

Wstęp

Pomimo że początek AI, najczęściej jest identyfikowany z tzw. testem Turinga¹, czyli eksperymentem, którego wyniki opublikowano w 1950 r. i będącym *de facto* próbą ustalenia zasad działania AI, a pojęcia AI użyto po raz pierwszy na świecie w 1956 r.², to dopiero w 2018 r. podjęto udaną próbę wprowadzenia pierwszej legalnej definicji AI w amerykańskim porządku prawnym, tj. w § 238 lit. (g) ustawy o autoryzacji wydatków na obronę narodową na rok fiskalny 2019 (NDAA 2019)³. Trudno mówić tu o opieszałości amerykańskiego ustawodawcy i któregośkolwiek innego na świecie, ponieważ trzeba wziąć pod uwagę fakt, że mamy do czynienia z jednym z najbardziej złożonych zagadnień w historii ludzkości. Za przyjęciem tak postawionej tezy przemawia kilka argumentów.

Po pierwsze, J. McCarthy będący twórcą terminu AI uznał, że jest to nauka i inżynieria tworzenia inteligentnych maszyn zdolnych do wykorzystywania niektórych cech ludzkiego umysłu⁴. Inaczej mówiąc, AI ma w pewnym stopniu odzwierciedlać ogół aktywności najbardziej złożonego ludzkiego organu jakim jest mózg, takich jak rozumienie języka, rozpoznawanie obrazów, rozwiązywanie problemów, podejmowanie decyzji i uczenie się. Ponieważ budowa i funkcje mózgu są nadal niezupełnie rozpoznane, to badacze AI wciąż nie są zgodni w przedmiocie wyznaczenia granicy pola badawczego⁵. Po drugie, co

¹ A. Turing, Computing Machinery and Intelligence, *Mind* 1950, Nr LIX (236), s. 433–460, doi:10.1093/mind/LIX.236.433, dostęp: 8.2.2024 r.

² J. McCarthy, M.L. Minsky, N. Rochester, C.E. Shannon, A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence *AI Magazine* 1956, Nr 27(4), s. 12, doi.org/10.1609/aimag.v27i4.1904, dostęp: 8.2.2024 r.

³ *National Defense Authorization Act for Fiscal Year 2019*, 115th U.S. Congress Pub. L Tooltip, Public Law 2018, s. 115–232 [132 Stat. 1636 through 132 Stat. 2423].

⁴ J. McCarthy, M.L. Minsky, N. Rochester, C.E. Shannon, A Proposal for the Dartmouth Summer, s. 11.

⁵ Zob szerzej R. Penrose, *Makroświat, mikroświat i ludzki umysł*, Cambridge University Press 1997; R. Penrose, *The Emperor's New Mind: Concerning Computers, Minds, and the Laws of Physics*, Oxford University Press 2002; R. Penrose, *Shadows of the Mind: A Search for the Missing Science of Consciousness*, Vintage 1995. W tym miejscu zasadne jest przywołanie prac R. Penro-

prawda AI jest przedmiotem badań od niedawna, to dynamiczny proces ewolucji i rozwoju determinowany postępowaniem technologicznym nieustannie potęguje jej niepewność i nieprzewidywalność⁶. Po trzecie, AI charakteryzuje się unikalną interdyscyplinarnie tożsamością, a badania nad nią lub z jej udziałem są właściwie prowadzone we wszystkich dziedzinach naukowych⁷. Należy do-

se'a, w ocenie którego powstanie AI na wzór ludzkiego mózgu nie jest możliwe, ponieważ proces myślowy człowieka znacząco różni się od obliczeń wykonywanych przez komputery. Dzieje się tak z dwóch powodów. Po pierwsze, twórcze myślenie w mózgu człowieka uwzględnia wiele zmiennych i całokształt okoliczności, w których funkcjonuje człowiek, podczas gdy szczegółowe obliczenia komputerowe dokonują analizy tylko poszczególnych elementów. Po drugie, ograniczenia wynikające ze stosowania charakterystycznych dla maszyn liczących metod indukcyjnych doprowadziły R. Penrose'a do wniosku, że rozumienie, a tym bardziej świadomość istoty myślącej nie są wynikiem prowadzonych obliczeń, lecz mają charakter nieobliczeniowy. R. Penrose uważa, że procesy myślenia zachodzące w mózgu mogą być zrozumiane jedynie na gruncie mechaniki kwantowej, co w konsekwencji oznacza, że teoria kwantów odnosi się zwykle do obiektów mikroskopowych, a mózg człowieka zalicza się do makroświata. Mało tego, R. Penrose dochodzi do wniosku, że mechanika kwantowa, jaką z ogromnym sukcesem stosujemy do opisu atomów, cząsteczek i wielkiego bogactwa różnorodnych zjawisk, wymaga zasadniczej rewizji. Zob. S. Mrówczyński, Spór wokół sztucznej inteligencji – nieobliczalna świadomość, *Polityka* 1998, Nr 37. Z kolei zdaniem krytyka badań nad sztuczną inteligencją, H. Dreyfusa, ludzka inteligencja i wiedza specjalistyczna zależą przede wszystkim od nieświadomych procesów, a nie świadomej manipulacji symbolicznej, i że te nieświadome umiejętności nigdy nie mogą zostać w pełni ujęte w formalnych regułach. Krytykę H. Dreyfusa uznano za niepoprawnie wrogą, a zdaniem P. McCorduck, jego sztyrdstwo było tak prowokujące, że zraził do siebie każdego, kogo mógł oświecić, co należy uznać za szkodę. Daniel Crevier stwierdził, że: „czas udowodnił trafność i przenikliwość niektórych komentarzy H. Dreyfusa. Gdyby sformułował je mniej agresywnie, konstruktywne działania, które sugerował, mogłyby zostać podjęte znacznie wcześniej. Przykładem, który niejako potwierdza tezę H. Dreyfusa jest praca D. Kahnemanna i A. Tverskyego, w której zebrali dowody na to, że ludzie używają dwóch bardzo różnych metod rozwiązywania problemów, system znany jako „adaptacyjna nieświadomość”, jest szybki, intuicyjny i nieświadomy, a system 2 jest powolny, logiczny i rozważny. Zob. szerzej H. Dreyfus, *Alchemy and AI*, RAND Corporation 1965; H. Dreyfus, *What Computers Can't Do*, New York 1972; P. McCorduck, *Machines Who Think. A Personal Inquiry into the History and Prospects of Artificial Intelligence*, Natick, Massachusetts 2004, s. 236; D. Crevier, *AI: The Tumultuous Search for Artificial Intelligence*, New York 1993, s. 125; D. Kahneman, *Thinking, Fast and Slow*, Farrar, Straus and Giroux 2011. Zdecydowanie odmienne stanowisko na rozwój AI zaprezentował L. Aschenbrenner, według którego, technologia sztucznej inteligencji może wykonać i zautomatyzować zestaw złożonych zadań opartych na wiedzy na poziomie ludzkim lub wyższym już w 2027 r., co będzie miało konsekwencje kształtujące epokę i momentem krytycznym dla przetrwania wolnego świata. Zob. L. Aschenbrenner, *Situational Awareness: The Decade Ahead*, situational-awareness.ai, dostęp: 6.9.2024 r.

