

UN DIRITTO PER L'INTELLIGENZA ARTIFICIALE: *QUIS, QUID, QUANDO, UBI, CUR?*

1. Lo scenario - 2. *Quis*: chi regola e chi è regolato? - 3. *Quid*: cosa regolare? - 4. *Quando*: la regolazione tra passato e futuro - 5. *Ubi*: il perimetro applicativo del regolamento - 6. *Cur*: perché regolare l'IA? - 7. Troppe domande sono già una risposta?

Abstract

Il dibattito sulla regolazione dell'Intelligenza Artificiale (IA) è cresciuto rapidamente negli ultimi anni, spinto dai progressi tecnologici e dall'incremento dell'adozione dell'IA in svariati settori. Questo contributo si propone di esaminare le principali questioni riguardanti la regolamentazione dell'IA, utilizzando l'espedito della retorica classica dei *loci argumentorum*. Le molteplici domande sollevate dalla regolazione dell'IA generano risposte complesse e frammentarie, che rischiano di essere controproducenti per gli innovatori, per coloro che temono l'eccesso di precauzione del legislatore europeo, e persino per quanti insistono sui rischi della rivoluzione in corso.

The debate on Artificial Intelligence (AI) regulation has been intensifying in recent years, driven by technological advancements and the widespread adoption of AI across various sectors. This paper aims to examine the key issues surrounding AI regulation by employing the classical rhetorical tool of *loci argumentorum*. Several questions raised by AI regulation generate complex and fragmented responses, which risk being counterproductive for innovators, for those wary of the European legislator's excessive caution, and even for those highlighting the risks of the ongoing revolution.

Keywords: Artificial Intelligence Act (AIA), Innovation and Legislation, Precautionary Approach, Ethical AI

1. Lo scenario

Quando si riflette sull'Intelligenza Artificiale, soprattutto nell'ambito delle scienze sociali e delle discipline umanistiche, le trattazioni indulgono sulla storia di imprese scientifiche e tecnologiche che hanno determinato l'attuale stato dell'arte, dal test di Turing alla conferenza di Dartmouth¹, dal *perceptron* di Rosenblatt alla sfida a scacchi tra Deep Blue e Garry Kasparov. Sono tappe significative di una grandiosa impresa collettiva; nessuna delle quali, però, ha mai guadagnato

¹ Cfr. J. MOOR, *The Dartmouth College Artificial Intelligence Conference: The Next Fifty Years*, in *AI Magazine*, 2006, 4, pp. 87-89. Le finalità programmatiche in J. MCCARTHY, M.L. MINSKY, N. ROCHESTER, C.E. SHANNON, *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, August 31, 1955*, ora in *AI Magazine*, 2006, 4. Per il test, v. A. TURING, *Computing Machinery and Intelligence*, in *Mind*, 1950, 59, pp. 433-460.

grande fortuna nel dibattito pubblico, meno che mai sino al punto di promuovere una disciplina giuridica e – forse – una nuova etica per l'Intelligenza Artificiale. Solo nel corso degli ultimi tre anni, grazie ai progressi nel *deep learning*, nella potenza computazionale e nella disponibilità dei dati, si è assistito all'esplosione del fenomeno, soprattutto nei campi della visione artificiale e dell'elaborazione del linguaggio naturale. Le prestazioni delle GPU (*Graphics Processing Unit*), che sono fondamentali per l'addestramento dei modelli, sono cresciute di circa settemila volte dal 2003 a oggi; dal 2012 in poi la tendenza supera quella descritta dalla legge di Moore, secondo la quale il numero di *transistor* in un circuito raddoppia(va) circa ogni due anni (oggi il tempo di raddoppio, se riferito a determinati parametri quali la capacità computazionale o l'investimento in risorse, è ben più ristretto)². Si tratta, cioè, di uno di quei pochi casi in cui è corretta la spendita dell'aggettivo “esponenziale” per descrivere l'attuale curva di crescita. Quanto ai possibili usi della tecnologia, non sembrano darsi attività produttive che non si avvantaggino dell'innovazione. E, anche per tale ragione, sono emerse preoccupazioni sull'impatto della rivoluzione tecnologica per quasi ogni impiego prospettato.

È questo lo scenario sul quale si è rapidamente concentrato il dibattito pubblico e, nei confronti di un *ensemble* di dispositivi e applicazioni di recente comparsa, si sono elevate le proposte tese a promuovere un'IA etica, antropocentrica, sostenibile³, o costituzionalmente orientata⁴, sino a giungere all'*AI Act* (AIA), la prima disciplina di rilievo che si occupa dei sistemi di IA nell'Unione Europea⁵. Secondo alcuni si tratterebbe di un risultato insoddisfacente perché l'approccio regolatorio è troppo cauto e permissivo: si giunge persino ad affermare che l'IA stessa non sia etica e che – quindi – andrebbe bandita⁶. Sul fronte opposto, si reputano infondate le ragioni di un approccio così tanto precauzionale e, come accade nel recente report di Mario Draghi sulla competitività⁷, si evidenzia che l'eccesso di produzione normativa non potrà che aumentare il ritardo (già a tratti incolmabile) degli attori economici europei da quelli americani o asiatici.

² <https://openai.com/index/ai-and-compute>.

³ Una prospettiva critica in A. PIROZZOLI, *Intelligenza artificiale, sviluppo sostenibile e ambiente*, in *Consulta OnLine*, 2024, 1, pp. 111 ss.

⁴ Cfr. A. SIMONCINI, *L'algoritmo incostituzionale: intelligenza artificiale e il futuro delle libertà*, in *BioLaw Journal*, 2019, 1, pp. 63 ss. V. ampiamente in A. STERPA (a cura di), *L'ordine giuridico dell'algoritmo*, Napoli, 2024.

⁵ I.e. Regolamento (UE) 2024/1689, 13 giugno 2024.

⁶ Un esempio: D. TAFANI, *L'etica» come specchio per le allodole. Sistemi di intelligenza artificiale e violazioni dei diritti*, in *Bollettino telematico di filosofia politica*, 2023, pp. 11 ss.

⁷ M. DRAGHI, *The Future of European Competitiveness – A Competitiveness Strategy for Europe*, 9 settembre 2024, p. 26: «For example, the AI Act imposes additional regulatory requirements on general purpose AI models that exceed a pre-defined threshold of computational power – a threshold which some state-of-the-art models already

Questo contributo si propone di analizzare le questioni della regolazione dell'Intelligenza Artificiale secondo un ordine espositivo che segue un adattamento dei *loci argumentorum* della tradizione classica (se si vuole, secondo l'ordine delle *five Ws*, in cui si ritrovano alcune delle *septem circumstantiae* della retorica classica)⁸, per dimostrare quanto sia ambiguo il riferimento ai soggetti della regolazione (*quis*, § 2); alla definizione dell'IA (*quid*, § 3); al rapporto temporale tra regolazione e innovazione (*quando*, § 4); all'ambito cui si rivolgono potenzialmente le regole (*ubi*, § 5); alle ragioni della regolazione (*cur*, § 6). Si argomenterà che nessuno dei piani problematici è stato adeguatamente risolto e che il risultato complessivo dell'iniziativa, al moltiplicarsi combinatorio degli interrogativi, sia alquanto indecifrabile, se non anche nocivo: certamente per chi avversa l'approccio precauzionale del legislatore europeo e invece intenderebbe favorire l'innovazione, ma persino per quanti paventano i rischi più significativi.

2. *Quis*: chi regola e chi è regolato?

A un primo sguardo parrebbe pacifico che l'iniziativa regolatoria abbia un termine attivo e uno passivo del rapporto: da un lato le istituzioni europee che intervengono con un regolamento sui sistemi di Intelligenza Artificiale; dall'altro le società che sviluppano e immettono sul mercato le applicazioni, o quelle che le adottano nel corso di attività commerciali (AIA, art. 2, § 1, e 3, nn. 3-8). Tuttavia, l'identificazione dei controllori e dei controllati è tutt'altro che semplice.

Assumiamo che i sistemi di IA possano essere impiegati in qualsiasi settore produttivo, come peraltro affermano tanto gli "apocalittici", quanto gli "integrati"⁹. Esisterebbe un'istituzione dotata di una competenza "legislativa" così estesa da regolare in un unico atto normativo tutti i settori della società? In teoria, parrebbe, una simile competenza non sussiste né in capo alle istituzioni

exceed. [...] Digital companies are deterred from doing business across the EU via subsidiaries, as they face heterogeneous requirements, a proliferation of regulatory agencies and "gold plating" of EU legislation by national authorities» (su quest'ultimo concetto, v. C. BURELLI, *Il fenomeno del gold-plating: un ostacolo (sottostimato) all'effettività del diritto dell'Unione europea*, in *AISDUE*, 2023, IV, pp. 674 ss.). Tra l'estremo dell'approccio inibitorio e le istanze più libertarie, va sottolineata la presenza di originali posizioni compatibiliste, secondo cui «l'interazione con le intelligenze artificiali può renderci più consapevoli di come ragioniamo, più avvertiti dei nostri pregiudizi, più capaci di motivare le nostre decisioni» (A. PUNZI, *L'umanesimo digitale: verso un nuovo principio di responsabilità?*, in questa *Rivista*, 2023, 1, p. 32).

⁸ *Inter alia*, M.C. SLOAN, *Aristotle's Nicomachean Ethics as the Original Locus for the Septem Circumstantiae*, in *Classical Philology*, 2010, 3, pp. 236 ss.

