

Andrzej Radomski
Katedra Informatyki i Kultury Cyfrowej UMCS
e-mail: andrzejradomski64@gmail.com
ORCID: 0000-0002-1735-605X

O możliwości zastosowania algorytmów sztucznej inteligencji do badań humanistycznych (na wybranych przykładach)

DOI: <http://dx.doi.org/10.12775/ZN.2020.009>

Abstrakt. Sztuczna inteligencja stała się jednym z najważniejszych elementów współczesnego świata. Jest wykorzystywana w różnych przejawach ludzkiej działalności. Odnosi także sukcesy w naukach przyrodniczych. W artykule podejmuje się problem zastosowania algorytmów sztucznej inteligencji do badania świata kultury, czyli tradycyjnego przedmiotu zainteresowań humanistów. Badacze kultury mają też coraz częściej do czynienia z wielkimi kolekcjami danych: *big data*. Dużych zestawów danych kulturowych nie da się badać za pomocą klasycznych metod. Autor postuluje, aby do badania współczesnych danych kulturowych użyć sztucznej inteligencji. Pokazuje przykłady takich badań – także z własnej praktyki badawczej. Omawia wybrane algorytmy i aplikacje, które można zastosować w humanistyce i w naukach społecznych. Autor uważa, że zastosowanie sztucznej inteligencji do badań w humanistyce może przynieść nowe rezultaty poznawcze i zmianę wielu panujących do tej pory poglądów.

Słowa kluczowe: sztuczna inteligencja; humanistyka; kultura; *big data*; uczenie maszynowe

About the Possibility of Applying Artificial Intelligence Algorithms to Humanities Research (on Selected Examples)

Abstract. Artificial intelligence has become one of the most important elements of the modern world. It is used in various human practices. It is also successful in natural sciences. The article deals with the problem of using artificial intelligence algorithms to study the world of culture. Culture researchers are also increasingly dealing with large collections of data called *big data*. Large cultural data sets cannot be tested using classical methods. The author postulates that artificial intelligence should be used to study contemporary cultural data. He shows examples of such research – also from his own research practice – and discusses selected algorithms and applications that can be applied in the humanities and social sciences. The author thinks that the use of artificial intelligence for research in humanities can bring new cognitive results and a change in many prevailing views.

Keywords: artificial intelligence; humanities; culture; *big data*; machine learning

Wstęp

O tzw. sztucznej inteligencji (AI)¹ pisze się ostatnio bardzo dużo. Stała się ona nie tylko ważną dziedziną badań, ale także przedmiotem dyskusji filozoficznych, rozważań praktyków życia gospodarczego czy modnym tematem prezentowanym w różnych mediach. Niewątpliwie na ten stan rzeczy złożyły się spektakularne sukcesy AI w wielu dziedzinach życia społecznego. Pierwszym takim sukcesem, który odbił się szerokim echem w świecie i zwrócił uwagę opinii publicznej na to zjawisko, było zwycięstwo superkomputera Deep Blue w grze szachy nad mistrzem świata Garim Kasparowem (1997)². Od 2010 r. odnotowuje się szereg kolejnych sukcesów AI. Potrafi ona komponować muzykę, malować obrazy, dobierać spersonalizowane reklamy – nawet dla pojedynczych odbiorców – czy komunikować się z internautami na tzw. chatbotach. I są to tylko wybrane przykłady jej zastosowań. W końcu AI trafiła także do praktyki naukowej.

Sztuczna inteligencja (jeśli chodzi o naukę) jest wykorzystywana m.in. do poszukiwania nowych planet, diagnozowania nowotworów, ustalania źródeł epidemii czy odkrywania nowych leków³. Jej spore osiągnięcia na polu dyscyplin przyrodniczych stwarzają okazję do postawienia pytania o możliwość zastosowania AI do poznawania świata kultury, czyli tradycyjnego przedmiotu badania różnych dyscyplin humanistycznych (zwłaszcza że do tej pory o tej kwestii praktycznie nie dyskutowano). Pytanie to można rozbić na dwa, bardziej szczegółowe zagadnienia, a mianowicie: 1) w jakich obszarach badań humanistycznych AI może być szczególnie przydatna (obecnie, rzecz jasna); 2) za pomocą jakich algorytmów sztucznej inteligencji można badać określone problemy z zakresu nauk humanistycznych? I odpowiedziom na powyższe dwa pytania będzie poświęcony niniejszy tekst.

Nie da się omówić w krótkim artykule wszystkich możliwych zastosowań AI w humanistyce. Stąd uwaga zostanie skoncentrowana na projektach dających największe nadzieje i mających już spore osiągnięcia, tj. analizie języka naturalnego (NLP) oraz rozpoznawaniu obrazów (*computer vision*).

Podstawą do zaprezentowania możliwości poznawczych AI w wybranych dziedzinach humanistycznych będą przede wszystkim badania własne autora. Ich

¹ Ang. *artificial intelligence*.

² Jeszcze większym osiągnięciem z tego zakresu była wygrana komputera AlphaGo w 2016 r. z Lee Sedolem, mistrzem w chińskiej grze GO, która uchodzi za bardziej skomplikowaną od szachów.