⁶ E. Schmidt (przew.), *Final Report. National Security Commission on Artificial Intelligence* z 1.3.2021 r., s. 10–11.

⁷ Y.K. Dwivedi, L. Hughes, E. Ismagilova, G. Aarts, C. Coombs, T. Crick, Y. Duan, R. Dwivedi, J. Edwards, A. Eirug, V. Galanos, P.V. Ilavarasan, M. Janssen, P. Jones, A. Kumar Kar, H. Kizgin, B. Kronemann, B. Lal, B. Lucini, R. Medaglia, K. Le Meunier-FitzHugh, L.C. Le Meunier-FitzHugh,

dać, że AI jest niejednokrotnie rozumiana jako eufemizm dla szerszego trendu obejmującego wykorzystanie zaawansowanej technologii do wykonywania zadań, które jeszcze w niedalekiej przeszłości były realizowane wyłącznie przez człowieka, ale niezawierającej komponentu AI.

W niniejszej monografii uwaga koncentruje się wokół perspektywy kierunku przeobrażenia się modelu zarządzania danymi w administracji federalnej Stanów Zjednoczonych, polegającego na wykorzystaniu technologii AI i ustanowieniu tej technologii fundamentem inteligentnego e-administrowania państwem. O ile zasadniczo technologia AI nie wpływa na zmianę koncepcji demokratycznego państwa prawa oraz zmianę struktury administracji federalnej Stanów Zjednoczonych, to w obszarze zarządzania danymi zmiana ma charakter fundamentalny. Polega ona na odejściu od dotychczasowego scentralizowanego i silosowego modelu zarządzania danymi, który był determinowany hierarchicznością i „resortowością” struktury administracyjnej, w kierunku współdzielenia danych publicznych z wykorzystaniem AI, z uprzednim ich rozproszonym przetwarzaniem⁸.

Celem wdrożenia rozproszonego przetwarzania jest integracja danych w administracji federalnej Stanów Zjednoczonych, która polegałaby na ich wzajemnym kojarzeniu w każdej możliwej konfiguracji⁹. Rozproszone przetwarzanie byłoby oparte na wielkoskalowych klastrach obliczeniowych, mających odpowiednią moc obliczeniową, umieszczonych blisko źródła danych, w którym te dane są generowane, co w konsekwencji pozwoliłoby na zwiększoną dostępność i skrócony czas przesyłania danych, a docelowo zmniejszyłoby liczbę obecnie wymaganych procesów administracyjnych i zwiększyłoby wydajność amerykańskiej administracji federalnej.

S. Misra, E. Mogaji, S. Sharma, J. Bahadur Singh, V. Raghavan, R. Raman, N.P. Rana, S. Samothrakis, J. Spencer, K. Tamilmani, A. Tubadji, P. Walton, M.D. Williams, Artificial Intelligence (AI), Multi-disciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy, *International Journal of Information Management* 2021, vol. 57, doi.org/10.1016/j.ijinfomgt.2019.08.002, dostęp: 8.8.2024 r.

⁸ Por. M. Sakowska-Baryła, Sztuczna inteligencja w sektorze publicznym, w: B. Fischer, A. Pązik, M. Świerczyński (red.), *Prawo sztucznej inteligencji i nowych technologii* 3, Warszawa 2024, s. 29–49.

⁹ Zob. szerzej B. Fischer, Współdzielenie danych jako niezbędny warunek rozwoju sztucznej inteligencji, w: B. Fischer, A. Pązik, M. Świerczyński (red.), *Prawo sztucznej inteligencji i nowych technologii*, Warszawa 2021, s. 91–111; E. Traple, Granice eksploracji tekstów i danych na potrzeby maszynowego uczenia się przez systemy sztucznej inteligencji, w: B. Fischer, A. Pązik, M. Świerczyński (red.), *Prawo sztucznej inteligencji i nowych technologii*, s. 19–41.