⁹ La dicotomia è apparsa in U. ECO, *Apocalittici e integrati: comunicazioni di massa e teorie della cultura di massa*, Milano, 1964.

europee (cfr. art. 2 TFUE), né nel diritto interno – in Italia, anche per via del riparto Stato-Regioni dettato dall'art. 117 Cost. La proposta della Commissione¹⁰, sin dalla relazione esplicativa dell'iniziativa dell'AIA (cfr. § 2.1), individuava la base giuridica per l'intervento nella forza espansiva della clausola di armonizzazione (art. 114 TFUE)¹¹, esplicitamente richiamata nel preambolo del testo approvato, insieme con l'art. 16 TFUE, che invece assegna al Parlamento e al Consiglio europei, secondo procedura ordinaria, la competenza a stabilire le norme relative alla protezione dei dati personali (la competenza esercitata per l'adozione del GDPR, il regolamento generale sulla protezione dei dati, n. 679/2016).

Si tratta, quanto all'armonizzazione, di un meccanismo espansivo delle competenze dell'Unione, peraltro presente in altri sistemi propriamente federali (in USA, la c.d. *commerce clause*)¹², che si avvantaggia per di più del solo voto a maggioranza qualificata, e non per esempio dell'unanimità dei Paesi membri richiesta per la clausola di flessibilità delle competenze (art. 352 TFUE). Nondimeno, l'espansione che l'armonizzazione consente non è illimitata¹³, benché sia controversa l'individuazione di un limite certo e condiviso¹⁴. Nella giurisprudenza europea in materia di proibizione del tabacco, per esempio, si afferma che l'armonizzazione mira al miglioramento del funzionamento o dell'instaurazione del mercato interno, e può legittimare un'espansione *ex art. 114 TFUE*, solo se «contribuisce ad eliminare gli ostacoli alla libera circolazione delle merci e alla libera prestazione dei servizi, nonché all'eliminazione delle distorsioni della concorrenza»¹⁵.

¹⁰ COM(2021) 206 final, 21 aprile 2021.

¹¹ In generale, v. T.M. MOSCHETTA, *Il ravvicinamento delle normative nazionali per il mercato interno. Riflessioni sul sistema delle fonti alla luce dell'art. 114 TFUE*, Bari, 2018. In senso critico, cfr. S. POLI, *Il rafforzamento della sovranità tecnologica europea e il problema delle basi giuridiche*, in *AISDUE*, 2021, III, in particolare pp. 78-79; M. VARJU, *5G Networks, (Cyber)Security Harmonisation and the Internal Market: The Limits of Article 114 TFEU*, in *European Law Review*, 2020, 4, pp. 471-486 (con riferimento alla cibersicurezza ma, parrebbe, *mutatis mutandis*, alla materia in questione). In senso favorevole alla sussistenza della base giuridica: M. INGLESE, *Il regolamento sull'intelligenza artificiale come atto per il completamento e il buon funzionamento del mercato interno?*, in *AISDUE*, 2021, f.s. 2, pp. 12 ss.

¹² Tra i tanti, D.L. FRANKLIN, *Facial Challenges, Legislative Purpose, and the Commerce Clause*, in *Iowa Law Review*, 2006, 92, pp. 41 ss. pp.

¹³ Cfr. T.M. MOSCHETTA, *Il ravvicinamento*, cit., pp. 47 ss.

¹⁴ Peraltro, è difficile delineare i margini di un effettivo controllo di legittimità del diritto derivato rispetto ai trattati (*ex art. 263 TFUE*), oppure un controllo di costituzionalità rispetto al diritto interno (v. vicenda *Taricco* e il dibattito sui c.d. controlimiti): cfr. già T. GROPPI, *I diritti fondamentali in Europa e la giurisprudenza multilivello*, in E. PACIOTTI (a cura di), *Diritti fondamentali in Europa*, 2011, pp. 23-39. Di recente, v. sent. CGUE, 22 novembre 2022, *WM and Sovim SA v Luxembourg Business Registers*, C-37/20 e C-601/20.

¹⁵ Sent. CGUE, 5 ottobre 2000, *Germany v European Parliament and Council*, C-376/98. Cfr. V. DELHOMME, *Internal Market as an Excuse: The Case of EU Anti-Tobacco Legislation*, in *European Legal Studies*, 2017, p. 2; S. WEATHERILL, *The Limits of Legislative Harmonization Ten Years after Tobacco Advertising: How the*

D'altronde è l'esegesi testuale dell'art. 114 che consegna questo primo limite (§ 1), prevedendo che l'intervento si basi «su un livello di protezione elevato, tenuto conto, in particolare, degli eventuali nuovi sviluppi fondati su riscontri scientifici» (§ 2) e consentendo agli Stati di intervenire qualora si ritenga «necessario introdurre disposizioni nazionali fondate su nuove prove scientifiche inerenti alla protezione dell'ambiente o dell'ambiente di lavoro, giustificate da un problema specifico» (§§ 3-4), salvo che la Commissione non ravvisi un'ingiustificata restrizione del mercato interno (§ 6). Queste considerazioni evidenziano, innanzitutto, che non sia chiaro fino a che punto le istituzioni europee possano rivendicare una competenza così generale, ma che tale espansione debba comunque essere finalizzata alla creazione del mercato interno e non già alla sua limitazione; limitazione che potrebbe essere attivata – sulla base di prove scientifiche – da parte dei singoli Stati membri (i quali, però, ben potrebbero superare le regolazioni europee con norme più permissive e che non siano di ostacolo).

A simili conclusioni si perviene se si osservano vicende di segno contrario, come per la giurisprudenza sugli OGM, nella quale la CGUE ha affermato la validità di un'espansione delle competenze per un settore in cui prove scientifiche e necessità di uniformare il mercato legittimavano l'intervento europeo¹⁶. Dunque, si può espandere la competenza in un settore determinato e solo per eliminare gli ostacoli alla creazione del mercato, non già in un “settore” onnicomprensivo, per creare nuovi ostacoli – più o meno ragionevoli – allo sviluppo del mercato.

Ciò non implica necessariamente una (parziale?) mancanza di base giuridica dell'*AI Act*, ma ne potrebbe circoscrivere l'applicabilità; donde sorge una seconda questione, quella del margine di intervento dei Paesi membri¹⁷. In Italia è in esame al Senato il DDL di “adeguamento” alla normativa¹⁸ che potrebbe introdurre nuovi limiti ai sistemi di AI, lasciando peraltro aperto il problema del riparto interno di competenze tra Stato e Regioni, visto che anche queste ultime potrebbero individuare spazi di competenza. A tacer d'altro, lo stesso regolamento prevede spazi di

Court's Case Law has become a “Drafting Guide”, in German Law Journal, 3, 2019, pp. 827 ss., il quale evidenzia «the perils of over-hasty centralization».

¹⁶ T.M. MOSCHETTA, *Il ravvicinamento*, cit., pp. 56 ss.

¹⁷ M. EBERS *et al.*, *The European Commission's Proposal for an Artificial Intelligence Act—A Critical Assessment by Members of the Robotics and AI Law Society*, in *Multidisciplinary Scientific Journal*, 2021, 4, p. 590. M. VALTA, J. VASEL, *Kommissionsvorschlag für eine Verordnung über Künstliche Intelligenz: Mit viel Bürokratie und wenig Risiko zum KI-Standort?*, in *Zeitschrift für Rechtspolitik*, 2021, 5, pp. 142 ss.

¹⁸ Il c.d. DDL Intelligenza Artificiale annunciato in Senato nella seduta parlamentare n. 191 del 21 maggio 2024 (A.S. 1146, XIX Legislatura). Un quadro ampio, con *focus* sul contesto italiano, in G. TADDEI ELMI, S. MARCHIAFAVA, *Sviluppi recenti in tema di Intelligenza Artificiale e diritto. Una rassegna di legislazione, giurisprudenza e dottrina*, in *Rivista italiana di informatica e diritto*, 2024, 1, pp. 242 ss.

sperimentazione normativa per l'IA che gli Stati devono creare al fine di promuovere l'innovazione (art. 57 AIA), imponendo così una ulteriore normazione direzionata ad ambiti più ristretti.

A ciò si aggiunga che l'*AI Act* coinvolge nella regolazione un novero di attori molto ampio, in linea con quanto già accade con il regolamento generale sulla protezione dei dati (GDPR). Non è questa la sede per una disamina nel dettaglio, ma si pensi al ruolo di Commissione, autorità di notifica (art. 28 AIA), organismi di valutazione della conformità (art. 29), organismi notificati (art. 31), ufficio per l'IA (art. 64), consiglio per l'IA europeo per l'intelligenza artificiale (art. 65), forum consultivo (art. 67), gruppo di esperti scientifici indipendenti (art. 68), autorità nazionali, di notifica, di vigilanza del mercato (art. 70), ecc. Un quadro molto articolato, che si sovrappone alle normative già in vigore, dalla protezione dei dati al pacchetto sui servizi digitali, dalla cibersicurezza alla concorrenza, ciascuna dotata di un proprio sistema di *governance* e di relativi attori istituzionali. Non da ultimo, il sistema include nella già corposa platea dei controllori gli stessi *stakeholder*, i destinatari delle norme, chiamati a un compito autoregolatorio attraverso, per esempio, lo strumento dei codici di condotta (come già accade nel "sistema" del GDPR)¹⁹.

In tal modo, senza soluzione di continuità, solo perlustrando l'insieme dei controllori si è già pervenuti a introdurre quello dei controllati, il termine "passivo" – si fa per dire – della regolazione. Quanto all'identificazione dei destinatari delle regole, alcune questioni saranno trattate subito *infra* (§ 2) in relazione alla definizione di IA e altre in seguito, con riferimento all'ambito applicativo della disciplina (§ 4).