³ Brytyjska firma farmaceutyczna Exscientia w marcu 2020 r. przystąpiła do testowania leku na pacjentach z zaburzeniami obsesyjno-kompulsywnymi, który został wynaleziony przez AI. W opracowywaniu leku wykorzystano *active learning*, rodzaj maszynowego uczenia nadzorowanego, w którym algorytm może sam poprosić naukowca o opisanie potrzebnych danych. Wykorzystuje się tę technikę głównie tam, gdzie danych jest bardzo dużo, a robienie tego ręcznie nie opłacałoby się lub zajmowało zbyt wiele czasu. *Active learning* pozwala systemowi nadawać priorytet najbardziej obiecującym związkom i uczyć się szybciej niż ludzie. Zob. <https://www.sztucznainteligencja.org.pl/pierwszy-lek-od-si-bedzie-testowany-na-ludziach/?fbclid=IwAR2Br-VdiFxpZvLhHrpbzdt0bCiz212PH8GNKnCjcATSC8wxDJRZlvcEhnM> (dostęp: 1.10.2020).

wyniki zostaną także uzupełnione przykładami z analogicznych analiz przeprowadzonych przez innych badaczy – tak, aby czytelnik miał jak najbardziej szeroki przegląd możliwości uczenia maszynowego⁴ odnoszącego się do humanistyki.

Narzędzia i metody

Sztuczna inteligencja zalicza się do tzw. uczenia maszynowego. Obecnie wyróżnia się jego dwa podstawowe typy: klasyczne (oparte na statystyce) i tzw. głębokie (ang. *deep learning*). W przypadku klasycznego uczenia maszynowego człowiek wprowadza dane, a także oczekiwane odpowiedzi (sic!), a komputer ma utworzyć reguły pozwalające na ich uzyskanie. Reguły te mogą zostać użyte w celu przetworzenia nowych danych i uzyskania oczekiwanych odpowiedzi (Flasiński 2018, s. 23). Klasyczne uczenie maszynowe wykorzystuje się zazwyczaj do tworzenia modeli regresji liniowej, klasyfikacji czy drzew decyzyjnych. Zakres *deep learningu* jest trochę szerszy – począwszy od sztuki (zwłaszcza generatywnej), poprzez usługi biznesowe, a skończywszy na różnych technologiach, jak np. autonomiczne samochody czy roboty.

Głębokie uczenie maszynowe może mieć postać nadzorowaną i nienadzorowaną. W przypadku tej pierwszej, tworząc model ML (*machine learning*), wprowadzamy dane z etykietami (np. człowiek czy samochód) i model uczy się rozpoznawać nowe obiekty (np. samochody), opierając się na przekazanych przez nas informacjach. W przypadku maszynowego uczenia nienadzorowanego sieć nie otrzymuje wyniku końcowego (np. etykiety samochodu), tylko przeanalizowawszy dane wejściowe, musi „odgadnąć” albo wyłonić grupy obiektów lub/i jakieś występujące w tym zbiorze wzorce (np. pojazdy określonego koloru). Mówiąc jeszcze inaczej: dwie główne metody stosowane w uczeniu nienadzorowanym to analiza składowych głównych oraz analiza skupień. Pierwsza z nich jest wykorzystywana do zmniejszania wymiarowości danych poprzez odkrywanie i odrzucanie cech, które niosą ze sobą najmniej informacji. Analiza skupień jest używana do grupowania lub segmentowania zestawów danych ze wspólnymi atrybutami w celu eks-trapolacji występujących w nich zależności. Identyfikuje podobieństwa w danych i pozwala na grupowanie danych, które nie zostały oznaczone, sklasyfikowane ani skategoryzowane. Ponieważ analiza skupień bazuje na obecności lub braku takich podobieństw w nowych danych, może być wykorzystana, aby wykryć nietypowe dane, które nie pasują do żadnej grupy.

Głębokie uczenie maszynowe, w skład którego wchodzi sztuczna inteligencja, składa się z trzech głównych działań: a) piszemy określony program bądź wyko-

⁴ Programy sztucznej inteligencji zalicza się do uczenia maszynowego, a mówiąc ściślej: tzw. głębokiego uczenia maszynowego.

rzystujemy konkretną bibliotekę (w rodzaju Scikit-Learn czy Keras), b) następnie wprowadzamy pewne dane i c) na wyjściu otrzymujemy jakąś wiedzę. Zatem w celu wykorzystania AI musimy zgromadzić: a) dane wejściowe (np. obrazy, dźwięki czy słowa), b) oczekiwane wyniki (np. etykiety samochodu, budynku, kobiety masajskiej etc.), c) opracować sposób automatycznego pomiaru działania algorytmu (np. procentową wartość poprawnie sklasyfikowanych zdjęć). Uczenie głębokie tym będzie różnić się od tego „klasycznego” (np. jednowarstwowego), że będzie zawierało wiele warstw. Współcześnie mogą już być ich setki (Chollet 2018, s. 24–25). Analiza danych przez każdą kolejną warstwę ma zapewnić lepszy wynik końcowy (czyli oczekiwany na wejściu, np. poprawne odgadnięcie naszego samochodu czy kobiety masajskiej). Uczenie tego typu składa się tu z wielu cykli (*epoche*) i kończy się, gdy otrzymamy automatyczną informację, że nasz model z dużym stopniem prawdopodobieństwa (np. 90-procentowym) rozpoznaje dane obiekty.

Za pomocą algorytmów uczenia maszynowego (niezależnie od jego wariantu) możemy: a) gromadzić i „oczyszczać” różne kolekcje danych, i to niezależnie od ich skali, b) automatycznie je analizować, c) automatycznie je wizualizować. Automatyka maszynowa jest tutaj kluczowym elementem, gdyż „ręcznie” w żaden sposób nie ogarniemy tysięcy (nie mówiąc już o większych kolekcjach) tekstów, zdjęć czy filmów.

