profondo e all'IA generativa, ha esteso in modo considerevole la capacità di modellare le interazioni tra farmaco e bersaglio, di prevedere le proprietà farmacocinetiche e tossicologiche e di esplorare porzioni dello spazio chimico altrimenti inaccessibili per via sperimentale. Da questa analisi emerge un quadro complesso, in cui accanto a risultati concreti e promettenti permangono limiti importanti e questioni ancora aperte.

Sul piano delle potenzialità, le evidenze raccolte indicano che l'IA incide in modo significativo sulle fasi precoci dello sviluppo. La sua efficacia è massima nell'identificazione dei bersagli e nell'ottimizzazione del capostipite, fasi nelle quali consente di selezionare una molecola candidata in tempi nettamente inferiori a quelli del percorso tradizionale. Il caso dell'inibitore di TNIK ne rappresenta l'esempio più avanzato. L'impiego combinato degli strumenti di identificazione del bersaglio e di progettazione molecolare ha permesso di completare il percorso dalla scelta del bersaglio alla nomina del candidato preclinico in circa diciotto mesi. La successiva progressione della molecola, denominata rentosertib, fino al

completamento di uno studio di fase IIa con un miglioramento dose-dipendente della funzione polmonare, costituisce la prima validazione clinica di un farmaco il cui bersaglio è stato identificato e il cui composto è stato progettato mediante IA. Questo risultato rappresenta un traguardo importante, perché dimostra che l'approccio computazionale è in grado di generare molecole con buone proprietà farmacologiche e di portarle fino alla sperimentazione sull'uomo. Va tuttavia precisato che nessun farmaco progettato mediante IA ha finora ottenuto l'approvazione delle agenzie regolatorie: le molecole più avanzate, incluso il rentosertib, si trovano ancora nelle fasi della sperimentazione clinica.

Le tecnologie ancora in sviluppo lasciano intravedere ulteriori progressi nella stessa direzione, cioè una scoperta dei farmaci più rapida ed efficiente. I modelli generativi e in particolare le architetture basate sulle reti avversarie e sul reinforcement learning, ampliano lo spazio chimico esplorabile e consentono di progettare strutture molecolari ottimizzate rispetto a proprietà desiderate, superando l'approccio per tentativi ed errori che ha storicamente caratterizzato la ricerca del lead compound. L'avvento dei modelli linguistici di grandi dimensioni e dei modelli di diffusione estende questa capacità alla progettazione di strutture proteiche e anticorpali, con applicazioni che riguardano in modo particolare le malattie a meccanismo molecolare complesso, quali i tumori, le patologie neurodegenerative e le malattie rare. Accanto a queste tecnologie, la medicina di precisione computazionale prospetta la possibilità di calibrare terapie e misure preventive sulle caratteristiche del singolo paziente, attraverso la stratificazione delle popolazioni in sottogruppi definiti non più sulla base dei soli sintomi clinici, ma dei meccanismi molecolari sottostanti. I digital twins, ossia le repliche virtuali di entità biologiche, aprono inoltre la prospettiva di simulare le traiettorie cliniche dei pazienti e di ridurre la necessità di arruolamento negli studi, con un impatto potenziale sui tempi, sui costi e sulle implicazioni etiche della sperimentazione. L'automazione dei laboratori, infine, mediante l'integrazione tra IA e sistemi robotici, lascia prevedere un processo di scoperta progressivamente più rapido e meno dipendente dall'intervento manuale, nel quale la sintesi e la validazione sperimentale, oggi principale collo di bottiglia, potrebbero essere accelerate in misura sostanziale. Queste prospettive, per quanto promettenti, devono ancora essere confermate nella pratica. Molti di questi strumenti hanno infatti dato buoni risultati computazionali o in esperimenti limitati, ma resta da dimostrare che siano in grado di portare a farmaci realmente efficaci nei pazienti.