Se volessimo procedere con un approccio empirico, si potrebbe notare che il mercato delle IA sia dominato per oltre i tre quarti da quattro *stakeholder* principali, con altre grandi aziende che detengono ciascuna circa l'1% delle quote di mercato, e che il fenomeno delle *startup* sia ancora marginale in termini quantitativi²⁰. Sono sei, invece, le categorie di soggetti individuate dall'art. 1 del regolamento (con l'impiego di quasi centocinquanta parole). Non si tratta, nemmeno in questo caso, di una circostanza nuova o ignota: per esempio, il *Digital Markets Act* prevede che la Commissione individui i *gatekeeper* cui si rivolge la disciplina, e allo stato figurano soltanto sette

¹⁹ Artt. 40-41 GDPR e art. 95 *AI Act*.

²⁰ Le stime possono essere recuperate, con discrasie non significative quanto al nostro argomento, da più fonti. Cfr. A.A. BABANIYAZOVICH, *The Impact of Artificial Intelligence on Human Life in the Future*, in *International Multidisciplinary Journal for Research & Development*, 2023, 10, pp. 445 ss.; J.-H. SYU, J.C.-W. LIN, G. SRIVASTAVA, K. YU, *A Comprehensive Survey on Artificial Intelligence Empowered Edge Computing on Consumer Electronics*, in *IEEE Transactions on Consumer Electronics*, 2023, 4, pp. 1023 ss.

società designate²¹. È un aspetto alquanto singolare che le normative digitali – in quanto normative, generali e astratte – ambiscano a regolare attori economici ben determinati o determinabili, anche perché si tratta di settori di mercato *winner-take-all*, in cui il prodotto che viene preferito rispetto ai concorrenti, anche se solo per poco, tende a ricevere una quota enorme dei ricavi riferiti a quella stessa classe di prodotti.

A titolo ancora esemplificativo, il regolamento prevede un determinato regime per «i modelli di IA per finalità generali con rischio sistemico» e, dopo aver tentato una definizione delle condizioni da soddisfare per rientrare in questa categoria (cfr. *infra*, § 2), adotta un singolare criterio: «Si presume che un modello di IA per finalità generali abbia capacità di impatto elevato [...], quando la quantità cumulativa di calcolo utilizzata per il suo addestramento misurata in operazioni in virgola mobile è superiore a 10^{25} » (art. 51, § 2). La soglia di 10^{25} FLOPs è in parte basata su valutazioni tecniche e pratiche consolidate nel campo dell'IA, ma è anche frutto di una decisione politica, di una mediazione tra Parlamento e Consiglio, che avrebbero mirato il primo a un abbassamento e il secondo a innalzarla a 10^{26} . Il parametro, in parte tecnico e in parte arbitrario, determina (*rectius*, concorre a determinare) il perimetro dei destinatari, di modo che la norma, pur presentandosi nella sua generalità, in realtà si rivolge a un numero di destinatari molto ristretto e determinato, in cui la soglia è arbitrariamente individuata al fine di includere o escludere specifiche società.

Si tratta, detto altrimenti, di atti normativi che hanno molti aspetti in comune con un provvedimento amministrativo o una sentenza, nel senso che sono particolari (cioè, diretti a un novero determinato o determinabile di soggetti) e perciò, seppure generici nell'approccio, non generali (non rivolti alla generalità dei consociati).

3. *Quid*: cosa regolare?

Un problema non marginale consiste nella mancanza di una definizione condivisa di IA, un problema che non è risolto dall'*AI Act*. Assumiamo, per ipotesi, che i sistemi di IA siano innanzitutto quelli recentemente immessi sul mercato, con elencazione indicale delle applicazioni. Ciò consegnerebbe una definizione “estensionale”, in cui, cioè, si specifica una classe di oggetti

²¹ https://digital-markets-act.ec.europa.eu/gatekeepers_en.

enumerandone esplicitamente tutti i membri, anziché descriverne le caratteristiche o le proprietà che li definiscono, che è invece tipico delle definizioni “intensionali”, se si vuole per genere e differenza specifica, insite nelle norme giuridiche²². Si può certamente obiettare che il “dogma” della distinzione tra verità sintetiche e analitiche sia privo di un autonomo fondamento²³, cioè che non si possa distinguere chiaramente tra le prime definizioni e le seconde²⁴, oppure che le definizioni nelle argomentazioni giuridiche siano (sempre e solo) “quasi-logiche”, perché non si può supporre «la chiarezza assoluta di tutti i termini messi a confronto»²⁵, cosicché esse mescolino necessariamente aspetti intensionali con altri estensionali. Ma, in ogni caso, le definizioni presupposte a una disciplina normativa dovranno mostrare un carattere – anche solo tendenziale o regolativo – di “intensionalità”, che nel caso dell’IA, come vedremo, manca.

La proposta di regolamento della Commissione recava una prima definizione (art. 1, n. 1): «“sistema di intelligenza artificiale” (sistema di IA): un software sviluppato con una o più delle tecniche e degli approcci elencati nell’allegato I, che può, per una determinata serie di obiettivi definiti dall’uomo, generare output quali contenuti, previsioni, raccomandazioni o decisioni che influenzano gli ambienti con cui interagiscono», con elencazione estensionale delle tecniche – alquanto dubbie e forse perciò emendate²⁶ – nel citato allegato.

Quanto al testo approvato, nel cons. 12 ritroviamo una dichiarazione di intenti condivisibile: «la nozione di «sistema di IA» di cui al presente regolamento dovrebbe essere definita in maniera chiara»; «[i]noltre, la definizione dovrebbe essere basata sulle principali caratteristiche dei sistemi di IA» (dunque, una definizione intensionale). Pertanto, si introduce nella versione definitiva una nuova definizione di IA – significativamente diversa da quella della proposta originaria – la quale perde il carattere di estensionalità insito nell’elencazione delle tecniche informatiche: «“sistema di

²² Per esempio: H.L.A HART, *Il concetto di diritto* [1963], Torino, 2002, p. 18; ma si pensi all’approccio *intensionale* di H. Kelsen, *Lineamenti di dottrina pura del diritto* [1934], Torino, 2000, pp. 50 ss.

²³ Come fa W.V. QUINE, *Due dogmi dell’empirismo*, in *Id.*, *Da un punto di vista logico. Saggi logico-filosofici* [1953], Milano, 2004, pp. 35-43, adottando il lessico kantiano.

²⁴ Come vorrebbe R. CARNAP, *Meaning and Necessity: A Study in Semantics and Modal Logic*, Chicago, 1956.

²⁵ C. PERELMAN, L.OLBRECHTS-TYTECA, *Trattato dell’argomentazione. La nuova retorica* [1958], Torino, 2013, p. 228.

²⁶ I seguenti “approcci” e “tecniche”: «a) Approcci di apprendimento automatico, compresi l’apprendimento supervisionato, l’apprendimento non supervisionato e l’apprendimento per rinforzo, con utilizzo di un’ampia gamma di metodi, tra cui l’apprendimento profondo (deep learning); b)approcci basati sulla logica e approcci basati sulla conoscenza, compresi la rappresentazione della conoscenza, la programmazione induttiva (logica), le basi di conoscenze, i motori inferenziali e deduttivi, il ragionamento (simbolico) e i sistemi esperti; c)approcci statistici, stima bayesiana, metodi di ricerca e ottimizzazione» (all. 1 della proposta). Cfr. M. VEALE, F. ZUIDERVEEN BORGESIU, *Demystifying the Draft EU Artificial Intelligence Act*, in *Computer Law Review International*, 2021, 4, *passim*.

IA”: un sistema automatizzato progettato per funzionare con livelli di autonomia variabili e che può presentare adattabilità dopo la diffusione e che, per obiettivi espliciti o impliciti, deduce dall’input che riceve come generare output quali previsioni, contenuti, raccomandazioni o decisioni che possono influenzare ambienti fisici o virtuali» (art. 1, n. 1 AIA). Volendo trascurare il criterio dell’adattabilità dopo la diffusione, in quanto requisito solo eventuale (il sistema «può presentare», quindi non “deve”), gli obiettivi, impliciti o espliciti, e l’elencazione esemplificativa degli *output*, che è comunque aperta a una classe indefinita di esempi, resta una definizione molto meno chiara di quanto preannunciato nel cons. 12.

Un sistema di IA è, secondo l’*AI Act*, un sistema automatizzato «con livelli di autonomia variabili», che «deduce dall’input che riceve come generare output». Da questa definizione, epurata delle prime ridondanze, occorre sottrarre anche il riferimento al processo *input-output*, che caratterizza in genere i sistemi in quanto tali e, certamente, tutti i sistemi artificiali automatizzati. Pertanto, la definizione “essenziale” – il nucleo residuo al netto di ridondanze, tautologie e complementi esemplificativi – dovrebbe essere quella di un sistema automatizzato e autonomo, se non fosse che una simile caratterizzazione è così generica da includere tanto un termostato programmabile, quanto un *software* basato su *large language model*²⁷.

Le cose non sembrano migliorare se rivolgiamo lo sguardo alla definizione di «modello di IA per finalità generali»: «un modello di IA, anche laddove tale modello di IA sia addestrato con grandi quantità di dati utilizzando l’autosupervisione su larga scala, che sia caratterizzato una generalità significativa e sia in grado di svolgere con competenza un’ampia gamma di compiti distinti, indipendentemente dalle modalità con cui il modello è immesso sul mercato, e che può essere integrato in una varietà di sistemi o applicazioni a valle, ad eccezione dei modelli di IA utilizzati per attività di ricerca, sviluppo o prototipazione prima di essere immessi sul mercato» (art. 1, n. 63). Una definizione secondo cui, al netto delle opportune rimozioni, un modello di IA con finalità generali è caratterizzato da (i) «una generalità significativa» e (i≈ii) «sia in grado di svolgere con competenza un’ampia gamma di compiti distinti». È anch’essa una definizione caratterizzata da una generalità problematica²⁸; ma, quantomeno, ha il pregio di esprimere una distinzione tra il sistema

²⁷ Anche sul fronte “apocalittico” si registrano forti critiche alla definizione di IA, che sarebbe un concetto ideologico, privo di fondamenti scientifici: A. HOLLAND MICHEL, *Recalibrating Assumptions on AI Towards an Evidence-Based and Inclusive AI Policy Discourse*, in *Digital Society Initiative*, 2023, pp. 7-8.