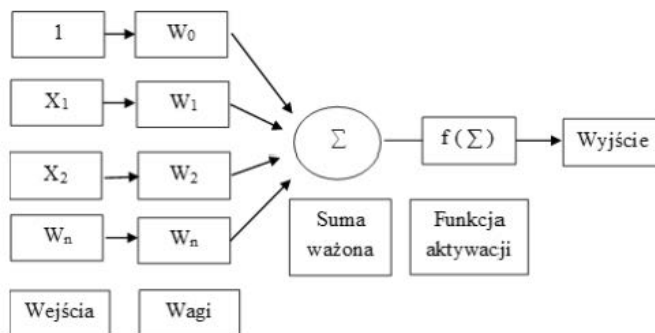
Należy stwierdzić, że dotychczasowy model zarządzania danymi w administracji federalnej przeobraża się w neuronowy model zarządzania danymi z wykorzystaniem AI, w którym pojedyncza jednostka organizacyjna administracji federalnej miałaby funkcjonować w przyszłości – w dużym uproszczeniu – według wzorca sztucznego neuronu, skonstruowanego na wzór neuronu naturalnego, a cała administracja federalna miałaby odpowiadać konstrukcji sztucznej sieci neuronowej. Sztuczny neuron, zwany również perceptronem, to podstawowa jednostka przetwarzająca w sztucznych sieciach neuronowych, które są modelowane na podstawie biologicznego neuronu i stanowią podstawową jednostkę przetwarzania informacji w mózgu człowieka.

Podstawową ideą sztucznego neuronu jest zdolność do przyjmowania danych z otoczenia przez wiele wejść, przetwarzania ich i generowania jednego wyjścia na podstawie określonej funkcji aktywacji. Dla przypomnienia, schemat opracowanego w 1943 r. przez *W. McCulloch'a* i *W. Pitts'a* pierwszego w historii sztucznego neuronu składa się z:

- 1) wejścia neuronu (dendryty), będącego odbiorcą gromadzonych sygnałów (danych) z otoczenia;
- 2) wag (synaps), będących wartościami poddanymi działaniu mnożenia przez odpowiednie współczynniki;
- 3) bloku sumującego lub sumy ważonej (jądra), będącego centrum obliczeniowym neuronu, w którym dochodzi do sumowania pomnożonych sygnałów;
- 4) bloku aktywacji (aksonu), będącego sygnałem wyjściowym neuronu;
- 5) wyjścia, będącego wartością funkcji aktywacji propagowaną do następnej warstwy neuronów¹⁰.

¹⁰ *W.S. McCulloch, W. Pitts, A logical calculus of the ideas immanent in nervous activity, The Bulletin of Mathematical Biophysics 1943, Nr 5(4), s. 115–133, doi:10.1007/BF02478259, dostęp: 8.8.2024 r.*

Rys. 1. Schemat perceptronu McCullocha-Pittsa (źródło: W.S. McCulloch, W. Pitts, A logical calculus).



Porównując powyższy wzorzec do administracji federalnej Stanów Zjednoczonych, poszczególne jej jednostki organizacyjne (np. Departament Energii) mogłyby odpowiadać wejściom dla odbieranych z otoczenia danych (dendryty), następnie danym tym nadawane byłyby wartości (wagi modelu AI), które z kolei byłyby wykorzystane w działaniach matematycznych (algorytmy AI) w rozproszonych wielkoskalowych klastrach obliczeniowych lub centralnych wielkoskalowych klastrach obliczeniowych, za które odpowiadałby człowiek i technologia AI, a wynik każdego takiego działania (model AI) mógłby stanowić podstawę inicjatywy legislacyjnej, w następstwie której rolę funkcji aktywacji pełniłby Kongres lub Prezydent Stanów Zjednoczonych, ze względu na mającą inicjatywę ustawodawczą. W dalszej kolejności zmienione przepisy prawa stanowiłyby podstawę do odbierania nowych danych i poprawy uprzednio zastosowanych wag, algorytmów i modeli AI, które wcześniej mogłyby być uznane za nieefektywne lub niepożądane.

Neuronowy model zarządzania danymi z wykorzystaniem AI wyodrębniono na podstawie przeglądu prawa federalnego bezpośrednio lub pośrednio powiązanego z AI, które przyjęto w latach 2018–2013, zwłaszcza wymienionych poniżej sześciu ustaw federalnych i trzech rozporządzeń wykonawczych Prezydenta Stanów Zjednoczonych z mocą ustawy. Zasadne jest więc udzielenie odpowiedzi na pytanie, w jaki sposób technologia AI, stając się strategicznym komponentem państwa, którym są Stany Zjednoczone, wpłynęła i nadal wpływa na zmiany w modelu zarządzania danymi w administracji federalnej. Otoczenie organizacyjne to administracja federalna Stanów Zjednoczonych, a przez pojęcie administracji federalnej należy rozumieć Prezydenta Stanów

Zjednoczonych oraz rząd i jego członków, a także podległe im agencje oraz inne jednostki organizacyjne lub podmioty. Z kolei próbę badawczą stanowią zgromadzone, wyselekcjonowane, opisane i zinterpretowane akty prawa federalnego dotyczące AI, tj.:

- 1) ustawa o identyfikowaniu danych wyjściowych generatywnych sieci przeciwstawnych z 2020 r. (IOGAN Act)¹¹;
- 2) ustawa o Narodowej Inicjatywie AI z 2020 r. (NAII)¹²;
- 3) ustawa o wykorzystaniu AI w sprawowaniu władzy federalnej z 2020 r. (AI GOV Act)¹³;
- 4) ustawa o powołaniu Krajowej Grupy Zadaniowej ds. Zasobów Badań nad AI z 2021 r. (NAIRR)¹⁴;
- 5) ustawa o szkoleniu z AI dla urzędników federalnych z działów zakupów z 2022 r. (AI Training Act)¹⁵;
- 6) ustawa o rozwijaniu AI w Stanach Zjednoczonych z 2022 r. (AAAA)¹⁶ oraz 3 rozporządzenia wykonawcze (ang. *executive order*) Prezydenta Stanów Zjednoczonych:
 - 1) rozporządzenie wykonawcze Nr 13859 w sprawie utrzymania amerykańskiego przywództwa w obszarze AI z 2019 r. (EO 13859)¹⁷;
 - 2) rozporządzenie wykonawcze Nr 13960 promujące wykorzystanie godnej zaufania AI przez rząd federalny z 2020 r. (EO 13960)¹⁸;
 - 3) rozporządzenie wykonawcze Nr 14110 w sprawie bezpiecznej i godnej zaufania AI z 2023 r. (EO 14110)¹⁹.

Próbkę badawczą uzupełniono o wiele innych norm *hard law* i *soft law* oraz orzecnictwem sądów amerykańskich, które w sposób bezpośredni lub po-

¹¹ *Identifying Outputs of Generative Adversarial Networks Act*, 116th Congress, Public Law 2020, Nr 116–258, s. 2904.

¹² *National Artificial Intelligence Initiative Act of 2020* (NAII), H.R.6216, 116th Congress, 3.12.2020 r.

¹³ *AI in Government Act of 2020*, H.R.2575, 116th Congress, 14.9.2020 r.

¹⁴ *National AI Research Resource Task Force Act of 2020*, 116th Congress, Public Law 2021, Nr 116–283, s. 3890.

¹⁵ *Artificial Intelligence Training for the Acquisition Workforce Act*, 117th Congress Public Law, Public Law 117–207, s. 2551, 17.10.2022 r.

¹⁶ *Advancing American AI Act*, 117th Congress, Public Law 2022, Nr 117–270, s. 1353.

¹⁷ *Executive Order 13859 z 11.2.2019 r., Maintaining American Leadership in Artificial Intelligence*.

¹⁸ *Executive Order 13960 z 3.12.2020 r., Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government*.

¹⁹ *Executive Order 14110 z 30.10.2023 r., Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*.

średni są powiązane z AI. Ponadto w celu ukazania skali zainteresowania technologią AI w obszarze tworzenia prawa w Stanach Zjednoczonych, dodatkowo przedstawiono dane dotyczące przyjętych ustaw na szczeblu stanowym i projektów ustaw federalnych, które wpłynęły do komisji Kongresu Stanów Zjednoczonych w samym tylko 2023 r. Co istotne, w wyodrębnionych do badań przepisach prawa powszechnie obowiązującego wprowadzono po raz pierwszy na świecie 3 różniące się od siebie definicje AI. Definicję zawartą w § 238 lit. (g) NDAA 2019, której zakres przedmiotowy jest najbardziej rozbudowany, omówiono w sposób szczegółowy, nie tylko ze względu na jej nowatorskość i wysoki stopień skomplikowania, ale również ułatwienie zrozumienia przedmiotu badań²⁰.

System prawa federalnego Stanów Zjednoczonych związanego z AI składa się z wielu norm prawnych, w których zawarto reguły postępowania, czyli norm *compliance*. Termin ten można rozumieć jako zbiór powiązanych ze sobą wzorców zachowań w określonych formach, w tym *hard law* i *soft law*, nakładających na adresatów obowiązek dostosowania się do nich²¹. Celem zgodności jest dążenie do tego, aby adresat normy prawnej zachowywał się w sposób oczekiwany przez jej nadawcę, zarówno w odniesieniu do konkretnego zachowania określonego w brzmieniu przepisu prawa, systemu prawa, jak również normy *compliance* o charakterze *soft law*, którymi mogą być wytyczne. W celu przedstawienia próby badawczej w sposób bardziej szczegółowy, w pracy przedstawiono 230 norm *compliance* wyodrębnionych z 6 ustaw federalnych i 3 rozporządzeń wykonawczych Prezydenta Stanów Zjednoczonych z mocą ustawy, które pogrupowano według ich liczebności (od największej do najmniejszej), chronologii wprowadzenia do porządku prawnego, systematyki jednostek redakcyjnych poszczególnych aktów prawnych, aż wreszcie alfabetycznie w przypadku form o tej samej liczbie. *Prima facie* może wydawać

²⁰ Z przeprowadzonych badań dotyczących wyzwań prawnych w sektorze bankowym związanych z AI wynika, że zasadniczym wyzwaniem jest kwestia odpowiedniego skonstruowania i interpretowania definicji legalnej pojęcia „sztuczna inteligencja”. Zob. K. Koźmiński, S. Żółtek, M. Jabłoński, C.W.P. Lewis, J. Czarnocki, A. Leszczyński, M. Macidłowski, S. Sasin, Raport z badania „Aktualne wyzwania prawne związane z zastosowaniem sztucznej inteligencji” w polskim sektorze bankowym, Warszawa 2024, s. 88, stan prawny na 8.4.2024 r.

²¹ L.S. Paine, Governance Beyond the Boardroom: Assessing Risk and Assuring Compliance Throughout Your Organization, Presenter, HBS Audioconference. Lecture at the HBS Audioconference, Harvard Business School Publishing 2003; L.S. Paine, Value Shift: Is Compliance a Catalyst for Sustainable Business Performance? Panel moderator, Frontlines Forum: Shift to Value. Frontlines Forum, San Francisco 2005; L.S. Paine, Performance vs. Compliance: A Global Leader's Guide to Managing Business Conduct, London 2011.

się, że szczegółowe przedstawienie tak licznej grupy norm *compliance* wykracza poza ramy pracy i jest zbyt obszerne, ale ich wartością jest ułatwienie ujęcia zagadnienia w sposób możliwie precyzyjny zagadnieniowo i metodycznie.