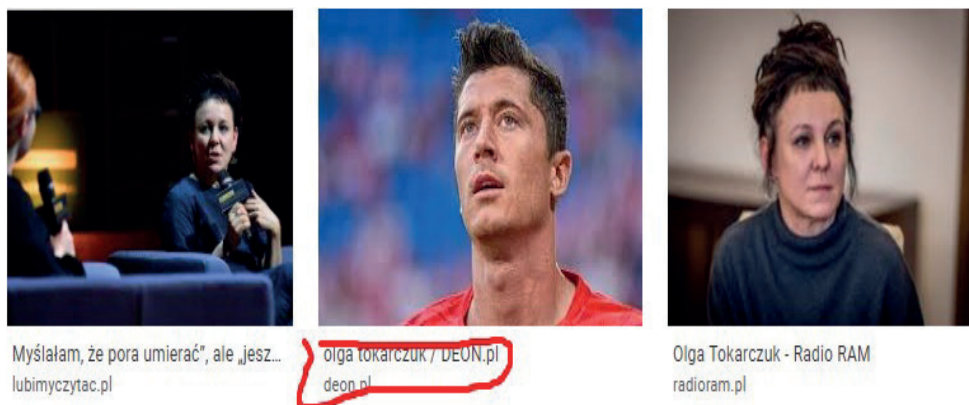
Obecnie dominują dwa podejścia do pracy z AI. Pierwsze polega na tradycyjnym programowaniu sieci neuronowych (bądź np. rojowskich czy rekurencyjnych) – z wykorzystaniem popularnych bibliotek programistycznych typu: Scikit-Learn, TensorFlow czy Keras. Drugie to rozwiązania typu AutoML, w ramach których ładujemy własne dane, a sieć trenuje na nich i analizuje je. To podejście nie wymaga umiejętności programistycznych po stronie użytkownika. W tym przypadku bowiem AI rozwija i udoskonala samą siebie. Na rynku pojawiło się wiele tego typu usług. Można tu wspomnieć o Microsoft Azure, Amazon Web Service (AWS) czy Google Vision.

Autor niniejszego artykułu preferował to drugie rozwiązanie – głównie po to, aby pokazać, że z algorytmów sztucznej inteligencji mogą korzystać nawet osoby nieposiadające praktycznie żadnych umiejętności programistycznych. Wykorzystano następujące usługi i programy: płatne typu: Amazon Web Services (AWS) i InfraNodus Lab, darmowe typu: Google Cloud Vision i Create ML.

Przeprowadzone badania dotyczyły dwóch dziedzin: rozpoznawania obrazów (*computer vision*) oraz analizy języka naturalnego (NLP) – czyli ważnych przedmiotów badań różnych dyscyplin humanistycznych, a także społecznych.

Computer vision można określić jako identyfikację i przetwarzanie obiektów na obrazach i filmach w taki sam sposób jak ludzie. Do niedawna widzenie komputerowe działało tylko w ograniczonym zakresie. Dzięki postępom sztucznej inteligencji i innowacjom w głębokim uczeniu się i sieciach neuronowych udało

się w ostatnich latach dokonać ogromnego skoku w tej dziedzinie i przewyższyć ludzi w niektórych zadaniach związanych z wykrywaniem i znakowaniem obiektów. W wyszukiwarkach internetowych bowiem, gdy poszukujemy określonego zdjęcia bądź całej ich kolekcji, to algorytmy szukają po tagach i innego typu etykietach. Jest to jednakże zawodna metoda. Dla przykładu, gdy chcemy znaleźć serię zdjęć określonej postaci – powiedzmy Olgi Tokarczuk (polskiej noblistki), to możemy otrzymać następujący rezultat, jak na zdjęciu nr 1.



Zdjęcie 1. Wynik poszukiwania zdjęcia Olgi Tokarczuk w wyszukiwarkach internetowych; Grafika Google, oprac. własne

Algorytmy Google’a bowiem „nie patrzą” na postać, tylko na to, jakim tagiem jest „podpisana” dana fotografia. Jeśli znany polski piłkarz Robert Lewandowski został opisany jako Olga Tokarczuk, to jego portret znajdzie się w kolekcji zawierającej fotografie polskiej noblistki. Musimy zatem posłużyć się inną metodą i co za tym idzie innymi algorytmami. Będą to: Create ML oraz Google Cloud Vision.

Jeśli np. w Create ML załadujemy do komputera miliony zdjęć kotów i zdjęcia zostaną „obrobione” przez odpowiednie algorytmy, które przeanalizują kolory na zdjęciu, kształty, odległości między kształtami, gdzie obiekty graniczą ze sobą itd., to po takiej analizie zidentyfikowany zostanie profil tego, co oznacza „kot”. Po takim treningu maszyna, gdy otrzyma nowe zdjęcia (już bez etykiety), będzie mogła z dużym prawdopodobieństwem zidentyfikować zwierzęta nazywane kotami. Mówiąc jeszcze inaczej: algorytmy spod znaku *computer vision* potrafią analizować sam obraz/obrazy i dostarczać określonych informacji o obiektach znajdujących się na danym zdjęciu/zdjęciach czy obrazach ruchomych. Przede wszystkim chodzi tutaj o tzw. wykrywanie twarzy i jej śledzenie. Wykrywanie twarzy to proces automatycznego lokalizowania ludzkich twarzy w mediach wizualnych (obrazy cyfrowe lub video). Rozpoznawanie twarzy automatycznie określa, czy dwie twarze prawdopodobnie odpowiadają tej samej osobie, a śledzenie twarzy

rozszerza wykrywanie twarzy na sekwencje wideo. Można śledzić dowolną twarz pojawiającą się na filmie przez dowolny czas. Oznacza to, że twarze wykryte w kolejnych klatkach wideo można zidentyfikować jako tę samą osobę.