²⁸ Parlano di «Uncertainty in Legal Definitions», B. GYEVNARA, N. FERGUSON, B. SCHAFER, *Bridging the Transparency Gap: What Can Explainable AI Learn from the AI Act?*, in *ECAI*, 2023, p. 968. È un pantano definitorio

autonomo – il termostato – e il sistema autonomo capace di svolgere compiti differenti, magari anche originariamente non previsti dallo sviluppatore (circostanza che potrebbe verificarsi in linea teorica anche per i termostati programmabili, nel caso in cui venissero impiegati per finalità diverse da quelle per cui sono prodotti). L'assunto di fondo, allora, sarebbe che il sistema – *rectius*, «il modello di IA per finalità generali» – sviluppi competenze ulteriori e imprevedute (forse, fuori controllo?) e che, pertanto, questo sviluppo incontrollato possa determinare rischi significativi. Detto altrimenti, è l'immagine del *Zauberlehrling* di Goethe, che attraversa la letteratura e la saggistica sino ai nostri giorni²⁹, lo sviluppatore impotente a dominare le indefinibili potenze sotterranee che lui stesso abbia evocate.

«So riconoscere un elefante quando lo vedo ma non sono capace di definirlo»³⁰? Non pare affatto questo il caso. Il fatto è che “Artificial Intelligence” è stato a lungo il nome di un programma di ricerche, proprio a partire da Dartmouth (1956), il quale è efficace non tanto per la sua precisione semantica, quanto per la sua forza nel *marketing*. E tutt'ora, volendosi rivolgere alle definizioni delle scienze dure e dell'informatica, manca una definizione condivisa, perché l'IA è un termine ombrello che si riferisce ad applicazioni volte a riprodurre singole attitudini o competenze, e ci riferisce ai sistemi artificiali “generalisti” allorché essi non si propongano di svolgere soltanto funzioni determinate, ma di affrontare compiti eterogenei, se non proprio di replicare gli *output* di qualsiasi attività intellettuale umana³¹. Peraltro, pur accogliendo in modo aporofetico la dicotomia artificiale/naturale, neppure sulla nozione di intelligenza (quella umana o degli altri animali) esiste un consenso³².

Non è detto, chiaramente, che questo termine ombrello si possa prestare a un'adozione indolore da parte di un sistema di norme propriamente giuridiche.

per C.I. GUTIERREZ *et al.*, *A Proposal for a Definition of General Purpose Artificial Intelligence Systems*, in *Digital Society*, 2023, p. 2.

²⁹ Perciò alcuni autori hanno proposto di esplicitare tale aspetto nella definizione. P.e., cfr. ancora in C.I. GUTIERREZ *et al.*, *A Proposal for a Definition of General Purpose Artificial Intelligence Systems*, in *Digital Society*, 2023, p. 5: «Un sistema di intelligenza artificiale capace di svolgere o adattarsi a svolgere una gamma di compiti distinti, inclusi alcuni per i quali non è stato intenzionalmente e specificamente addestrato» (trad. nostra).

³⁰ H.L.A HART, *Il concetto di diritto*, cit., p. 18, richiamando Agostino sul concetto di tempo.

³¹ Un quadro delle possibili distinzioni tra IA debole, forte e generale è in R. FJELLAND, *Why General Artificial Intelligence Will Not Be Realized*, in *Humanities and Social Sciences Communications*, 2020, 7. P.e., S. RUSSELL, P. NORVIG, *Intelligenza artificiale: un approccio moderno*, Milano, 2005, p. 4, propongono uno schema quadripartito delle possibili definizioni; ciò conferma che una singola definizione non c'è.

³² Si pensi soltanto alla proposta di individuare tipi differenti di intelligenza umana (c.d. teoria delle intelligenze multiple): H. GARDNER, *Intelligence Reframed: Multiple Intelligences for the 21st Century*, New York, 1999.

4. *Quando*: la regolazione tra passato e futuro

Il perimetro temporale di una normativa, se vogliamo il tempo al quale essa si rivolge, può essere ragionevolmente definito, sia pure contemplando eccezioni significative: in generale, le norme non dispongono che per l'avvenire³³ – questo è il carattere dell'irretroattività – salvo ipotesi tollerabili in cui tale principio recede. Sotto questo profilo, almeno apparentemente, la regolazione dell'IA non pone criticità significative. Il regolamento, pur stabilendo scadenze differenziate per l'efficacia delle singole parti³⁴, interviene e dispone per l'avvenire.

Tuttavia, in linea con quanto *supra* (§ 1), ci si attenderebbe che la norma disponga per l'avvenire regolando tipologie di attività che sono estratte (o astratte) dalla fenomenologia dei casi passati, “tipizzando”, cioè, in ragione dei fatti noti cui si intenda connettere conseguenze giuridiche³⁵. È in virtù di una casistica di sottrazioni di cose nell'altrui disponibilità che si può “tipizzare” il furto, e così decidere la relativa sanzione da fare conseguire. Ben più arduo sarebbe tipizzare un reato prima che un fatto accada, vale a dire delineare una classe di fatti futuri ed eventuali senza che la realtà fenomenica abbia offerto serie di eventi in cui apprezzare ricorrenze e comunanze che richiedono sanzioni o divieti.

Questo assunto proprio di ogni iniziativa normativa, e soprattutto dell'attività legislativa in senso stretto, è solo in parte compatibile con lo “spirito” e la “lettera” dell'*AI Act*. Innanzitutto, l'assunto contraddice il proposito di ogni regolazione dell'innovazione, in quanto la tipologia dei casi regolati non è indotta dai fatti passati, ma solo ipotizzata, forse, per abduzione: se tutte le norme regolano fatti del futuro, i casi che le normative sull'innovazione vorrebbero regolare sono eventi e condotte solo immaginabili e del tutto indefiniti. Non c'è nessuna tipizzazione, in senso proprio, nessuna astrazione dei tratti salienti di una casistica nota: solo ipotesi di casi futuri, al più riferite a fenomeni (più o meno) noti in Paesi tecnologicamente più avanzati. È così che si vieta il *social scoring*, la manipolazione cognitiva, la polizia predittiva, o il controllo algoritmico delle emozioni

³³ Gli echi napoleonici della formula presente nell'art. 11 delle preleggi. Cfr., da ultimo, Corte cost., sent. 4/2024 e i relativi richiami, la quale ribadisce «la centralità che assume il principio di non retroattività della legge, “inteso quale fondamentale valore di civiltà giuridica, non solo nella materia penale (art. 25 Cost.), ma anche in altri settori dell'ordinamento» (e certamente in materia di sanzioni amministrative).

³⁴ Cfr. art. 113 AIA, secondo cui il regolamento si applica a decorrere dal 2 agosto 2026, salvo tre classi di eccezioni (con scadenze 2 febbraio 2025, 2 agosto 2025 e 2 agosto 2024).

³⁵ Se si vuole, distinguendo tra fatto, fattispecie quale «classe di fatti futuri» (astratta o concreta) e caso, «ossia il fatto esistente pensato in riguardo di conformità o difformità alla fattispecie concreta» (M. ORLANDI, *Introduzione alla logica giuridica*, Bologna, 2021, pp. 18-32).

nel contesto scolastico; tutte pratiche presenti in Europa soltanto nel mondo della *science fiction*. Il legislatore, così ragionando, avrebbe potuto esprimere il suo sfavore anche per la ribellione dei robot, o vietare che essi sviluppino una coscienza e si pongano interrogativi filosofici, visto che ha scelto come punto di partenza – nel suo primo contatto con la materia – le tre leggi di Asimov³⁶. L'inesistenza o la mancata diffusione di alcune delle applicazioni che si vorrebbero vietare spiega il frequente richiamo alla letteratura futurologica, se non anche a quella fantascientifica, le quali costituiscono la principale fonte da cui estrapolare gli scenari di rischio³⁷.

Chiaramente, tale operazione prende solo le mosse da una incerta tipizzazione dei casi futuribili, ma presto abbandona il metodo tipologico per approdare a un approccio più allusivo. Si è anticipato che ciò si evince dallo “spirito” e dalla “lettera” dell'*AI Act*, e alcuni richiami testuali forniranno validi esempi. Si afferma che «[c]ome prerequisito, l'IA dovrebbe essere una tecnologia antropocentrica. Dovrebbe fungere da strumento per le persone» (cons. 6), senza che sia possibile decifrare una proiezione applicativa di questo auspicio. Nell'ambito del divieto di alcune pratiche di IA (art. 6), tra le disposizioni di maggior rilievo del regolamento, si escludono «le tecniche subliminali che agiscono senza che una persona ne sia consapevole o tecniche volutamente manipolative o ingannevoli aventi lo scopo o l'effetto di distorcere materialmente il comportamento di una persona»; l'all. III, nel definire i sistemi ad alto rischio, “individua” «i sistemi di IA destinati a essere utilizzati per la categorizzazione biometrica in base ad attributi

³⁶ Il riferimento è all'insuperata *Risoluzione del Parlamento europeo del 16 febbraio 2017 recante raccomandazioni alla Commissione concernenti norme di diritto civile sulla robotica*, (2018/C 252/25), della quale ci si limita a riportare il cons. *a*: «considerando che, dal mostro di Frankenstein ideato da Mary Shelley al mito classico di Pigmalione, passando per la storia del Golem di Praga e il robot di Karel Čapek, che ha coniato la parola, gli esseri umani hanno fantasticato sulla possibilità di costruire macchine intelligenti, spesso androidi con caratteristiche»; e *c*: «considerando che le leggi di Asimov devono essere considerate come rivolte ai progettisti, ai fabbricanti e agli utilizzatori di robot, compresi i robot con capacità di autonomia e di autoapprendimento integrate, dal momento che tali leggi non possono essere convertite in codice macchina». Cfr. A. D'ALOIA, *Il diritto verso “il mondo nuovo”. Le sfide dell'Intelligenza Artificiale*, in *BioLaw Journal*, 2019, 1, p. 6.