Przegląd prawa federalnego związanego z AI stanowi przyczynek do ustalenia kierunku przeobrażenia się sposobu funkcjonowania rządu federalnego Stanów Zjednoczonych zmierzającego do neuronowego modelu zarządzania danymi z wykorzystaniem AI, który docelowo ma doprowadzić do maksymalizacji wykorzystania potencjału tej przełomowej w historii ludzkości technologii, a w konsekwencji spotęgowania wydajności procesów administracyjnych. Co istotne, 14.1.2019 r. Prezydent Stanów Zjednoczonych podpisał ustawę o podstawach kształtowania polityki opartej na dowodach (*Evidence Act*)²². Ogólnie rzecz biorąc, w ustawie tej zobowiązano rząd federalny Stanów Zjednoczonych do ustanowienia procesów w zakresie modernizacji praktyk zarządzania danymi²³, w tym udostępniania danych przez agencje federalne, gromadzenia dowodów na podstawie danych oraz uzyskania efektywności statystycznej w celu wspierania kształtowania polityki i podejmowania decyzji politycznych²⁴.

Prima facie wydaje się, że rząd federalny Stanów Zjednoczonych podjął się jednej z największych inicjatyw w zakresie transformacji rządu i administracji federalnej od początku jej istnienia. Inicjatywa ta polega na przeobrażeniu się w taki sposób, aby umożliwić podejmowanie decyzji na podstawie dowodów opartych na danych. Zrealizowanie jej nie tylko ma umożliwić dostosowanie się do dynamicznie zmieniającego się otoczenia, zwiększyć efektywność i wydajność rządu i administracji federalnej, ale i w perspektywie długoterminowej korzystnie wpłynąć na gospodarkę i globalną konkurencyjność Stanów Zjednoczonych. Holistyczny przegląd prawa federalnego doprowadził do wniosku,

²² *Foundations for Evidence-Based Policymaking Act*, the 115th U.S. Congress, Public Law 2019, Nr 115, s. 435.

²³ W celu monitorowania ustanawiania procesów w zakresie modernizacji praktyk zarządzania danymi w Stanach Zjednoczonych uruchomiono stronę: *Evaluation.gov*, na której rząd federalny informuje o bieżącej ewaluacji programów federalnych i Rady Urzędników ds. Ewaluacji, źródło: *evaluation.gov*, dostęp: 8.8.2024 r.

²⁴ W 2017 r. amerykańska Komisja ds. Polityki Opartej na Dowodach (*Evidence Commission*) wydała raport dla Kongresu i Prezydenta Stanów Zjednoczonych zawierający 22 zalecenia, które dotyczyły poprawy dostępu do danych, zwiększeniu ochrony ich prywatności i zdolności do tworzenia polityki opartej na dowodach. W 2022 r. opublikowano raport dotyczący postępu wdrażania 22 zaleceń oraz modernizacji praktyk zarządzania danymi w Stanach Zjednoczonych: *N. Hart, S. Stefanik, Evidence Commission after 5 years: A Progress Report on the Promise for a More Evidence – Informed Society*, Data Foundation 2022.

że w Stanach Zjednoczonych wyodrębnia się neuronowy model zarządzania danymi w administracji federalnej z wykorzystaniem AI, który jest wzorowany na sztucznym neuronie i sztucznej sieci neuronowej.

Neuronowy model zarządzania danymi z wykorzystaniem AI może umożliwiać i ułatwiać zarządzanie publiczne w demokratycznym państwie prawa oparte na rozproszonym przetwarzaniu danych, tworząc zdolność do rozpoznawania ukrytych prawidłowości i korelacji w dostępnych danych, w tym nieprzetworzonych, grupowania ich i klasyfikowania, a także nieustannego uczenia się i doskonalenia, aż wreszcie potęgowania wydajności procesów administracyjnych. Wyróżniającą cechą sztucznej sieci neuronowej jako narzędzia informatycznego jest możliwość komputerowego rozwiązywania problemów bez ich uprzedniej matematycznej formalizacji, a także braku konieczności odwoływania się do jakichkolwiek teoretycznych założeń.

W pracy podjęto się udzielenia odpowiedzi na pytanie, w jaki sposób technologia AI stającą się strategicznym komponentem organizacji jaką są Stany Zjednoczone, wpłynęła dotychczas i nadal wpływa na zmiany w dotychczasowym modelu zarządzania danymi w administracji federalnej. W następstwie udzielenia odpowiedzi na tak postawione pytanie podjęto próbę zweryfikowania hipotezy głównej, że dotychczasowy model zarządzania danymi w administracji federalnej przeobraża się w neuronowy model zarządzania danymi z wykorzystaniem AI, w którym pojedyncza jednostka organizacyjna administracji federalnej miałaby funkcjonować w przyszłości – w dużym uproszczeniu – według wzorca sztucznego neuronu skonstruowanego na wzór neuronu naturalnego, a cała administracja federalna miałaby odpowiadać konstrukcji sztucznej sieci neuronowej.

W opracowaniu sformułowano 3 hipotezy pomocnicze, które mają uzupełnić i ułatwić zweryfikowanie hipotezy głównej. Pierwsza hipoteza pomocnicza ma na celu zweryfikowanie, że przyjęte prawo federalne w Stanach Zjednoczonych związane z AI odgrywa istotną rolę w zwiększeniu świadomości i podnoszeniu wiedzy na temat technologii AI, w tym:

- 1) osiąganiu politycznych celów strategicznych Stanów Zjednoczonych związanych z AI;
- 2) poprawie nastrojów społecznych względem AI;
- 3) wykorzystaniu federalnej matrycy instytucjonalnej w procesie rozwoju AI;
- 4) zrozumieniu technologicznych i infrastrukturalnych aspektów neuronowego modelu zarządzania danymi z wykorzystaniem AI.