Questi sviluppi non devono però far credere che la disciplina sia già a uno stadio più avanzato di quello reale. L'analisi delle criticità ha evidenziato come i limiti principali non siano di natura computazionale, ma riguardino la qualità, la coerenza e la quantità dei dati disponibili. La scarsità di dati biologici annotati, l'ambiguità delle etichette e la natura poco trasparente dei modelli più complessi costituiscono ostacoli che non si risolvono semplicemente aumentando la capacità di calcolo. A questi si affianca il problema della validazione, poiché l'accuratezza tecnica di un modello non equivale alla sua validità clinica e numerosi studi presentano disegni metodologici non affidabili. Il dato più indicativo di questo divario è il paradosso che caratterizza attualmente il settore. A fronte di un numero crescente di molecole progettate mediante IA giunte alla sperimentazione clinica, nessuna ha ad oggi ottenuto l'approvazione delle agenzie regolatorie. L'osservazione che le molecole scoperte mediante IA mostrino nella Fase I tassi di successo superiori alla media storica per poi ridursi ai valori tradizionali nella Fase II, conferma quanto già emerso dall'analisi del modello tradizionale. Il vero collo di bottiglia dello sviluppo non risiede nella tossicità o nelle proprietà chimiche delle molecole, bensì nella capacità di individuare bersagli realmente determinanti nella patologia umana e di tradurre l'attività su un bersaglio in un beneficio terapeutico concreto. È proprio nella Fase II, dedicata alla prima valutazione dell'efficacia su un'ampia casistica di pazienti, che l'IA non ha finora dimostrato di poter modificare l'esito del processo. Il limite, in altri termini, non è tecnologico ma biologico, e dipende dalla complessità intrinseca della risposta dell'organismo umano, che nessun modello computazionale è oggi in grado di riprodurre interamente.

Il bilancio complessivo è quindi quello di una tecnologia in una fase di transizione. L'IA ha dimostrato di poter accelerare e ottimizzare le fasi iniziali della scoperta, ma non ha ancora modificato in modo dimostrabile l'esito ultimo del processo, ossia l'approvazione di nuovi farmaci capaci di rispondere a un bisogno clinico. Il suo valore attuale risiede pertanto nel ruolo di strumento di supporto, capace di potenziare le capacità del ricercatore e del clinico senza sostituirne il giudizio. La piena realizzazione delle sue potenzialità dipenderà da tre condizioni convergenti. La prima è la disponibilità di dati condivisi, abbondanti e di qualità, che costituiscono al tempo stesso il principale punto di forza e il principale limite dell'approccio. La seconda è lo sviluppo di metodi più trasparenti e rigorosamente validati, in grado di soddisfare i requisiti di interpretabilità richiesti dalla valutazione regolatoria. La terza è la definizione di un quadro normativo ed etico adeguato, che affronti le questioni della responsabilità in caso di errore algoritmico, dell'equità nell'accesso alle cure e della

protezione dei dati dei pazienti, che le principali autorità regolatorie hanno cominciato ad affrontare solo di recente.

In questo scenario in evoluzione si colloca la figura del farmacista, la cui competenza è destinata ad ampliarsi anziché a ridursi. La diffusione di strumenti basati su IA nella ricerca, nella pratica clinica e nella farmacovigilanza richiede infatti professionisti in grado di interpretarne criticamente gli output, di valutarne la coerenza con il contesto clinico e con le caratteristiche del singolo paziente e di vigilare sulla sicurezza delle decisioni assunte sulla loro base. La natura poco trasparente dei modelli più complessi rende questa funzione di interpretazione particolarmente delicata. Un sistema che fornisce una predizione senza renderne esplicito il razionale non può essere adottato in modo acritico e richiede la presenza di un professionista capace di valutare la plausibilità del risultato alla luce delle proprie conoscenze farmacologiche. Il farmacista clinico, per la sua competenza specifica sul farmaco e sui suoi meccanismi d'azione, è la figura più indicata ad esercitare questo controllo nel momento in cui l'output algoritmico riguarda la scelta, il dosaggio o la sicurezza di una terapia.

Un primo ambito nel quale tale ruolo si concretizza è quello della farmacovigilanza assistita. I modelli linguistici di grandi dimensioni consentono il monitoraggio in tempo reale delle cartelle cliniche elettroniche e delle segnalazioni di eventi avversi, facilitando l'individuazione precoce di segnali di rischio e l'estrazione automatica di informazioni di tossicità da fonti testuali non strutturate. Questi strumenti, tuttavia, presentano limiti legati alla qualità dei dati e alla dipendenza da correlazioni statistiche prive di un'effettiva comprensione dei nessi causali. Ne deriva che la sorveglianza generata in modo automatico non può sostituire il giudizio del professionista, ma deve essere sottoposta ad una validazione che ne verifichi la plausibilità prima che si traduca in una decisione clinica. In questo contesto il farmacista assume un ruolo di supervisione, ponendosi come filtro critico tra la segnalazione prodotta dall'algoritmo e il suo impiego nella pratica.

Un secondo ambito riguarda il counseling personalizzato. L'affermarsi di una medicina di precisione, nella quale le terapie vengono progressivamente calibrate sulle caratteristiche molecolari del singolo paziente, comporta un coinvolgimento crescente della persona nella gestione della propria salute. Il farmacista contribuisce alla traduzione delle informazioni generate dai modelli predittivi in indicazioni comprensibili, rendendo accessibile al paziente il significato di una stima di rischio o di una raccomandazione posologica, verificandone la

pertinenza rispetto alla terapia in corso e sostenendo l'aderenza al trattamento. Questa funzione di mediazione assume particolare rilievo in una fase nella quale la comunicazione sul ruolo della tecnologia nel processo decisionale è considerata un requisito essenziale per un impiego equo e responsabile dell'IA. Il farmacista si configura così come l'interfaccia tra lo strumento computazionale e la persona, in grado di restituire una dimensione umana e relazionale a un processo altrimenti astratto.