³⁷ Cfr. il cons. 42 AIA: «[l]e persone fisiche non dovrebbero mai essere giudicate sulla base di un comportamento previsto dall'IA basato unicamente sulla profilazione, sui tratti della personalità o su caratteristiche quali la cittadinanza, il luogo di nascita, il luogo di residenza, il numero di figli, il livello di indebitamento o il tipo di automobile». Al momento, nel campo della profilazione bancaria e creditizia, esistono sistemi che utilizzano l'IA per decidere l'accesso al credito o calcolare i tassi d'interesse, spesso basandosi su dati come il livello di indebitamento o il luogo di residenza; così come esistono sistemi assicurativi che adottano l'IA per calcolare premi basati su fattori personali, come il luogo di residenza, il tipo di automobile o lo storico degli incidenti (un ambito che è già regolamentato: per esempio, in Italia, si pensi alle norme relative all'esercizio delle attività assicurative nel d. lgs. 209/2005 e ai relativi regolamenti dell'IVASS). Nondimeno, queste applicazioni non dovrebbero essere vietate dall'*AI Act*, che sembra rivolgersi a uno scenario in cui le valutazioni sono basate *unicamente sulla profilazione* operata dall'AI. Una casistica delle applicazioni vietate la ritroviamo piuttosto nella filmografia, come in *Gattaca* (1997), *Minority Report* (2002), nell'episodio *Nosedive* della serie *Black Mirrors* (2016), oppure negli esercizi futurologici di Y.N. HARARI, *Homo Deus. Breve storia del futuro* (2015), Milano, 2017, in particolare pp. 559 ss.

o caratteristiche sensibili protetti basati sulla deduzione di tali attributi o caratteristiche» oppure «i sistemi di IA destinati a essere utilizzati per il riconoscimento delle emozioni». Si vorrebbe impedire, in questa ed altre disposizioni, una forma di controllo e di manipolazione digitale, che appartiene più alla fantascienza – almeno per ora – che alla realtà.

Si pensi al caso del *social credit score* cinese (o alla sua narrazione nella *vulgata* occidentale³⁸) spesso descritto come un meccanismo centralizzato di sorveglianza totale che assegna punteggi a ogni cittadino in tempo reale, magari esclusivamente tramite l'IA. In realtà, il sistema è composto da una serie di programmi locali e settoriali frammentati, gestiti da autorità regionali o aziende private, senza un'infrastruttura centralizzata e unificata. I dati utilizzati provengono soprattutto da fonti tradizionali, come registri finanziari e giudiziari, e l'uso dell'IA è limitato principalmente all'automazione della raccolta e gestione di grandi quantità di informazioni.

Anche se l'IA potrà acquisire un ruolo sempre più pervasivo nello sviluppo di questi sistemi – che oggi appaiono più disfunzionali che distopici³⁹ –, l'aspetto che suscita le maggiori critiche è l'attribuzione di premi e sanzioni sociali sulla base di caratteristiche o comportamenti degli individui «che non sono collegati ai contesti in cui i dati sono stati originariamente generati o raccolti» (art. 5, § 1 AIA). Va però considerato che un sistema analogo, in cui sorveglianza e valutazioni fossero affidate solo alle decisioni di persone fisiche per conto di un'autorità pubblica, non risulterebbe meno discutibile o eticamente più valido. Anzi, l'attribuzione della valutazione ad algoritmi sembra spiegare il favore di molti utenti, che si sottopongono “volontariamente” al sistema, preferendo il calcolo algoritmico a giudizi discrezionali di funzionari o *manager* delle aziende private⁴⁰. Il problema, dunque, non è l'IA, ma il trattamento discriminatorio di istituzioni pubbliche o private, che sia compiuto tramite sistemi automatizzati o decisioni umane.

In ogni caso, nei limiti delineati dall'art. 5, si potranno introdurre applicazioni che attribuiscono premi sulla base di comportamenti «collegati ai contesti in cui i dati sono stati

³⁸ Cfr. le conclusioni dell'analisi di V. BRUSSEE (*China's Social Credit Score – Untangling Myth from Reality*, in *Merics.org*, 11 febbraio 2022): «Often, the SoCS is merely invoked as a metaphor: either to depict some technological threat at home or to portray a techno-dystopian China. This is symptomatic of a tendency to see China not as a real place with real people, but as an abstract “negative opposite” of “us”». V. già J. HORSLEY, *China's Orwellian Social Credit Score Isn't Real*, in *Foreign Policy*, 16 novembre 2018: «For now, while there are many things to be worried about in China, a single and all-pervasive ranking system isn't one of them—yet». Similmente, v. S. AHMED, *The Messy Truth About Social Credit*, in *Logic(s)*, 7, 2019.

³⁹ R. CHEUNG, *The Grand Experiment. Two Decades On, China's Social Credit System Is More Dysfunctional Than Dystopian*, in *The Wire China*, 17 dicembre 2023.

⁴⁰ G. KOSTKA, *China's Social Credit Systems and Public Opinion: Explaining High Levels of Approval*, in *New Media & Society*, 2019, 7, pp. 1565-1593.

originariamente generati o raccolti» (per esempio, premiare conducenti che rispettano le regole del traffico, senza utilizzare questi dati per altre finalità, come l'accesso a servizi non correlati, l'accesso al credito, ecc.). Dunque, l'*AI Act* non riesce a imporre un divieto esplicito per applicazioni reali (come quelle che opererebbero in Cina), ma si premura di vietare uno scenario che sembra non esistere nella realtà.

A conclusioni simili si giunge se si considera il caso dei sistemi di IA che sfruttano «le vulnerabilità dovute all'età, alla disabilità o a una specifica situazione sociale o economica» per distorcere il comportamento delle persone⁴¹. “Immaginiamo”, per ipotesi, un'IA che sfrutti vulnerabilità specifiche, come uno stato emotivo fragile o una condizione economica precaria: la promozione di microcrediti ad alto interesse o di giochi d'azzardo a individui identificati come vulnerabili dal punto di vista finanziario o psicologico; i meccanismi di *loot boxes* nei videogiochi, progettati per sfruttare i comportamenti impulsivi e la propensione al rischio, specialmente nei minori, rafforzati da IA manipolative; oppure, gli algoritmi progettati per identificare vulnerabilità emotive attraverso il linguaggio o i comportamenti *online* che possono essere utilizzati per indirizzare contenuti polarizzanti o ansiogeni, aumentando l'*engagement* e aggravando stati emotivi fragili, quali la depressione o l'ansia. A ben vedere, anche volendo ammettere che l'IA presenti aspetti specifici che richiedano considerazioni distinte rispetto ad altre pratiche simili (in parte comunque lecite), il confine tra persuasione legittima e manipolazione dannosa rimane alquanto indefinito. In particolare, l'effetto cui si riferisce il legislatore europeo, quello di «distorcere materialmente il comportamento», risulta problematico per la mancata e forse impossibile distinzione tra persuasione legittima, manipolazione e alterazione del comportamento⁴². Inoltre, l'assenza di criteri oggettivi per misurare la “materialità” di una distorsione comportamentale, unita alla soggettività della percezione dell'utente, apre la strada a interpretazioni potenzialmente arbitrarie. Da un lato, qualsiasi stimolo esterno potrebbe indurre un cambiamento di

⁴¹ *I.e.* «l'immissione sul mercato, la messa in servizio o l'uso di un sistema di IA che sfrutta le vulnerabilità di una persona fisica o di uno specifico gruppo di persone, dovute all'età, alla disabilità o a una specifica situazione sociale o economica, con l'obiettivo o l'effetto di distorcere materialmente il comportamento di tale persona o di una persona che appartiene a tale gruppo in un modo che provochi o possa ragionevolmente provocare a tale persona o a un'altra persona un danno significativo» (art. 5, § I AIA).

⁴² Sulla base di queste considerazioni, per esempio, la Corte costituzionale ha espunto dall'ordinamento il reato di plagio (nel senso della manipolazione mentale) in ragione della sua non verificabilità: «L'accertamento se l'attività psichica possa essere qualificata come persuasione o suggestione con gli eventuali effetti giuridici a questa connessi, nel caso del plagio non potrà che essere del tutto incerto e affidato all'arbitrio del giudice» (Corte cost., sent. 96/1981).

comportamento⁴³; dall'altro – soprattutto per chi afferma la libertà dell'arbitrio (che qui si vorrebbe tutelare) – nessuno stimolo esterno potrebbe distorcere fino in fondo l'autodeterminazione di una persona.