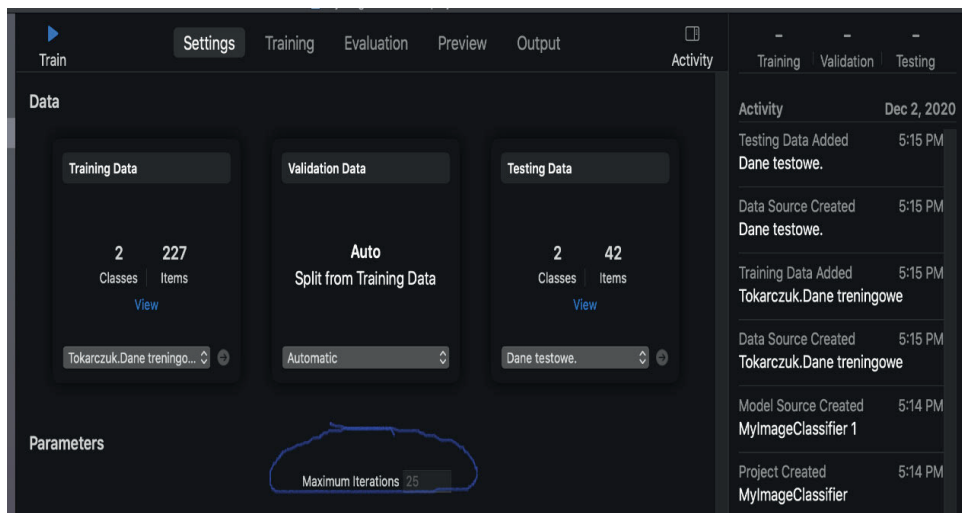
Z kolei przypadku NLP często chcemy wyjść poza zwykłą statystykę w analizie tekstu, jaką oferują programy w rodzaju darmowego Voyanta, i wydobyć coś więcej z takiego tekstu. Narzędzia w rodzaju Microsoft Azure czy AWS umożliwiają (w przypadku języka): rozpoznawanie mowy, rozumienie języka naturalnego oraz syntezę mowy. W kontekście rozpoznawania mowy zaowocowało to powstaniem coraz doskonalszych narzędzi tłumaczenia (Google Translate czy Deep L), asystentów głosowych typu: Cortana, Alexa bądź Siri, czy chatbotów. Jeśli chodzi o syntezę mowy, to powstały systemy przekładające mowę na tekst i tekst na mowę (np. Ivona). W przypadku rozumienia języka naturalnego udoskonalono analizę struktury gramatycznej zdań opartą na analizie syntaktycznej i przede wszystkim na analizie sentymentu. W programie InfraNodus Lab otrzymujemy dodatkowo możliwość wykrycia dominujących tematów w danym tekście czy, powiedzmy, przemówieniu.

Analiza obrazów i tekstu

W tej części zostaną zaprezentowane przykłady analiz obrazów i tekstu – z zastosowaniem algorytmów sztucznej inteligencji, które przeprowadził autor niniejszego artykułu. Zostały tu zastosowane narzędzia i metody opisane powyżej.

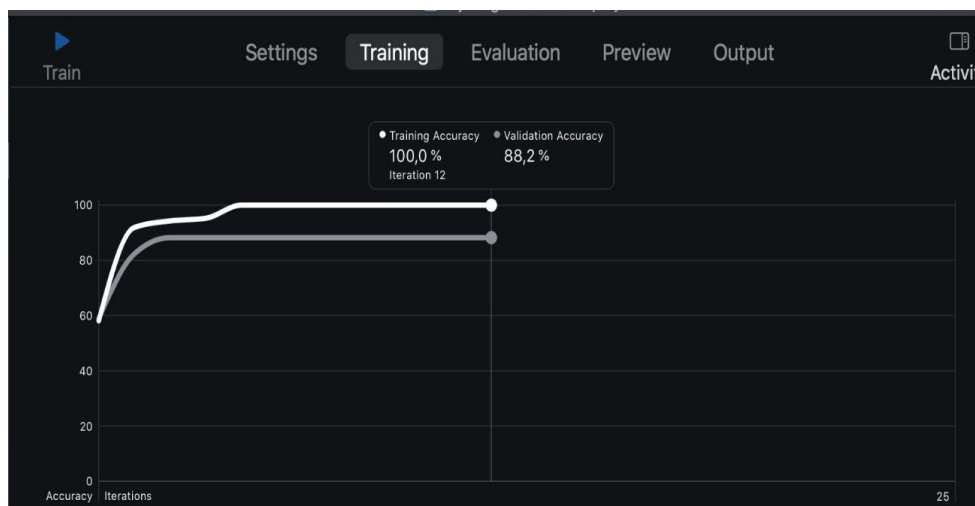
Pierwszy przykład dotyczy rozpoznawania obrazów przez AI. Celem było nauczenie AI w miarę niezawodnego wskazywania postaci ostatniej polskiej noblistki, czyli Olgi Tokarczuk. Jak pokazano w części poprzedniej artykułu, grafika Google potrafi jako Olę Tokarczuk wskazać np. Roberta Lewandowskiego. Aby nauczyć AI rozpoznawać, i do z dużym stopniem prawdopodobieństwa, postać Tokarczuk, wykorzystano środowisko developerskie Create ML. Program ten służy przede wszystkim do pisania programów w języku Swift, ale ma też funkcję trenowania AI.