A queste funzioni se ne affianca una di natura più ampia, connessa alla tutela dell'equità e alla responsabilità professionale. L'impiego diffuso di strumenti basati su IA comporta il rischio di amplificare le disuguaglianze sanitarie, qualora i modelli siano addestrati su dati che sottorappresentano determinate popolazioni, con la possibilità di generare raccomandazioni inaccurate per i gruppi esclusi. Il farmacista, per la sua collocazione a contatto diretto con un'utenza eterogenea, può contribuire a riconoscere e segnalare tali distorsioni, garantendo che l'adozione delle nuove tecnologie non si traduca in un trattamento non equo. La realizzazione di questo ruolo presuppone tuttavia un adeguamento della formazione professionale. L'acquisizione di competenze digitali e di una conoscenza di base dei principi che governano gli strumenti di IA diviene una condizione necessaria perché il farmacista possa esercitare in modo consapevole la propria funzione di supervisione critica, senza né rifiutare la tecnologia a prescindere, né accettarne i risultati senza verificarli.

L'intelligenza artificiale non rende superfluo il lavoro umano, ma lo trasforma. I compiti ripetitivi possono essere affidati alla tecnologia, mentre al professionista spettano le funzioni più importanti, cioè controllare i risultati, valutarne l'attendibilità e assumersi la responsabilità delle decisioni. La tecnologia potrà accelerare la scoperta, ampliare lo spazio delle possibilità terapeutiche e rendere il processo più predittivo, ma la valutazione del beneficio reale per il paziente, il giudizio sull'appropriatezza di una terapia e la responsabilità della sua sicurezza rimangono prerogative della competenza professionale. È in questa prospettiva che la professione farmaceutica, fondata sulla conoscenza del farmaco e sulla centralità del paziente, potrà accompagnare in modo consapevole la transizione verso una pratica clinica realmente integrata dall'IA, contribuendo a garantire che l'innovazione tecnologica si traduca in un beneficio concreto e sicuro per la salute.

## BIBLIOGRAFIA

1. Paul SM, Mytelka DS, Dunwiddie CT, Persinger CC, Munos BH, Lindborg SR, et al. How to improve R&D productivity: the pharmaceutical industry's grand challenge. *Nat Rev Drug Discov.* marzo 2010;9(3):203–14. doi:10.1038/nrd3078
2. DiMasi JA, Grabowski HG, Hansen RW. Innovation in the pharmaceutical industry: New estimates of R&D costs. *J Health Econ.* maggio 2016;47:20–33. doi:10.1016/j.jhealeco.2016.01.012
3. Wouters OJ, McKee M, Luyten J. Estimated Research and Development Investment Needed to Bring a New Medicine to Market, 2009-2018. *JAMA.* 3 marzo 2020;323(9):844. doi:10.1001/jama.2020.1166
4. Wong CH, Siah KW, Lo AW. Estimation of clinical trial success rates and related parameters. *Biostatistics.* 1 aprile 2019;20(2):273–86. doi:10.1093/biostatistics/kxx069
5. Kola I, Landis J. Can the pharmaceutical industry reduce attrition rates? *Nat Rev Drug Discov.* agosto 2004;3(8):711–6. doi:10.1038/nrd1470
6. Cook D, Brown D, Alexander R, March R, Morgan P, Satterthwaite G, et al. Lessons learned from the fate of AstraZeneca's drug pipeline: a five-dimensional framework. *Nat Rev Drug Discov.* giugno 2014;13(6):419–31. doi:10.1038/nrd4309
7. Waring MJ, Arrowsmith J, Leach AR, Leeson PD, Mandrell S, Owen RM, et al. An analysis of the attrition of drug candidates from four major pharmaceutical companies. *Nat Rev Drug Discov.* luglio 2015;14(7):475–86. doi:10.1038/nrd4609
8. Scannell JW, Blanckley A, Boldon H, Warrington B. Diagnosing the decline in pharmaceutical R&D efficiency. *Nat Rev Drug Discov.* marzo 2012;11(3):191–200. doi:10.1038/nrd3681
9. Vamathevan J, Clark D, Czodrowski P, Dunham I, Ferran E, Lee G, et al. Applications of machine learning in drug discovery and development. *Nat Rev Drug Discov.* giugno 2019;18(6):463–77. doi:10.1038/s41573-019-0024-5