Dunque, l'*AI Act* si rivolge alla generica induzione di comportamenti prodotta dall'odierna tecnologia, o a un potenziale futuro, ancora non presente sulla scena? Una possibile risposta è fornita dall'interpretazione autentica presente nei considerando. I sistemi di IA cui fa riferimento il regolamento «impiegano componenti subliminali quali stimoli audio, grafici e video che le persone non sono in grado di percepire poiché tali stimoli vanno al di là della percezione umana» (cons. 29)⁴⁴. In questo modo, il legislatore non si preoccupa tanto di vietare applicazioni già diffuse o esistenti, quanto piuttosto di anticipare un divieto per tecnologie (forse) futuribili. Si tratta, in sostanza, di vietare una distopia: un intento sul quale è facile trovare consenso (come dimostra il voto del Parlamento Europeo con 523 voti favorevoli, 46 contrari e 49 astenuti), anche perché non è affatto chiaro su quali basi si stia legiferando. Del resto, chi potrebbe dichiararsi favorevole a uno scenario distopico come quello evocato dai sostenitori più accaniti della regolamentazione⁴⁵?

Così facendo, l'iniziativa si palesa indifferente a due principi degli ordinamenti contemporanei. In primo luogo, sembra tradito il principio di legalità, proprio perché l'idea di regolare l'ignoto – non già vietando *apertis verbis* nessuna delle applicazioni oggi diffuse, ma rivolgendosi a scenari futuri – esclude che il fatto illecito sia tipizzato e sufficientemente determinato *ex ante*.

⁴³ Si pensi alla psicologia della persuasione, sin da R. B. CIALDINI, *Le armi della persuasione. Come e perché si finisce col dire di sì* [1984], Firenze, 2022. Oppure, si consideri la questione del *nudging* [R. H. THALER, C. R. SUNSTEIN, *Nudge. La spinta gentile* (2008), Milano, 2014]. La teoria – che è valsa al primo degli autori il premio Nobel per l'economia nel 2017 – si fonda sulla distinzione tra due sistemi che caratterizzerebbero il “pensiero” umano, l'uno riflessivo, l'altro automatico; altro non sono che il sistema 1, istintivo ed emozionale, e il sistema 2, deliberativo e logico, descritti poi nel noto lavoro di D. KAHNEMAN, *Pensieri lenti e veloci* (2011), Milano, 2012, a valle della collaborazione sorta già negli anni Sessanta con A. Tversky, e degli studi sulla c.d. *prospect theory* (dove il Nobel a entrambi nel 2002). L'approccio vorrebbe dimostrare – non senza critiche – che le persone non agiscono sempre in modo razionale, ma sono influenzate da *bias* cognitivi e contesti decisionali specifici, elementi che costituiscono anche il fondamento del *nudging*.

⁴⁴ «Recenti panoramiche sulle scoperte relative alla percezione subliminale hanno fornito solo prove molto deboli che degli stimoli sotto le soglie soggettive degli osservatori possano influenzarne le motivazioni, gli atteggiamenti, le credenze o le scelte. Nella maggior parte degli studi, gli stimoli non consistono in direttive, comandi o imperativi, né c'è alcuna prova affidabile che gli stimoli subliminali abbiano sulle intenzioni un qualsiasi impatto o effetto pragmatico», secondo T.E. MOORE, *Percezione subliminale: fatti e fantasie*, in *Cicap*, 28 marzo 2011.

⁴⁵ Interessante l'analisi della narrazione distopica in H. COOLS, B. VAN GORP, M. OPGENHAFFEN, *Where Exactly Between Utopia and Dystopia? A Framing Analysis of AI And Automation in US Newspapers*, in *Journalism*, 2024, 1, pp. 3 ss., forse esportabile al nostro dibattito pubblico.

È pur vero che l'ordinamento contempra numerose clausole generali del tutto atipiche, che sembrano declinare il principio di legalità in un senso più sfumato, o che quasi lo trascurano. Si pensi al caso della responsabilità aquiliana, sia nella formula tradizionale del *neminem laedere* che troviamo già in Ulpiano, sia nella clausola generale dell'art. 2043 cod. civ., in cui è atipica la classe dei fatti sanzionati con l'obbligazione risarcitoria. Intanto, la persistenza negli ordinamenti di simili clausole ha generato un dibattito ormai secolare⁴⁶, sino al punto che in sistemi giuridici diversi dal nostro si sia optato per una formula tipizzante (per esempio, nel § 823 BGB)⁴⁷. In ogni caso, la portata effettiva dell'art. 2043 cod. civ. resta piuttosto ancorata alla giurisprudenza che ne definisce progressivamente la casistica, e non si rivolge esclusivamente a classi di fatti futuri, dei quali quasi nessuno è già accaduto con ricorrenze tali da astrarne una forma tipica.

In secondo luogo, in conseguenza di questo sovvertimento dei principi di legalità e tipicità degli illeciti, è compromesso anche il principio di libertà, che vorrebbe permesso tutto ciò che non è vietato; detto altrimenti, è sovvertito il rapporto tra regola ed eccezione che si pone tra lo spazio di libertà e le sue limitazioni.

Infatti, salvo che si tratti di ricerca scientifica *non-profit* (e altre ipotesi contemplate dall'*AI Act*)⁴⁸, lo sviluppo e l'impiego di IA è permesso (solo) in quanto regolato e controllato, senza che

⁴⁶ Cfr. la ricostruzione degli ultimi decenni di dibattito sul tema della qualificazione dell'illecito extracontrattuale e della sua atipicità in G. VISINTINI, *Atipicità dei fatti illeciti e danno ingiusto*, in G. CONTE, A. FUSARO, A. SOMMA, V. ZENO-ZENCOVICH (a cura di), *Dialoghi con Guido Alpa. Un volume offerto in occasione del suo LXXI compleanno*, Roma, 2018, pp. 589 ss. V. già i lavori di S. RODOTÀ (*Il problema della responsabilità civile*, Milano, 1964) o, poi, di G. ALPA (*Il problema della atipicità dell'illecito*, Napoli, 1979) che risolsero la atipicità della disposizione (del '42, ma in altra formula già presente nelle codificazioni napoleoniche) con l'ancoramento ai valori costituzionali. Cfr., sempre sull'ingiustizia del danno e la sua atipicità in relazione ad ordinamenti stranieri, F. GALGANO, *Trattato di diritto civile*, Padova, p. 134. O, ancora, si pensi alle articolate ricostruzioni che si rinvencono nella giurisprudenza costituzionale: «L'art. 2043 cod. civ. è una sorta di “norma in bianco”: mentre nello stesso articolo è espressamente e chiaramente indicata l'obbligazione risarcitoria, che consegue al fatto doloso o colposo, non sono individuati i beni giuridici la cui lesione è vietata: l'illiceità oggettiva del fatto, che condiziona il sorgere dell'obbligazione risarcitoria, viene indicata unicamente attraverso la “ingiustizia” del danno prodotto dall'illecito. È stato affermato, quasi all'inizio di questo secolo (l'osservazione era riferita all'art. 1151 dell'abrogato codice civile ma vale, ovviamente, anche per il vigente art. 2043 cod. civ.) che l'articolo in esame “contiene una norma giuridica secondaria, la cui applicazione suppone l'esistenza d'una norma giuridica primaria, perché non fa che statuire le conseguenze dell'iniuria, dell'atto contra ius, cioè della violazione della norma di diritto obiettivo”» (Corte cost., sent. 184/1986).

⁴⁷ Che prevede espressamente i beni e i diritti protetti: «Wer vorsätzlich oder fahrlässig das Leben, den Körper, die Gesundheit, die Freiheit, das Eigentum oder ein sonstiges Recht eines anderen widerrechtlich verletzt, ist dem anderen zum Ersatz des daraus entstehenden Schadens verpflichtet». Un'utile analisi comparata dalla prospettiva tedesca in J. METZING, *Gefährdungshaftung: Einzelgesetze oder Generalklausel?*, in *Bucerius Law Journal*, 2, 2014, la quale pur pronunciandosi in favore di una ammissibilità della clausola generale, riprende le questioni relative alla deformazione dei rapporti tra potere legislativo e potere giudiziario nel caso delle norme in bianco.

⁴⁸ Semplifico l'eccezione dell'art. 2, § 6: «Il presente regolamento non si applica ai sistemi di IA o modelli di IA, ivi compresi i loro output, specificamente sviluppati e messi in servizio al solo scopo di ricerca e sviluppo scientifici»; l'eccezione più significativa, in quanto in ossequio alla libertà di ricerca. Le altre limitazioni dell'art. 2 si riferiscono

sia possibile predeterminare quale classe di fatti sia attratta nel perimetro della norma, prima che il decisore del caso lo affermi. Per altro verso, la ragione per cui il legislatore europeo intende che sia meglio vietare che permettere è da ricondurre senz'altro al principio di precauzione, che costituisce ormai da tempo una strategia retorica persuasiva e pervasiva, ma spesso indecifrabile⁴⁹.

Ma, oltre alle ridondanze e ai sovvertimenti di alcuni dei principi chiave dell'ordinamento, la disciplina si assegna un ruolo sorprendente (soprattutto, per gli sviluppatori dei *software*): «La normazione dovrebbe svolgere un ruolo fondamentale nel fornire soluzioni tecniche» (cons. 121). Non si tratta del *nudging* o della funzione promozionale delle norme: il diritto stesso potrà offrire soluzioni tecniche all'informatica.

Per queste ragioni, la normativa prospetta un rapporto col tempo del tutto diverso dagli approcci regolatori consueti, e si distacca da un modello fondato sui principi di libertà, di legalità, sulla c.d. calcolabilità del diritto⁵⁰; principi forse più regolativi che effettivi, ma raramente sovvertiti *apertis verbis* dal legislatore.