Aby wytrenować AI, zgromadzono dwa zbiory: treningowy i testowy. Pierwszy zawierał 127 zdjęć Olgi Tokarczuk oraz 100 zdjęć Wisławy Szymborskiej (z tego powodu, że potrzebne było utworzenie dwóch klas). Zbiór testowy zawierał odpowiednio: 20 zdjęć Tokarczuk oraz 22 Szymborskiej (zdjęcie nr 2).



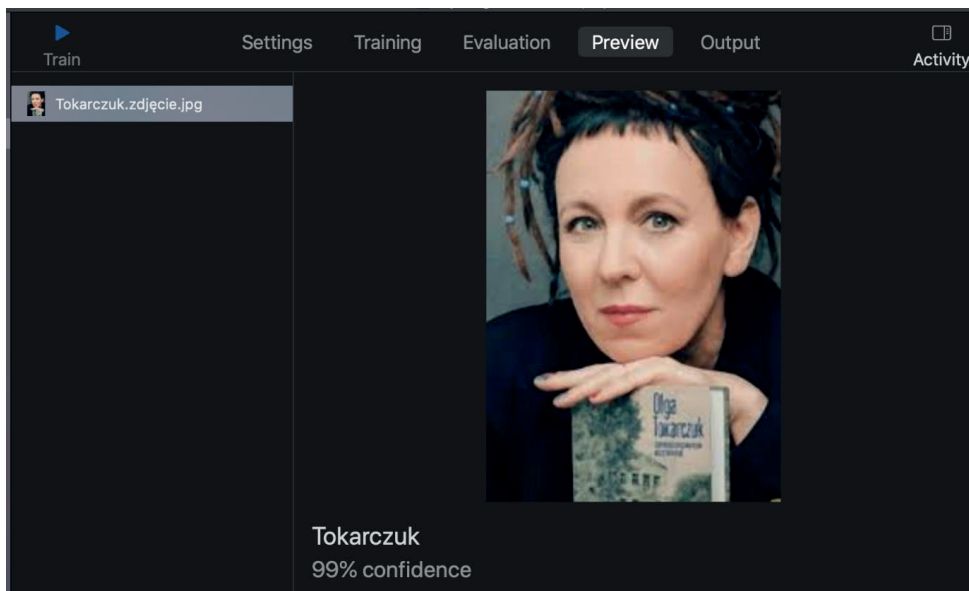
Zdjęcie 2. Zbiory treningowe i testowe w Create ML; oprac. własne

Po przeprowadzeniu sesji treningowej przetestowano skuteczność tej nauki – tj. na ile model był w stanie trafnie rozpoznać zdjęcia Tokarczuk po pokazaniu mu obrazów, które nie były w zbiorze treningowym. Gdyby okazało się, że ta skuteczność jest zbyt mała, to trening należałoby powtórzyć. Wynik ewaluacji był następujący:



Zdjęcie 3. Wynik ewaluacji rozpoznania zdjęć po treningu w Create ML; oprac. własne

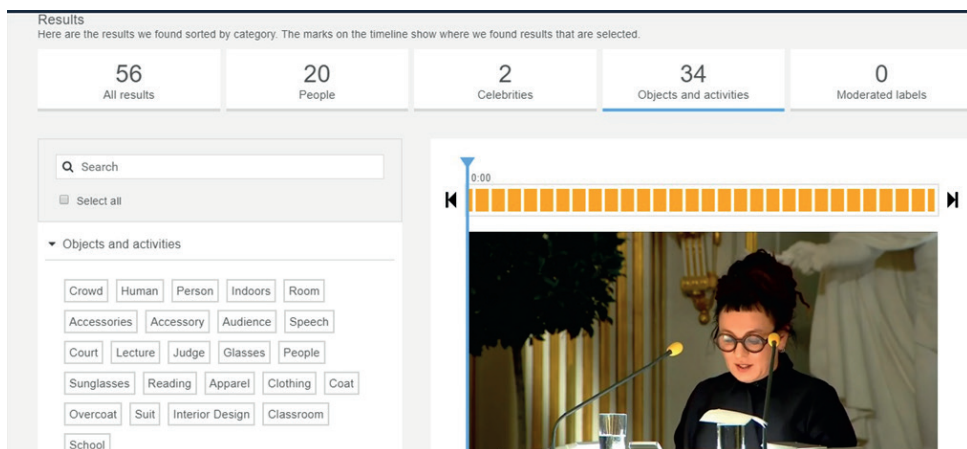
Jak widzimy, model całkiem dobrze poradził sobie z postawionym zadaniem, zatem można go było użyć do rozpoznawania dowolnych zdjęć Olgi Tokarczuk. Model rozpoznawał z dużym stopniem prawdopodobieństwa przekładane mu fotografie (zwykle powyżej 90%) – o czym świadczy kolejny wynik, pokazany na zdjęciu nr 4.