10. Wang H, Fu T, Du Y, Gao W, Huang K, Liu Z, et al. Scientific discovery in the age of artificial intelligence. *Nature*. 3 agosto 2023;620(7972):47–60. doi:10.1038/s41586-023-06221-2
11. Zhang K, Yang X, Wang Y, Yu Y, Huang N, Li G, et al. Artificial intelligence in drug development. *Nat Med*. gennaio 2025;31(1):45–59. doi:10.1038/s41591-024-03434-4
12. Conway A. Considerations for the Use of Artificial Intelligence To Support Regulatory Decision-Making for Drug and Biological Products.
13. Fia W. Reflection paper on the use of Artificial Intelligence (AI) in the medicinal product lifecycle\_240903.
14. Bordukova M, Makarov N, Rodriguez-Esteban R, Schmich F, Menden MP. Generative artificial intelligence empowers digital twins in drug discovery and clinical trials. *Expert Opin Drug Discov*. 2 gennaio 2024;19(1):33–42. doi:10.1080/17460441.2023.2273839 PubMed PMID: 37887266.
15. Vamathevan J, Clark D, Czodrowski P, Dunham I, Ferran E, Lee G, et al. Applications of machine learning in drug discovery and development. *Nat Rev Drug Discov*. giugno 2019;18(6):463–77. doi:10.1038/s41573-019-0024-5
16. Wang H, Fu T, Du Y, Gao W, Huang K, Liu Z, et al. Scientific discovery in the age of artificial intelligence. *Nature*. agosto 2023;620(7972):47–60. doi:10.1038/s41586-023-06221-2
17. Guo W, Liu J, Dong F, Song M, Li Z, Khan MKH, et al. Review of machine learning and deep learning models for toxicity prediction. *Exp Biol Med*. novembre 2023;248(21):1952–73. doi:10.1177/15353702231209421 PubMed PMID: 38057999; PubMed Central PMCID: PMC10798180.
18. Li Z, Jiang M, Wang S, Zhang S. Deep learning methods for molecular representation and property prediction. *Drug Discov Today*. dicembre 2022;27(12):103373. doi:10.1016/j.drudis.2022.103373

19. Zhang Y, Mastouri M, Zhang Y. Accelerating drug discovery, development, and clinical trials by artificial intelligence. *Med.* settembre 2024;5(9):1050–70. doi:10.1016/j.medj.2024.07.026
20. Vanhaelen Q, Lin YC, Zhavoronkov A. The Advent of Generative Chemistry. *ACS Med Chem Lett.* 13 agosto 2020;11(8):1496–505. doi:10.1021/acsmchemlett.0c00088
21. Tang Q, Ratnayake R, Seabra G, Jiang Z, Fang R, Cui L, et al. Morphological profiling for drug discovery in the era of deep learning. *Brief Bioinform.* 23 maggio 2024;25(4):bbae284. doi:10.1093/bib/bbae284
22. Skrede OJ, De Raedt S, Kleppe A, Hveem TS, Liestøl K, Maddison J, et al. Deep learning for prediction of colorectal cancer outcome: a discovery and validation study. *The Lancet.* febbraio 2020;395(10221):350–60. doi:10.1016/S0140-6736(19)32998-8
23. Bilal M, Raza SEA, Azam A, Graham S, Ilyas M, Cree IA, et al. Development and validation of a weakly supervised deep learning framework to predict the status of molecular pathways and key mutations in colorectal cancer from routine histology images: a retrospective study. *Lancet Digit Health.* dicembre 2021;3(12):e763–72. doi:10.1016/S2589-7500(21)00180-1
24. Rang HP. The receptor concept: pharmacology's big idea. *Br J Pharmacol.* gennaio 2006;147(S1). doi:10.1038/sj.bjp.0706457
25. Tuntland T, Ethell B, Kosaka T, Blasco F, Zang RX, Jain M, et al. Implementation of pharmacokinetic and pharmacodynamic strategies in early research phases of drug discovery and development at Novartis Institute of Biomedical Research. *Front Pharmacol.* 28 luglio 2014;5. doi:10.3389/fphar.2014.00174
26. Acosta JN, Falcone GJ, Rajpurkar P, Topol EJ. Multimodal biomedical AI. *Nat Med.* settembre 2022;28(9):1773–84. doi:10.1038/s41591-022-01981-2
27. Moingeon P, Kuenemann M, Guedj M. Artificial intelligence-enhanced drug design and development: Toward a computational precision medicine. *Drug Discov Today.* gennaio 2022;27(1):215–22. doi:10.1016/j.drudis.2021.09.006