Ancora con riferimento alla manipolazione cognitiva, al di là dell'indefinita nozione dei sistemi di IA, le questioni irrisolte si moltiplicano: restano oscuri i riferimenti alle «tecniche subliminali» o alle «tecniche volutamente manipolative o ingannevoli», allo «scopo o l'effetto di distorcere materialmente il comportamento», al pregiudizio considerevole della «capacità di prendere una decisione informata», all'induzione a prendere decisioni che altrimenti non si sarebbero prese, alla capacità del sistema di poter «ragionevolmente provocare» dei danni, alla significatività che dovrebbe caratterizzare tali danni, al concetto di sfruttamento delle vulnerabilità, o alla stessa nozione di vulnerabilità.

È forse vero che nel diritto non si possa aspirare alla stessa precisione definitoria propria delle scienze dure, e che le passioni per le definizioni debbano essere temperate dal buon senso e da un approccio pragmatico. Tuttavia, ciò può mai giustificare un abbandono a un'irriflessiva vaghezza?

perlopiù ad attività di contrasto, a scopi militari, di difesa o di sicurezza nazionale; cioè, non si tratta di spazi di libertà riconosciuti ai privati, ma spazi di manovra autoritativa che il legislatore europeo riserva agli Stati membri. In ogni caso, l'eccezione, visti gli attuali costi necessari per lo sviluppo di modelli con finalità generali, sembra abbastanza ingenua, oltre a lasciare intendere un presupposto dubbio (i ricercatori sono intimamente buoni; i mercanti probabilmente cattivi).

⁴⁹ In senso critico, C.R. SUNSTEIN, *The Paralyzing Principle*, in *Regulation*, 2002, 25, pp. 32 ss.; ID., *Il diritto della paura. Oltre il principio di precauzione* [2005], Bologna, 2010.

⁵⁰ La nozione di diritto *calcolabile*, in quanto razionale e formale, e come tale attribuita a Weber, è in N. IRTI, *Un diritto incalcolabile*, Torino, 2016, pp. 3 ss. Cfr. anche i diversi contributi in A. CARLEO, *Calcolabilità giuridica*, Bologna, 2017.

4. *Ubi*: il perimetro applicativo del regolamento

L'ambito di applicazione di una disciplina che si rivolge alle tecnologie digitali è un tema complesso, innanzitutto per via della “immaterialità” e della diffusione transnazionale dei *software*.

Tra i precedenti più emblematici si consideri proprio il GDPR, che stabilisce due criteri relativi all'ambito di applicazione, l'*establishment criterion* e il *targeting criterion*. Secondo il primo (art. 3, § 1 GDPR), il regolamento si applica «al trattamento dei dati personali effettuato nell'ambito delle attività di uno stabilimento», mentre il secondo prevede l'applicazione «al trattamento dei dati personali di interessati che si trovano nell'Unione, effettuato da un titolare del trattamento o da un responsabile del trattamento che non è stabilito nell'Unione» (art. 3, § 2), a patto che offra beni e servizi o monitori il comportamento di una persona nell'UE.

In breve, il GDPR si applica a chiunque compia attività di trattamento dati in uno stabilimento in Europa, o anche nel caso in cui – nell'ambito dell'attività dello stabilimento – i dati vengano trattati fuori dell'Unione (un'azienda stabilita in Italia, che delocalizzi determinate attività di trattamento dei dati fuori dei confini europei, indipendentemente dal fatto che i dati personali trattati siano di cittadini europei). E, in aggiunta, si applica anche nel caso in cui un'azienda straniera tratti i dati, offrendo servizi o con forme di monitoraggio, di un utente che si trovi in Europa (cittadino o meno). Si tratta di una semplificazione, perché si possono immaginare situazioni limite: nel caso di un lavoratore americano che si serva in Europa di un'*app* destinata al mercato americano, il regolamento non troverà applicazione (che sia o meno cittadino europeo); e neppure si applicherà al caso di un'azienda americana che tratti i dati di un dipendente che si trovi in Italia per un viaggio di lavoro; mentre troverà applicazione per una università americana che rivolga ai cittadini europei un'offerta formativa, pubblicizzando la piattaforma *online* in Europa, magari con pubblicità nelle lingue nazionali⁵¹.

Allora, non tutti i trattamenti dei dati di utenti in Europa rientreranno nel *territorial scope* del GDPR, e neppure tutti i trattamenti relativi ai cittadini europei, mentre alcuni trattamenti avvenuti “fuori” del territorio dell'Unione, persino se riferiti a cittadini non europei, saranno attratti dalla regolazione. Si configura, cioè, un “territorio” mobile e articolato, con ogni intuibile difficoltà

⁵¹ I casi sono tratti, con adattamenti, da EDPB, *Guidelines 3/2018 on the Territorial Scope of the GDPR (Article 3)*, vers. 2.1, 12 novembre 2019.

in ordine all'effettività delle regole; un'effettività affidata al più alla minaccia di sanzioni economiche e, nel caso limite, ai *ban* delle piattaforme.

Tutto ciò vale *a fortiori* per l'*AI Act*, nel quale l'art. 2 replica la logica dello *establishment criterion*, rivolgendosi a *providers*, *deployers*, importatori e distributori che siano stabiliti o locati nell'UE (o che propongano un sistema AI per il mercato europeo) e il criterio *ratione personae* («alle persone interessate che si trovano nell'Unione», art. 2, § 1, lett. *g*), estendendo per entrambi i criteri lo *scope* del GDPR⁵², con ogni conseguenza in relazione all'*enforcement* della disciplina.

5. *Cur*: perché regolare l'IA?

La domanda sulle ragioni della regolazione dell'IA può intendersi in modi differenti e ambire a risposte di ordine eterogeneo. Una cosa è chiedersi le ragioni in termini di politica del diritto⁵³, altro è riflettere sull'intenzione storica del legislatore, altro sulla c.d. *ratio legis* della disciplina⁵⁴, ecc. Tali piani sono distinguibili in modo analitico per via di approcci teorico-generalisti piuttosto condivisi.

Si può, in questa sede, far riferimento al fine in vista del quale le forze politiche e culturali che hanno ispirato e sostenuto l'iniziativa (diremmo, in sede di politica del diritto, in una prospettiva storica); si può poi far riferimento all'intenzione del legislatore, ripercorrendo i lavori preparatori e gli stessi *recital* che precedono l'articolato del regolamento; o far riferimento alla «volontà della legge», alla ricerca della c.d. *ratio legis* (sono questioni che si pongono sul piano dell'interpretazione giuridica propriamente intesa).

È del tutto evidente che ciascuno di questi piani possa essere frequentato alla luce delle opportune prudenze metodologiche: perché non è possibile, o quantomeno è discutibile, attribuire «intenzioni ad entità collettive» (il che precluderebbe la possibilità di ricostruire l'intenzione del legislatore); o perché la *ratio legis* è a sua volta il risultato di un'attività ermeneutica volta a ricostruire «lo scopo, il risultato razionale che la norma può oggettivamente perseguire»⁵⁵ (un'attività pur sempre soggettiva); o perché la stessa ricostruzione delle finalità politico-giuridiche

⁵² Cfr. le limitazioni di cui all'art. 3, § 2 GDPR, assenti nella lettera dell'art. art. 2, § 1, lett. *g*), AIA.

⁵³ O di sociologia del diritto come in H. Kelsen, *Lineamenti di dottrina pura del diritto*, cit., pp. 53-55.

⁵⁴ Si possono adottare le distinzioni di R. GUASTINI, *Interpretare e argomentare*, Milano, 2011, pp. 271-276.

⁵⁵ R. GUASTINI, *Interpretare e argomentare*, cit., p. 274.

dell'iniziativa si accompagna all'ingenua presunzione che forze politiche e culturali differenti abbiano agito per ragioni comuni (presunzione storicamente inconsistente).

Inoltre, in sede di interpretazione, i singoli piani del ragionamento possono produrre argomenti autonomi e finalizzati a sostenere una tesi specifica o argomenti ausiliari al sostegno «di una più ampia strategia argomentativa»⁵⁶; e la stessa ricostruzione delle finalità politico-giuridiche può essere tanto finalizzata a evidenziare i limiti dell'iniziativa, quanto a rimarcare i pregi. Senza trascurare che i singoli piani di analisi si pongono frequentemente in conflitto perché possono rispondere ad obiettivi antitetici e sostenere strategie argomentative opposte (*ratio legis* contro intenzione del legislatore, politica del diritto legislativa contro politica del diritto dell'interprete, ecc.).

Nessuna di queste prudenze, però, ci impedisce di notare gli elementi comuni ai risultati possibili delle diverse analisi, i quali risultati dimostrano contraddizioni sia quando convergono, che quando divergono.

Sembra impossibile una ricerca della cifra unitiva delle finalità politico-ideologiche: si intende perseguire uno sviluppo etico dell'IA, progettando sistemi che riflettano i valori umani⁵⁷; si mira a mitigare le implicazioni sociali e le disuguaglianze, anche fronteggiando il rischio dei *bias*⁵⁸; a proteggere gli individui dal capitalismo della sorveglianza⁵⁹; al contempo, si vorrebbe istituire un mercato capitalistico concorrenziale affinché si realizzi lo sviluppo dei sistemi di IA⁶⁰; o, ancora, garantire democraticità, trasparenza ed accessibilità⁶¹; garantire i consumatori; stimolare l'innovazione; assicurare la supervisione umana degli algoritmi; garantire la sostenibilità ambientale; assicurare un quadro di regole certe; garantire regole flessibili; ecc. Volendo prescindere dagli esiti, è raro che una proposta legislativa possa contestualmente promuovere finalità tanto diverse.

⁵⁶ *Ibid.*, p. 272.