Zdjęcie 4. Wynik z rozpoznania fotografii Olgi Tokarczuk w Create ML; oprac. własne

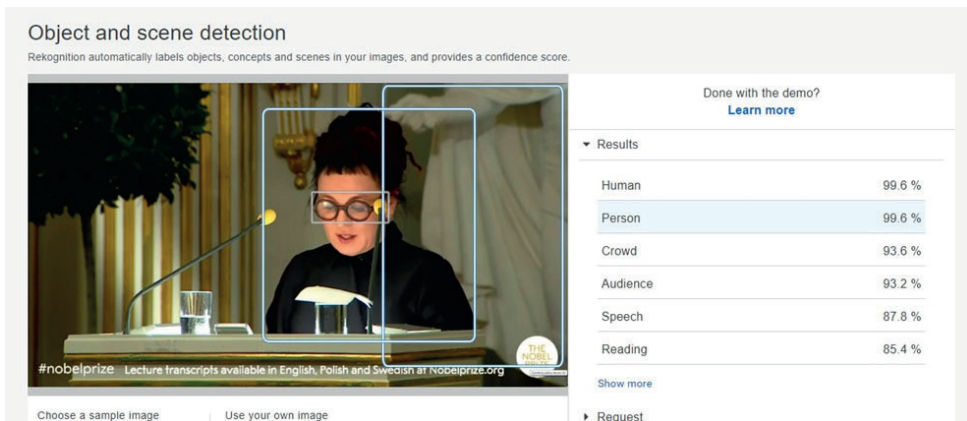
Jak widzimy (zdjęcie nr 4), załadowano do modelu fotografię, której nie było ani w zbiorze treningowym, ani testowym, i model z niemal całkowitą pewnością (99% *confidence*) stwierdził, że jest to Olga Tokarczuk.

Z kolei przedmiotem analizy za pomocą algorytmów sztucznej inteligencji oferowanych przez AWS i AutoML Vision była mowa noblowska Olgi Tokarczuk. Oczywiście mógł to być każdy inny film bądź zdjęcie. Najpierw fragment filmu został załadowany do ML AWS. Program przebadła pierwszą minutę filmu – analizując go klatka po klatce (niebieski kursor). Wynik owej analizy przedstawia się następująco:



Zdjęcie 5. Analiza wykładu noblowskiego Olgi Tokarczuk (Sztokholm, 2019) za pomocą AWS; oprac. własne

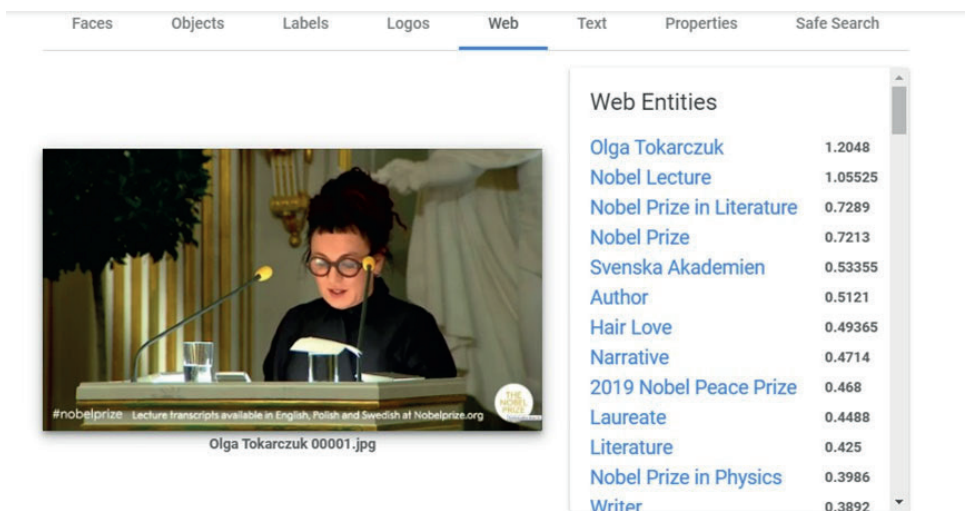
Jak widzimy powyżej, AWS zidentyfikował 56 obiektów – w tym 20 osób. Dwie określił jako celebrytów. Spis wszystkich obiektów widzimy po lewej stronie. Następnie przedmiotem analizy była główna bohaterka i wynik tej analizy przedstawia się następująco:



Zdjęcie 6. Analiza postaci za pomocą AWS; oprac. własne

Program z prawie 100-procentową pewnością wykrył, że na zdjęciu nr 6 mamy do czynienia z człowiekiem, osobą, która przemawia do publiczności, czytając tekst.

Algorytmy AI potrafią jednakże powiedzieć jeszcze więcej o zdjęciu i osobie/osobach na nich się znajdujących. Tym razem zobaczmy, jak to widzi AutoML Google'a. Pozostajemy cały czas przy tym samym zdjęciu. Wynik interpretacji googlowskiej jest następujący:

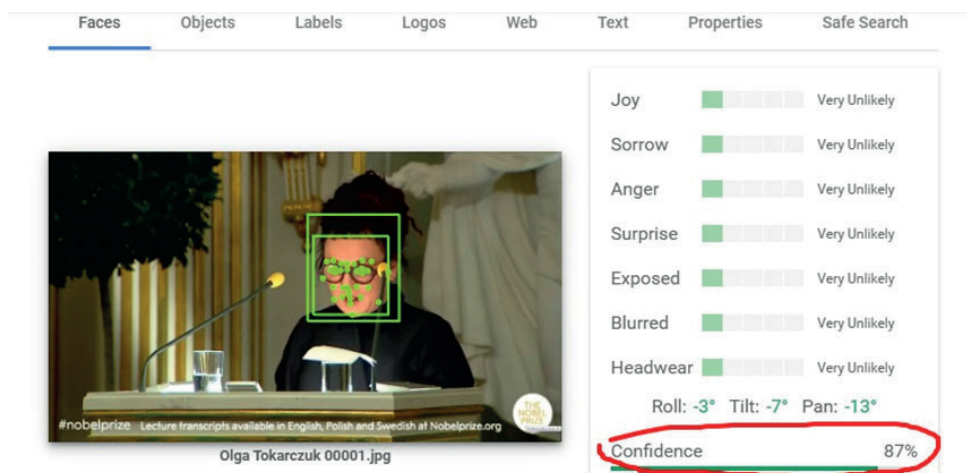


The screenshot shows the Google Image Search interface. The 'Web' tab is selected. On the left, there is a photo of Olga Tokarczuk speaking at a podium. Below the photo, the text reads: '#nobelprize Lecture transcripts available in English, Polish and Swedish at Nobelprize.org' and 'Olga Tokarczuk 00001.jpg'. On the right, a 'Web Entities' panel displays a list of entities and their scores:

Entity	Score
Olga Tokarczuk	1.2048
Nobel Lecture	1.05525
Nobel Prize in Literature	0.7289
Nobel Prize	0.7213
Svenska Akademien	0.53355
Author	0.5121
Hair Love	0.49365
Narrative	0.4714
2019 Nobel Peace Prize	0.468
Laureate	0.4488
Literature	0.425
Nobel Prize in Physics	0.3986
Writer	0.3892

Zdjęcie 7. Analiza mowy Tokarczuk za pomocą AutoML; oprac. własne

Program niemal bezbłędnie wykrył, że to Olga Tokarczuk, która wygłasza mowę noblowską w Sztokholmie w trakcie dorocznej uroczystości. To jednak nie wszystko. Ten sam program potrafi określić stan emocji danej postaci. Zwraca uwagę na takie cechy, jak: radość, smutek, złość itp. – o czym świadczy następny wynik (zdjęcie nr 8):



Zdjęcie 8. Rozpoznanie emocji przez program typu AutoML. Google Cloud Vision; oprac. własne

Widzimy, że określone stany emocjonalne Olgi Tokarczuk zostały zidentyfikowane z pewnością sięgającą 87%. Na marginesie należy dodać, że w przypadku innych fotografii i postaci ten wynik czasami wynosi 100%.

Zaprezentowane powyżej przykłady zastosowania AI do analizy obrazów nie wyczerpują wszystkich wchodzących w grę możliwości, które mogłyby być użyteczne w humanistyce. O to jeszcze inny przykład z zakresu *computer vision*.

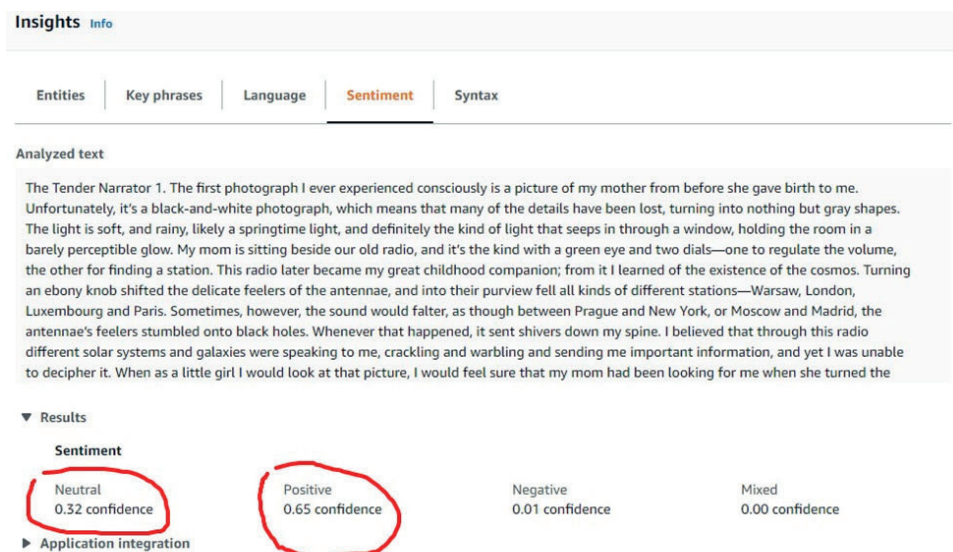
W 2019 r. archeolodzy poinformowali, że odkryli 143 nieznane wcześniej neoglify (przedstawienia ludzi i zwierząt na pustynnych kamieniach) na płaskowyżu Nazca w Peru. Niektóre z konstrukcji były bardzo małe. Najmniejsze z nowo odkrytych mierzą poniżej 5 m. To pokazuje, że poszukiwania nieznanych geoglifów są trudnym zadaniem, szczególnie w połączeniu z ogromną przestrzenią regionu Nazca. Dodatkowo z biegiem czasu warunki atmosferyczne zredukowały wiele linii i figur do lekkich zagłębień w terenie. Dlatego naukowcy skorzystali z pomocy sztucznej inteligencji opracowanej przez inżynierów z IBM.