Zweryfikowanie tak postawionej hipotezy pomocniczej jest o tyle istotne, że, z jednej strony, AI powinna odpowiadać potrzebom państwa i obywateli, a z drugiej, że umiejętność i efektywność jej praktycznego wykorzystania jest pochodną świadomości i wiedzy na temat AI oraz interakcji pomiędzy człowiekiem i technologią AI.

Druga hipoteza pomocnicza ma na celu zweryfikowanie, że ochrona i bezpieczeństwo wag modelu AI, które są matematycznym wyrażeniem połączeń między sztucznymi neuronami tworzącymi sztuczną sieć neuronową, ma krytyczne znaczenie nie tylko dla wydajności wytrenowanego modelu i wpływu na efektywność neuronowego modelu zarządzania danymi z wykorzystaniem AI, ale przede wszystkim dla ryzyk i wyzwań z nimi związanych. Z jednej strony, wagi modelu AI mogą wspomagać zarządzanie przez dokonywanie bardziej precyzyjnych prognoz i podejmowanie trafniejszych decyzji, a z drugiej strony, mogą generować błędną ocenę wartości danych. Co istotne, owa błędna ocena może nie tylko zmniejszać efektywność działań rządu i administracji federalnej, ale nawet zagrażać bezpieczeństwu narodowemu Stanów Zjednoczonych.

Trzecia hipoteza pomocnicza koncentruje się na zweryfikowaniu założenia, że neuronowy model zarządzania danymi z wykorzystaniem AI przyjęty w latach 2018–2023 w 6 ustawach federalnych i 3 rozporządzeniach wykonawczych Prezydenta Stanów Zjednoczonych z mocą ustawy związane z AI ma umożliwić współdzielenie danych w następstwie wdrożenia rozproszonego przetwarzania, w którym umieszczenie mocy obliczeniowej i przechowywanie danych następuje blisko źródła ich generowania, co zwiększa dostępność i skraca czas przesyłania danych, ale nie wyklucza centralnego ich wykorzystania²⁵. Powyższe założenie stanowi płaszczyznę porównawczą względem dotychczasowego scentralizowanego modelu zarządzania danymi w administracji federalnej, co wynikało z hierarchicznej oraz „resortowej” struktury i kompetencji administracji federalnej. O ile struktura i zakres kompetencji administracji federalnej nie uległy zmianie, to w obszarze zarządzania danymi następuje rozproszenie mocy obliczeniowej i zasobów bazodanowych w celu umożliwienia ich rozproszonego przetwarzania w wielkoskalowych i lokalnych

²⁵ S.S. Sonone, K. Saini, S. Jadhav, M.S. Sankhla, V. Nagar, Chapter Twelve – The future of edge computing, *Advances in Computers* 2022, vol. 127, s. 333–358, doi.org/10.1016/bs.adcom.2022.02.009, dostęp: 8.8.2024 r.

klastrach obliczeniowych, z możliwością ich integrowania i wzajemnego kojarzenia w każdej możliwej konfiguracji²⁶.

W monografii zastosowano dobór nieprobabilistyczny, z relatywnie nie-
zbyt dużą liczebnością próby badawczej, czyli przyjęte w latach 2018–2023
6 ustaw federalnych i 3 rozporządzenia wykonawcze Prezydenta Stanów Zjed-
noczonych z mocą ustawy. Niewątpliwie liczba dziewięciu aktów prawa fede-
ralnego stanowi znaczne ograniczenie badawcze, ale zakwalifikowanie prze-
pisów prawa federalnego do próby badawczej może decydować o jej nowa-
torskości i przesądzać o wyższej wartości uogólniającej wyników badania. Za
przyjęciem takiego stanowiska przemawia również fakt, że owe 6 ustaw i 3
rozporządzenia są jednymi z pierwszych aktów prawa powszechnie obowiąz-
ującego na świecie, które w sposób tak szczegółowy regulują podejście do tech-
nologii AI.

W badaniu wykorzystano metodę monograficzną, w tym dokonano holi-
stycznego przeglądu prawa federalnego, w którym uregulowano zagadnienie
AI. Szczególną uwagę poświęcono pierwszej na świecie legalnej definicji AI
zawartej w § 238 lit. (g) NDAA 2019. Ponadto przeglądowi poddano szereg
innych norm prawa *hard law* i *soft law*, które w sposób bezpośredni lub po-
średni są powiązane ze AI. Posłużono się również metodą badania dokumen-
tów o znaczeniu strategicznym oraz urzędowych i naukowych dokumentów,
co miało umożliwić ustalenie podejścia do AI, w tym wyodrębnić szanse i za-
grożenia wynikające z wykorzystania AI.