28. Bender A, Cortes-Ciriano I. Artificial intelligence in drug discovery: what is realistic, what are illusions? Part 2: a discussion of chemical and biological data. *Drug Discov Today*. aprile 2021;26(4):1040–52. doi:10.1016/j.drudis.2020.11.037
29. Zhang Y, Mastouri M, Zhang Y. Accelerating drug discovery, development, and clinical trials by artificial intelligence. *Med*. 13 settembre 2024;5(9):1050–70. doi:10.1016/j.medj.2024.07.026 PubMed PMID: 39173629.
30. Moingeon P, Kuenemann M, Guedj M. Artificial intelligence-enhanced drug design and development: Toward a computational precision medicine. *Drug Discov Today*. gennaio 2022;27(1):215–22. doi:10.1016/j.drudis.2021.09.006 PubMed PMID: 34555509.
31. Qureshi R, Irfan M, Gondal TM, Khan S, Wu J, Hadi MU, et al. AI in drug discovery and its clinical relevance. *Heliyon*. luglio 2023;9(7):e17575. doi:10.1016/j.heliyon.2023.e17575
32. Xu Z, Ren F, Wang P, Cao J, Tan C, Ma D, et al. A generative AI-discovered TNIK inhibitor for idiopathic pulmonary fibrosis: a randomized phase 2a trial. *Nat Med*. agosto 2025;31(8):2602–10. doi:10.1038/s41591-025-03743-2
33. Miller EB, Hwang H, Shelley M, Placzek A, Rodrigues JPGLM, Suto RK, et al. Enabling structure-based drug discovery utilizing predicted models. *Cell*. febbraio 2024;187(3):521–5. doi:10.1016/j.cell.2023.12.034
34. Fang X, Liu L, Lei J, He D, Zhang S, Zhou J, et al. Geometry-enhanced molecular representation learning for property prediction. *Nat Mach Intell*. febbraio 2022;4(2):127–34. doi:10.1038/s42256-021-00438-4
35. Vanhaelen Q, Lin YC, Zhavoronkov A. The Advent of Generative Chemistry. *ACS Med Chem Lett*. 13 agosto 2020;11(8):1496–505. doi:10.1021/acsmedchemlett.0c00088
36. Gong H, Wen J, Luo R, Feng Y, Guo J, Fu H, et al. Integrated mRNA sequence optimization using deep learning. *Brief Bioinform*. 19 gennaio 2023;24(1):bbad001. doi:10.1093/bib/bbad001

37. Göller AH, Kuhnke L, Montanari F, Bonin A, Schneckener S, Ter Laak A, et al. Bayer's in silico ADMET platform: a journey of machine learning over the past two decades. *Drug Discov Today*. settembre 2020;25(9):1702–9. doi:10.1016/j.drudis.2020.07.001
38. Wu Y, Ma L, Li X, Yang J, Rao X, Hu Y, et al. The role of artificial intelligence in drug screening, drug design, and clinical trials. *Front Pharmacol*. 29 novembre 2024;15:1459954. doi:10.3389/fphar.2024.1459954 PubMed PMID: 39679365; PubMed Central PMCID: PMC11637864.
39. Li Y, Zhang L, Wang Y, Zou J, Yang R, Luo X, et al. Generative deep learning enables the discovery of a potent and selective RIPK1 inhibitor. *Nat Commun*. 12 novembre 2022;13(1):6891. doi:10.1038/s41467-022-34692-w
40. Chatterjee A, Walters R, Shafi Z, Ahmed OS, Sebek M, Gysi D, et al. Improving the generalizability of protein-ligand binding predictions with AI-Bind. *Nat Commun*. 8 aprile 2023;14(1):1989. doi:10.1038/s41467-023-37572-z
41. Tang Q, Ratnayake R, Seabra G, Jiang Z, Fang R, Cui L, et al. Morphological profiling for drug discovery in the era of deep learning. *Brief Bioinform*. 23 maggio 2024;25(4):bbae284. doi:10.1093/bib/bbae284
42. Kp Jayatunga M, Ayers M, Bruens L, Jayanth D, Meier C. How successful are AI-discovered drugs in clinical trials? A first analysis and emerging lessons. *Drug Discov Today*. giugno 2024;29(6):104009. doi:10.1016/j.drudis.2024.104009 PubMed PMID: 38692505.
43. Zhang J, Li H, Zhang Y, Huang J, Ren L, Zhang C, et al. Computational toxicology in drug discovery: applications of artificial intelligence in ADMET and toxicity prediction. *Brief Bioinform*. 31 agosto 2025;26(5):bbaf533. doi:10.1093/bib/bbaf533 PubMed PMID: 41052279; PubMed Central PMCID: PMC12499773.
44. Onakpoya IJ, Heneghan CJ, Aronson JK. Post-marketing withdrawal of 462 medicinal products because of adverse drug reactions: a systematic review of the world literature. *BMC Med*. 4 febbraio 2016;14:10. doi:10.1186/s12916-016-0553-2 PubMed PMID: 26843061; PubMed Central PMCID: PMC4740994.