⁵⁷ Le ragioni riassunte recentemente in L. FLORIDI, *The Ethics of Artificial Intelligence: Principles, Challenges, and Opportunities*, Oxford, 2023. Oppure, cfr. V. DIGNUM, *Responsible Artificial Intelligence. How to Develop and Use AI in a Responsible Way*, Cham, 2019.

⁵⁸ K. CRAWFORD, *Power, Politics, and the Planetary Costs of Artificial Intelligence*, New Haven, 2021, in particolare pp. 53 ss.; T. GEBRU, *Race and Gender*, in M.D. DUBBER, F. PASQUALE, S. DAS (a cura di), *The Oxford Handbook of Ethics of AI*, Oxford, 2020, pp. 251-269.

⁵⁹ S. ZUBOFF, *Il capitalismo della sorveglianza* [2019], Roma, 2019, in particolare pp. 369 ss.

⁶⁰ Argomenti recentemente ribaditi in H. VÖPEL, *The AI Revolution: A New Paradigm of Economic Order*, in *The Economists' Voice*, 2024.

⁶¹ Cfr. J. LAUX, *Institutionalised Distrust and Human Oversight of Artificial Intelligence: Towards a Democratic Design of AI Governance Under the European Union AI Act*, in *AI & Society*, 2023, pp. 1-14.

Invece, proprio l'analisi dei lavori preparatori conferma una fiducia da parte della Commissione nel raggiungimento di una sintesi possibile⁶². Una fiducia ribadita dalle altre istituzioni interpellate, le quali – nel salutare con favore l'iniziativa – hanno espresso perplessità (proprio e soltanto) per i relativi ambiti di competenza: la BCE per quel che concerne i limiti all'IA nel settore creditizio⁶³; il Comitato economico e sociale europeo per la troppa cautela nella protezione dei diritti⁶⁴; il Comitato europeo delle regioni esprimendo perplessità sugli oneri burocratici e sui limiti posti all'IA nel contesto degli enti locali⁶⁵; il Parlamento europeo con proposte di emendamento eterogenee, in tema di diritti, sostenibilità, cibersicurezza, qualità, *privacy*, ecc.⁶⁶.

Sulla base di queste premesse, non potrà stupire che i tentativi di (ri)costruire la *ratio legis* delle norme ricavabili saranno alquanto erratici: perché andrà combinato l'approccio basato sul rischio (*bottom-up*) con il divieto esplicito di taluni impieghi (dove il rischio è predeterminato, in una logica *top-down*)⁶⁷; bisognerà contemperare i divieti con le norme promozionali; i riferimenti all'approccio antropocentrico della tecnologia e le sue svalutazioni nell'adesione a una sensibilità ambientalista (anti-antropocentrica?); l'istituzione di un mercato interno, dove sia anche promossa la concorrenza, e i diritti degli individui (ora fondamentali, ora declinati quali diritti del consumatore), la democrazia, se non anche l'eguaglianza sostanziale.

6. Troppe domande sono già una risposta?

Abbiamo tentato di illustrare che la regolazione dell'IA si fondi su un quadro teorico ambiguo e si connoti per criticità tutt'altro che trascurabili.

In primo luogo, si è osservato come sia del tutto opaca l'identificazione delle competenze normative su un ambito così generale (se non generico); un aspetto critico che non è risolto, ma al più aggravato dalla moltiplicazione dei soggetti coinvolti nella produzione delle regole e, da ultimo,

⁶² Cfr. § 1 dell'*Explanatory Memorandum* della proposta COM(2021) 206 *final*; v. già Commissione europea, *Libro bianco sull'intelligenza artificiale - Un approccio europeo all'eccellenza e alla fiducia*, COM(2020) 65 *final*.

⁶³ Parere 2022/C 115/05, 11 marzo 2021.

⁶⁴ Parere 2021/C 517/09, 22 dicembre 2021.

⁶⁵ Parere 2022/C 97/12, 28 febbraio 2021.

⁶⁶ C/2024/506, 23 gennaio 2024.

⁶⁷ Cfr. l'analisi (compatibilista) di C. NOVELLI, *L'Artificial Intelligence Act Europeo: alcune questioni di implementazione*, in *federalismi.it*, 2024, 2, pp. 102-105.

dall'inclusione dei destinatari stessi nella normazione, secondo un approccio (anche) autoregolatorio.

In secondo luogo, abbiamo evidenziato che il quadro definitorio dell'oggetto stesso della regolazione sia del tutto carente e che ciò determini grande incertezza in relazione all'ambito di applicazione, oltre che il rischio dell'arbitrio dei controllori.

Poi, in terzo luogo, abbiamo esaminato le criticità in termini di legalità e prevedibilità, se non anche di tipicità e di determinatezza delle fattispecie illecite, in relazione a un'iniziativa che non è volta a regolare fatti e condotte prevedibili in quanto già verificatisi in passato, ma che mira piuttosto ad afferrare rischi ed abusi di tecnologie che ancora non esistono e che forse neppure esisteranno.

Sul piano dell'ambito applicativo delle norme, nel quarto paragrafo, abbiamo proposto una ricostruzione critica del combinarsi dei criteri che consentono al regolamento di rivolgersi a condotte realizzate al di fuori dell'UE, con ogni contraddizione insita nel proposito di regolare in Europa quello che accade altrove (e nella speranza che l'altrove si adegui alle regole europee, o che ne sia ispirato).

Da ultimo, nel quinto paragrafo, abbiamo tentato di ricostruire le "ragioni" dell'intervento regolatorio, sia sul piano della politica del diritto che sul fronte dell'interpretazione delle disposizioni, concludendo che esse, allorché presentino elementi comuni, appaiono fortemente autocontraddittorie, così come appaiono in contraddizione gli assunti che le sorreggono.

L'*AI Act*, insieme con l'intero *corpus* europeo di normative digitali, è certamente un valido esempio di postdiritto⁶⁸: nondimeno, è dubbio che all'indebolimento delle categorie e degli assoluti della modernità debba conseguire necessariamente una rincorsa quasi intenzionale al disordine e all'arbitrio. Ciò non risponderebbe né al proposito di ricondurre l'ambiente digitale a una nuova forma di legalità (potremmo chiamare "legalità" lo scenario descritto?), né di promuovere l'innovazione e istituire un mercato europeo. Detto altrimenti, non è una soluzione né per gli

⁶⁸ Cioè, distante da qualsiasi modello idealtipico di *rule of law* (à la Fuller, Dicey, ecc.) o di *Rechtsstaat*: non si individuano fonti e competenze (né gerarchia delle fonti, né norme di riconoscimento, ecc.); le norme non sono generali ed astratte; oppure, sono indeterminate in modo tale che la norma realmente applicabile al caso è solo conoscibile *ex post*; sono contraddittorie; "impongono" prestazioni di dubbia fattibilità (se non impossibili); sono mutevoli; richiedono il complemento di *soft law* (e di atti via via sempre meno ascrivibili al *genus* delle fonti). Cfr., per tutti, A. SCHIAVELLO, *Postdiritto: una breve guida per i perplessi*, in *Rivista di filosofia del diritto*, 2023, 2, pp. 261 ss.; v. anche G. ZACCARIA, *Postdiritto. Nuove fonti, nuove categorie*, Bologna, 2022, pp. 12 ss. Alcune osservazioni critiche sull'iniziativa e le sue contraddizioni in V. ZENO-ZENCOVICH, *Artificial Intelligence, Natural Stupidity and Other Legal Idiocias*, in *MediaLaws*, 2024, 1, p. 12.

“apocalittici”, né per gli “integrati”, e neppure un compromesso di sintesi tra le istanze di entrambi⁶⁹.

Sia chiaro, nessuna di queste conclusioni provvisorie ci preclude di ritenere che il regolamento farà il suo corso: l'*AI Act* è già applicabile e, si vedrà come e con quali effetti, sarà applicato. Soltanto il tempo potrà dirci quanto fondata fosse la scommessa delle politiche europee: si avrà il c.d. effetto Bruxelles, che vedrebbe l'Europa ergersi a faro della protezione dei diritti, in modo da orientare globalmente la tecnologia verso i “suoi” valori⁷⁰?

Prevedere gli esiti di una scommessa così ambiziosa non rientra tra gli obiettivi possibili dell'analisi. Tuttavia, rileviamo indizi rilevanti della presenza di ostacoli che potranno allontanare o impedire la realizzazione del programma. Siamo davvero convinti che le realtà dell'innovazione in atto siano alla ricerca di una guida etico-giuridica di matrice europea? Siamo in grado di definire quali siano i “nostri” valori e di assicurarci che non ci sia un conflitto proprio dentro le nostre mura? Quanto forte è il rischio che il più vecchio⁷¹ dei continenti, più che un faro per il progresso, diventi soprattutto un banco di sabbie burocratiche?

FRANCESCO CIRILLO
Università degli Studi della Tuscia

⁶⁹ E, in ogni caso, si converrà sul punto che sia «superfluo speculare su come arginare il progresso» (L. CORSO, *Intelligenza collettiva, intelligenza artificiale e principio democratico*, in AA. VV., *Il diritto nell'era digitale. Persona, Mercato, Amministrazione, Giustizia*, Milano, 2022, p. 459).

⁷⁰ L'immagine olistica dei valori europei imporrebbe un interrogativo: quali valori? «Nulla garantisce che questi valori indichino, concordemente, un'unica direzione alle nostre scelte» (B. CELANO, *Diritti, principi e valori nello Stato costituzionale di diritto: tre ipotesi di ricostruzione*, in *Diritto & Questioni pubbliche*, 2004, 4, p. 15).

⁷¹ L'attribuzione è riferita alla demografia e non alla storia.