W pracy dokonano także analizy i krytyki piśmiennictwa, w tym litera-
tury branżowej i naukowej z wielu obszarów, reprezentujących niejednokrot-
nie uzupełniające się wzajemnie paradygmaty i punkty widzenia. Wspomniana
analiza i krytyka miały ukazać kierunek dyskusji naukowej w obszarze AI,
ze szczególnym uwzględnieniem Stanów Zjednoczonych, a także umożliwić
uzupełnienie wiedzy o treści i komunikaty zawarte w opracowaniach nauko-
wych. Należy wspomnieć, że praca ma charakter prognozujący, a właściwie
stanowi koncepcję perspektywistyczną, której celem jest niejako przewidywa-
nie przyszłego zdarzenia, dotychczas nieznanego i nienazwanego, na podsta-

²⁶ Por. G. Zyskind, O. Nathan, A. Pentland, Decentralizing privacy: Using blockchain to protect personal data, Conference 2015 IEEE Security and Privacy Workshops.2015, s. 180–184; A. Pentland, A. Lipton, T. Hardjono, Building the New Economy. Data as Capital, Cambridge 2021; K. Miller, Digital Economy Lab Fellow Alex „Sandy” Pentland: Building a Better, Safer Digital Ecosystem. The MIT professor describes practical approaches to improving the internet, Stanford University 2024.

wie zdarzeń znanych i nazwanych, które wyodrębniono z reguł postępowania określonych w prawie federalnym Stanów Zjednoczonych.

Chciałbym złożyć serdeczne podziękowania dla Pana *Malcolma K. Sparrow'a*, Profesora Praktyki Zarządzania Publicznego w Harvard John F. Kennedy School of Government, który zafascynował mnie odkrywaniem synergii między rzemiosłem regulacyjnym i zarządzaniem wydajnością, a także nieustannym dostrzeganiem perspektywy strategicznej. Nie mógłbym pominąć podziękowań dla Pana dr hab., Prof. UW, *Sebastiana Skuzy* z Wydziału Zarządzania Uniwersytetu Warszawskiego, za silne bodźce do dalszego rozwoju i rozszerzania pola badawczego, wspólnie realizowane projekty zawodowe i naukowe, które wpłynęły na moje wybory i dążyły do nadania im waloru aplikacyjnego, aż wreszcie za wzorzec ludzkiej pracowitości, rzetelności i życzliwości.

Książkę pragnę zadedykować moim rodzicom: *Wando i Andrzej* – *memini tui, memento mei*.

Rozdział I. Wprowadzenie do AI, organizacji administracji federalnej, zarządzania publicznego i neuronowego modelu zarządzania danymi

§ 1. *Artificial intelligence* (AI) – geneza powstania i jej istota

W opublikowanym w 1950 r. przełomowym artykule A. Turing zadał pytanie, czy maszyny mogą myśleć. W celu udzielenia odpowiedzi na tak postawione pytanie nie podjął się słownikowej deskrypcji terminów „maszyna” i „myśleć”, ale przeprowadził eksperyment, który nazwał imitacją gry (*imitation game*)¹. Eksperyment polegał na umożliwieniu komunikowania się mężczyzny (A), kobiety (B) i przesłuchującego (C), w celu ustalenia przez przesłuchującego, będącego w innym pokoju niż pozostałe dwie osoby, która z nich jest mężczyzną, a która kobietą. Jeżeli przesłuchujący wykorzystując język naturalny nie rozpoznał, że jedną ze stron dialogu jest maszyna, spośród kobiety i mężczyzny, to w przyjętym założeniu wynik testu stanowił dowód, iż owa maszyna reagowała w sposób możliwie zbliżony do ludzkiego. Według A. Turinga, nierozpoznanie maszyny przez co najmniej 30% przesłuchujących, podczas 5-minutowego dialogu prowadzonego przez każdego z nich, miało przesądzać o tym, że maszyna potrafi myśleć. Test Turinga zakończył się niepowodzeniem, czyli nie udało się osiągnąć założonego pułapu 30%, ale w 2014 r. pomyślnie udało się przekonać 33% przesłuchujących, że mają do czynienia z człowiekiem, a nie z maszyną².

¹ A.M. Turing, *Computing machinery*, s. 433–434.

² *Turing Test success marks milestone in computing history*, University of Reading 2014, archive.reading.ac.uk/news-events/2014/June/pr583836.html, dostęp: 9.9.2024 r.; K. Warwick, H. Shah, *Can machines think? A report on Turing test experiments at the Royal Society*, Jo-

Co prawda test *Turinga* stanowił podstawę do opracowania definicji AI, a dokładniej jednej z jej kategorii, charakteryzującej się działaniem jak człowiek, jednak S. Russell i P. Norvig dokonali przeglądu szeregu innych definicji i łącznie wyodrębnili 4 uniwersalne cechy AI:

- 1) myślenie jak człowiek (*Thinking humanly*) – automatyzowanie czynności, które są kojarzone z ludzkim myśleniem i takimi czynnościami jak podejmowanie decyzji, rozwiązywanie problemów, uczenie się, *etc.*;
- 2) myślenie racjonalne (*Thinking rationally*) – badanie zdolności umysłowych przez wykorzystanie modeli obliczeniowych;
- 3) działanie jak człowiek (*Acting humanly*) – naśladowanie zachowania człowieka przez maszyny (inteligentny agent programowy lub robot humanoidalny) do realizowania zadań z użyciem inteligencji, w sposób, który byłby realizowany przez człowieka;
- 4) działanie racjonalne (*Acting rationally*) – wyjaśnianie i naśladowanie zachowania człowieka przez maszyny przez obliczenia inteligentne³.

Jak można zauważyć, dwie pierwsze cechy obejmują procesy myślowe i rozumowanie, natomiast dwie kolejne dotyczą zachowania. Z kolei cechy zawarte w pkt 1 i 3 decydują o wierności odzwierciedlenia ludzkiego zachowania, podczas gdy cechy określone w pkt 2 i 4 są związane z miarą idealnej wydajności, zwanej racjonalnością, czyli postępowaniem w sposób „właściwy” w danych okolicznościach, uzależnionym od zakresu udostępnionych systemowi danych i mocy obliczeniowej. Im większa moc obliczeniowa i zakres danych, tym bardziej prawdopodobne właściwe zachowanie.