45. Lexchin J. Drugs Withdrawn from the Canadian Market for Safety and Effectiveness Reasons, 1990–2024: A Cross-Sectional Study. *Drug Saf.* 1 marzo 2026;49(3):325–35. doi:10.1007/s40264-025-01612-y
46. Shah P, Kendall F, Khozin S, Goosen R, Hu J, Laramie J, et al. Artificial intelligence and machine learning in clinical development: a translational perspective. *Npj Digit Med.* 26 luglio 2019;2(1):69. doi:10.1038/s41746-019-0148-3
47. Cunningham JW, Abraham WT, Bhatt AS, Dunn J, Felker GM, Jain SS, et al. Artificial Intelligence in Cardiovascular Clinical Trials. *J Am Coll Cardiol.* 12 novembre 2024;84(20):2051–62. doi:10.1016/j.jacc.2024.08.069 PubMed PMID: 39505413; PubMed Central PMCID: PMC12178241.
48. Muti HS, Heij LR, Keller G, Kohlruss M, Langer R, Dislich B, et al. Development and validation of deep learning classifiers to detect Epstein-Barr virus and microsatellite instability status in gastric cancer: a retrospective multicentre cohort study. *Lancet Digit Health.* ottobre 2021;3(10):e654–64. doi:10.1016/S2589-7500(21)00133-3
49. Farahmand S, Fernandez AI, Ahmed FS, Rimm DL, Chuang JH, Reisenbichler E, et al. Deep learning trained on hematoxylin and eosin tumor region of Interest predicts HER2 status and trastuzumab treatment response in HER2+ breast cancer. *Mod Pathol.* gennaio 2022;35(1):44–51. doi:10.1038/s41379-021-00911-w
50. AbdelHamid SG, Halawa EM, Ibrahim EM, ElHefnawi M. Artificial intelligence-powered liquid biopsy in cancer: a paradigm shift in cancer detection and personalized care. *Cancer Cell Int.* 2 aprile 2026;26(1):160. doi:10.1186/s12935-026-04263-w
51. Shi Z, Wu B, Hu B, Zhong J, Li Z, Zhang F, et al. A large language model for clinical outcome adjudication from telephone follow-up interviews: a secondary analysis of a multicenter randomized clinical trial. *Nat Commun.* 1 dicembre 2025;17(1):211. doi:10.1038/s41467-025-66910-6
52. Dimitsaki S, Natsiavas P, Jaulent MC. Applying AI to Structured Real-World Data for Pharmacovigilance Purposes: Scoping Review. *J Med Internet Res.* 30 dicembre 2024;26:e57824. doi:10.2196/57824 PubMed PMID: 39753222; PubMed Central PMCID: PMC11729787.

53. Conway A. Considerations for the Use of Artificial Intelligence To Support Regulatory Decision-Making for Drug and Biological Products.
54. Fia W. Reflection paper on the use of Artificial Intelligence (AI) in the medicinal product lifecycle\_240903.
55. Landscape Analysis of the Application of Artificial Intelligence and Machine Learning in Regulatory Submissions for Drug Development From 2016 to 2021 [Internet]. [citato 19 giugno 2026]. Disponibile su: <https://ascpt.onlinelibrary.wiley.com/doi/epdf/10.1002/cpt.2668> doi:10.1002/cpt.2668
56. U.S. Food and Drug Administration. Using Artificial Intelligence & Machine Learning in the Development of Drug & Biological Products — Discussion Paper and Request for Feedback, maggio 2023 (revisione febbraio 2025) [Internet]. FDA; gennaio 2025 [citato 22 giugno 2026]. Disponibile su: <https://www.fda.gov/media/167973/download>
57. Research C for DE and. FDA [Internet]. FDA; 2026 [citato 19 giugno 2026]. Artificial Intelligence for Drug Development. Disponibile su: <https://www.fda.gov/about-fda/center-drug-evaluation-and-research-cder/artificial-intelligence-drug-development>
58. Research C for DE and. FDA Qualifies First AI Drug Development Tool, Will Be Used in «MASH» Clinical Trials. FDA [Internet]. 4 agosto 2026 [citato 19 giugno 2026]. Disponibile su: <https://www.fda.gov/drugs/drug-safety-and-availability/fda-qualifies-first-ai-drug-development-tool-will-be-used-mash-clinical-trials>
59. Xu Z, Ren F, Wang P, Cao J, Tan C, Ma D, et al. A generative AI-discovered TNIK inhibitor for idiopathic pulmonary fibrosis: a randomized phase 2a trial. *Nat Med.* 2025;31(8):2602–10. doi:10.1038/s41591-025-03743-2 PubMed PMID: 40461817; PubMed Central PMCID: PMC12353801.
60. Wayback Machine [Internet]. 2026 [citato 21 giugno 2026]. Disponibile su: [https://web.archive.org/web/20260226051720/https://www.ema.europa.eu/en/documents/other/seizing-opportunities-changing-medicines-landscape-european-medicines-agencies-network-strategy-2028-final\\_en.pdf](https://web.archive.org/web/20260226051720/https://www.ema.europa.eu/en/documents/other/seizing-opportunities-changing-medicines-landscape-european-medicines-agencies-network-strategy-2028-final_en.pdf)