Działanie maszyny w sposób podobny do myślenia człowieka ma miejsce wtedy, gdy system analizuje własne zasoby danych i obliczenia wykonane na ich podstawie oraz bada związki przyczynowo-skutkowe, a także weryfikuje sposób realizowania wybranych procesów⁴. Myślenie racjonalne jest niczym innym jak logicznym rozumowaniem, czyli zdolnością do integrowania i ana-

urnal of Experimental & Theoretical Artificial Intelligence 2015, Nr 28(6), s. 990–993, doi:10.1080/0952813X.2015.1055826, dostęp: 8.8.2024 r.

³ S.J. Russell, P. Norvig, *Artificial Intelligence. A Modern Approach*, Pearson Education 2010, s. 1–5 i powołana tam literatura. Zob. szerzej R.J. Schalkoff, *Artificial Intelligence: An Engineering Approach*, McGraw-Hill 1990; E. Charniak, D.V. McDermott, *Introduction to Artificial Intelligence*, Addison–Wesley 1985; R. Kurzweil, *The Age of Intelligent Machines*, Cambridge 1990; R. Bellman, *An Introduction to Artificial Intelligence: Can Computers Think?*, San Francisco 1978; P. McCorduck, *Machines Who Think*.

⁴ Zob. N. Wiener, *Cybernetics or Control and Communication in the Animal and the Machine*, Cambridge 2019, s. 233–249, doi.org/10.7551/mitpress/11810.001.0001, dostęp: 8.8.2024 r.

lizi dostępnych danych. W ramach testu *Turinga* użyto systemu działającego w sposób podobny do działania człowieka, na który to składały się:

- 1) przetwarzanie języka naturalnego (*natural language processing*)⁵;
- 2) reprezentacja wiedzy (*knowledge representation*)⁶;
- 3) zautomatyzowane rozumowanie (*automated reasoning*);
- 4) uczenie maszynowe (*machine learning*)⁷.

O ile w ramach testu *Turinga* celowo uniemożliwiono bezpośrednią fizyczną interakcję między przesłuchującym a przesłuchiwanym, ponieważ kontakt fizyczny nie był konieczny dla potwierdzenia lub zaprzeczenia postawionej hipotezy, to gdyby do takiego kontaktu doszło, wówczas niezbędnymi zdolnościami systemu byłyby również rozpoznawanie obrazu (*computer vision*)⁸ i robotyka (*robotics*)⁹. Racjonalne działanie oznacza działanie w taki sposób, aby osiągnąć wyznaczone cele, biorąc pod uwagę zakres danych i dostępu do nich oraz moc obliczeniową, poczynwszy od percepcji, planowania, rozumowania, uczenia się, komunikacji, podejmowania decyzji, a skończywszy na podjęciu adekwatnego działania. Mówiąc inaczej, inteligentny agent programowy lub robot humanoidalny działa racjonalnie, jeśli ma dostęp do danych i zdol-

⁵ Zob. szerzej R.C. Schank, L. Tesler, A Conceptual Dependency Parser for Natural Language, International Conference on Computational Linguistics COLING 1969, Nr 2, Sweden 1969, acanthology.org/C69-0201.pdf, dostęp: 9.8.2024 r.

⁶ R. Davis, H. Shrobe, P. Szolovits, What is a Knowledge Representation?, AI Magazine 1993, Nr 14(1), s. 17–33, groups.csail.mit.edu/medg/people/psz/ftp/k-rep.html, dostęp: 9.8.2024 r.

⁷ A. Samuel, Some Studies in Machine Learning Using the Game of Checkers, IBM Journal of Research and Development 1959, vol. 3(3), s. 210–229, doi:10.1147/rd.33.0210, dostęp: 8.8.2024 r.; R. Kohavi, F. Provost, Glossary of terms, Machine Learning 1998, vol. 30, Nr 2–3, s. 271–274.

⁸ Zob. szerzej R. Klette, Concise Computer Vision, Berlin 2014; L.G. Shapiro, G.C. Stockman Computer Vision, Prentice Hall 2001; T. Morris, Computer Vision and Image Processing, New York 2004; B. Jähne, H. Haußecker, Computer Vision and Applications, A Guide for Students and Practitioners, Academic Press 2000, s. 1–5.

⁹ H.V. Daniel, Smart Robots, A Handbook of Intelligent Robotic Systems, Chapman and Hall 1985, s. 141.

W przypadku robotyki warto wspomnieć o trzech kwestiach. Po pierwsze, że termin „robot” został użyty po raz pierwszy w sztuce teatralnej: J. Capek, R.U.R. (Rossum’s Universal Robots), Prague 1921; por. A. Roberts, The History of Science Fiction, New York 2006, s. 168. Po drugie, w latach 40. XX w. I. Asimov wprowadził koncepcję robotyki w swojej serii opowiadań i książce o robotach, w których ujął słynne „Trzy Prawa Robotyki”, które miały regulować zachowanie robotów i wpłynęły na późniejsze rozważania etyczne i techniczne dotyczące robotyki (por. I. Asimov, I, Robot, Doubleday 1950). Po trzecie, w latach 50. XX w. amerykański wynalazca G. Devol stworzył pierwszego przemysłowego robota, nazwanego „Unimate”, a jego prace nad automatyzacją i robotyką doprowadziły do powstania pierwszych robotów przemysłowych używanych w produkcji (por. G. Devol, Industrial Robot, U.S. Patent Nr 2,988,237 wprowadzony w 2011 r., invent.org/inductees/george-devol), dostęp: 11.8.2024 r.