61. Artificial Intelligence in healthcare - Public Health - European Commission [Internet]. 2026 [citato 21 giugno 2026]. Disponibile su: [https://health.ec.europa.eu/ehealth-digital-health-and-care/artificial-intelligence-healthcare\\_en](https://health.ec.europa.eu/ehealth-digital-health-and-care/artificial-intelligence-healthcare_en)
62. Agenzia Italiana del Farmaco (AIFA). Intelligenza artificiale e salute. Come l'IA sta rivoluzionando la ricerca farmaceutica, la medicina di precisione e il futuro della salute globale [dossier] [Internet]. Agenzia Italiana del Farmaco (AIFA); marzo 2026. Disponibile su: [https://www.aifa.gov.it/documents/20142/3346516/Dossier\\_stamp\\_a\\_IA\\_e\\_Salute.pdf](https://www.aifa.gov.it/documents/20142/3346516/Dossier_stamp_a_IA_e_Salute.pdf)
63. Conway A. Considerations for the Use of Artificial Intelligence To Support Regulatory Decision-Making for Drug and Biological Products.
64. Guiding Principles of Good AI Practice in Drug Development | January 2026.
65. Health C for D and R. Good Machine Learning Practice for Medical Device Development: Guiding Principles. FDA [Internet]. 19 dicembre 2025 [citato 21 giugno 2026]. Disponibile su: <https://www.fda.gov/medical-devices/software-medical-device-samd/good-machine-learning-practice-medical-device-development-guiding-principles>
66. Wong CH, Siah KW, Lo AW. Estimation of clinical trial success rates and related parameters. *Biostat Oxf Engl*. aprile 2019;20(2):273–86. doi:10.1093/biostatistics/kxx069 PubMed PMID: 29394327; PubMed Central PMCID: PMC6409418.
67. Marketing Submission Recommendations for a Predetermined Change Control Plan for Artificial Intelligence-Enabled Device Software Functions.
68. REGULATION (EU) 2016/ 679 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL - of 27 April 2016 - on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/ 46/ EC (General Data Protection Regulation).
69. European Medicines Agency. Seizing opportunities in a changing medicines landscape: the European medicines agencies network strategy 2028. [Internet]. LU:

- Publications Office; 2025 [citato 21 giugno 2026]. Disponibile su:  
<https://data.europa.eu/doi/10.2809/8994389> doi:10.2809/8994389
70. Kp Jayatunga M, Ayers M, Bruens L, Jayanth D, Meier C. How successful are AI-discovered drugs in clinical trials? A first analysis and emerging lessons. *Drug Discov Today*. giugno 2024;29(6):104009. doi:10.1016/j.drudis.2024.104009
71. Dealmakers B. Generative AI platforms drive drug discovery dealmaking. *Biopharma Deal*. 16 agosto 2024. doi:10.1038/d43747-024-00084-w
72. Ren F, Aliper A, Chen J, Zhao H, Rao S, Kuppe C, et al. A small-molecule TNIK inhibitor targets fibrosis in preclinical and clinical models. *Nat Biotechnol*. gennaio 2025;43(1):63–75. doi:10.1038/s41587-024-02143-0
73. Aliper A, Plis S, Artemov A, Ulloa A, Mamoshina P, Zhavoronkov A. Deep Learning Applications for Predicting Pharmacological Properties of Drugs and Drug Repurposing Using Transcriptomic Data. *Mol Pharm*. 5 luglio 2016;13(7):2524–30. doi:10.1021/acs.molpharmaceut.6b00248
74. Aliper A, Kudrin R, Polykovskiy D, Kamya P, Tutubalina E, Chen S, et al. Prediction of Clinical Trials Outcomes Based on Target Choice and Clinical Trial Design with Multi-Modal Artificial Intelligence. *Clin Pharmacol Ther*. 2023;114(5):972–80. doi:10.1002/cpt.3008
75. Exelixis and Insilico Medicine Enter into Exclusive Global License Agreement for ISM3091, a Potentially Best-in-Class USP1 Inhibitor.
76. menarini-com-en [Internet]. [citato 29 giugno 2026]. Menarini Group and Insilico Medicine Enter Global Exclusive License Agreement for Novel KAT6 Inhibitor for Potential Breast Cancer Treatment and Other Oncology Indications. Disponibile su: <https://www.menarini.com/en-us/news/news-detail/menarini-group-and-insilico-medicine-enter-global-exclusive-license-agreement-for-novel-kat6-inhibitor-for-potential-breast-cancer-treatment-and-other-oncology-indications.html>
77. Jeon J, Nim S, Teyra J, Datti A, Wrana JL, Sidhu SS, et al. A systematic approach to identify novel cancer drug targets using machine learning, inhibitor design and high-

- throughput screening. *Genome Med.* 30 luglio 2014;6(7):57. doi:10.1186/s13073-014-0057-7
78. Stokes JM, Yang K, Swanson K, Jin W, Cubillos-Ruiz A, Donghia NM, et al. A Deep Learning Approach to Antibiotic Discovery. *Cell.* febbraio 2020;180(4):688-702.e13. doi:10.1016/j.cell.2020.01.021
  79. Catacutan DB, Alexander J, Arnold A, Stokes JM. Machine learning in preclinical drug discovery. *Nat Chem Biol.* agosto 2024;20(8):960–73. doi:10.1038/s41589-024-01679-1
  80. Simm J, Klambauer G, Arany A, Steijaert M, Wegner JK, Gustin E, et al. Repurposing High-Throughput Image Assays Enables Biological Activity Prediction for Drug Discovery. *Cell Chem Biol.* maggio 2018;25(5):611-618.e3. doi:10.1016/j.chembiol.2018.01.015
  81. Richardson P, Griffin I, Tucker C, Smith D, Oechsle O, Phelan A, et al. Baricitinib as potential treatment for 2019-nCoV acute respiratory disease. *The Lancet.* febbraio 2020;395(10223):e30–1. doi:10.1016/S0140-6736(20)30304-4
  82. Stebbing J, Krishnan V, De Bono S, Ottaviani S, Casalini G, Richardson PJ, et al. Mechanism of baricitinib supports artificial intelligence-predicted testing in COVID-19 patients. *EMBO Mol Med.* 7 agosto 2020;12(8):EMMM202012697. doi:10.15252/emmm.202012697
  83. Kalil AC, Patterson TF, Mehta AK, Tomashek KM, Wolfe CR, Ghazaryan V, et al. Baricitinib plus Remdesivir for Hospitalized Adults with Covid-19. *N Engl J Med.* 4 marzo 2021;384(9):795–807. doi:10.1056/NEJMoa2031994
  84. Pushpakom S, Iorio F, Eyers PA, Escott KJ, Hopper S, Wells A, et al. Drug repurposing: progress, challenges and recommendations. *Nat Rev Drug Discov.* gennaio 2019;18(1):41–58. doi:10.1038/nrd.2018.168
  85. Echle A, Grabsch HI, Quirke P, Van Den Brandt PA, West NP, Hutchins GGA, et al. Clinical-Grade Detection of Microsatellite Instability in Colorectal Tumors by Deep Learning. *Gastroenterology.* ottobre 2020;159(4):1406-1416.e11. doi:10.1053/j.gastro.2020.06.021

86. Wang X, Barrera C, Bera K, Viswanathan VS, Azarianpour-Esfahani S, Koyuncu C, et al. Spatial interplay patterns of cancer nuclei and tumor-infiltrating lymphocytes (TILs) predict clinical benefit for immune checkpoint inhibitors. *Sci Adv.* 3 giugno 2022;8(22):eabn3966. doi:10.1126/sciadv.abn3966
87. Briganti G, Le Moine O. Artificial Intelligence in Medicine: Today and Tomorrow. *Front Med.* 5 febbraio 2020;7:27. doi:10.3389/fmed.2020.00027
88. Bretthauer M, Løberg M, Holme Ø, Adami HO, Kalager M. Deep learning and cancer biomarkers: recognising lead-time bias. *The Lancet.* gennaio 2021;397(10270):194. doi:10.1016/S0140-6736(20)32402-8
89. Lu X, Yang C, Liang L, Hu G, Zhong Z, Jiang Z. Artificial intelligence for optimizing recruitment and retention in clinical trials: a scoping review. *J Am Med Inform Assoc JAMIA.* 11 settembre 2024;31(11):2749–59. doi:10.1093/jamia/ocae243 PubMed PMID: 39259922; PubMed Central PMCID: PMC11491624.