Algorytmy wykorzystane w badaniach działały na systemie analizy geoprzestrzennej o nazwie IBM PAIRS Geoscope. Sztuczna inteligencja – IBM Watson Machine Learning Accelerator (WMLA) – analizowała ogromne ilości zdjęć satelitarnych oraz pozyskanych z dronów, aby sprawdzić, czy może wykryć jakieś ukryte, pominięte przez ludzi rysunki. System znalazł jedno dopasowanie – geoglif typu B. To wyblakły zarys małej humanoidalnej postaci w postawie stojącej. Pomógł też wizualizować inne, odkryte już przez badaczy rysunki. Ponadto w ostatnich latach mieliśmy przykłady użycia algorytmów AI m.in. do odtwo-

rzenia wyglądu twarzy Mikołaja Kopernika czy rekonstrukcji głosu egipskiego kapłana Nesyamuna z czasów faraona Ramzesa XI.

Wróćmy jeszcze do AWS. Integralną częścią każdego filmu czy utworu multimedialnego jest warstwa audialna – w postaci mowy czy muzyki. Nie inaczej jest w przypadku wystąpienia noblowskiej Olgi Tokarczuk. Jej przemówienie zostało przekonwertowane na tekst i dzięki temu możliwa była analiza sentymentu. Zalicza się ona do NLP. Analizę sentymentu można oczywiście przeprowadzić na każdym tekście – nie tylko wydobytym z utworu filmowego czy nagrania audialnego.

W ramach AWS istnieje usługa Amazon Comprehend (Google i Microsoft Azure oferują podobne rozwiązania), która pozwala ustalić emocjonalny wydźwięk samego przemówienia (oczywiście także każdego tekstu). Nazywa się to właśnie analizą sentymentu. Jako ludzie bowiem mamy zdolność do rozpoznawania cudzych emocji na podstawie treści wypowiedzianych słów, a także pewnych czynników pozawerbalnych, jak ton głosu czy ogólna ekspresja ciała. Ograniczając nasze zainteresowanie do prostszego problemu, tj. rozpoznania jedynie ogólnego stosunku autora do opisywanej treści, a nie konkretnej emocji, możemy starać się taki proces zautomatyzować – tym właśnie zajmuje się analiza sentymentu. Zatem polega ona na określeniu stosunku autora danej wypowiedzi/tekstu do czegoś bądź kogoś. Może on przybierać cztery postaci: pozytywny, negatywny, neutralny i mieszany. W przypadku przemówienia Olgi Tokarczuk otrzymaliśmy wynik następujący:



Zdjęcie 9. Analiza sentymentu mowy noblowskiej Olgi Tokarczuk za pomocą AWS; oprac. własne

Jak widzimy, algorytmy ML wskazują, że istnieje 65% prawdopodobieństwa pozytywnego wydźwięku tej mowy – natomiast neutralny stosunek określony został na 32%. W tym więc przypadku nie jest to jednoznaczna opinia na temat emocjonalnego wymowy tego przekazu.

Jeszcze innym interesującym rozwiązaniem na polu NLP są analizy korpusów tekstualnych. Chodzi tu o ich interpretację statystyczną, stylometrię i znajdowanie relacji wewnątrztekstualnych. Pokażmy to na przykładzie programu InfraNodus. Projekt został stworzony przez Dmitrya Paranyushkina.

Narzędzie (InfraNodus) może być wykorzystywane przez badaczy do porządkowania i lepszego zrozumienia tekstów, do pomiaru poziomu stroniczości, do identyfikacji tych części dyskursu, które zawierają nowe pomysły, oraz do klasycznej statystyki (liczba słów, przeciętna długość zdań itp.). Metoda ta opiera się na algorytmie analizy sieci tekstowej, który reprezentuje dowolny tekst jako sieć i identyfikuje najbardziej wpływowe słowa w dyskursie w oparciu o współwystępowanie terminów. Następnie stosuje się algorytm wykrywania zbiorowości grafów w celu zidentyfikowania różnych klastrów tematycznych, które reprezentują główne tematy w tekście, jak również relacji między nimi. Struktura społeczności jest używana w połączeniu z innymi środkami w celu zidentyfikowania poziomu stroniczości lub zróżnicowania poznawczego dyskursu. Wreszcie luki strukturalne na wykresie mogą wskazywać te części dyskursu, w których brak powiązań, podkreślając tym samym obszary, gdzie istnieje potencjał dla nowych pomysłów. Oto przykład takiej analizy:



Zdjęcie 10. Wizualizacja inauguracyjnego przemówienia prezydenta USA Jimmy'ego Cartera ze stycznia 1977 r. za pomocą InfraNodus; oprac. własne

Jest to analiza i wizualizacja tradycyjnego przemówienia – wygłaszanego przez prezydenta USA obejmującego swój urząd (w tym przypadku przez Cartera). Na zdjęciu numer 10 otrzymaliśmy, oprócz statystki, wizualizację przedstawiającą dokument jako sieć. Wizualizacja zaprezentowanego przemówienia jest oparta na wspomnianej teorii grafów. Jest ona interaktywna. Widzimy na niej węzły dominujące i relacje łączące je z innymi węzłami. Co więcej, możemy zaobserwować, że pewne węzły są na tyle podobne, że tworzą odrębne grupy tematyczne – co jest zaznaczone odmiennymi kolorami. Otrzymane wyniki podają nam: główne pojęcia (podzielone na cztery grupy), najbardziej wpływowe słowa oraz podstawowe zestawienia charakterystyczne dla grafów, typu: gęstość grafu, średni stopień czy modularność. Rysuje nam się także struktura dyskursu analizowanej wypowiedzi – podzielonej na określone tematy i z dominującymi pojęciami, takimi jak: naród, siła, duch i rząd.

Sztuczną inteligencję stosuje się też do badań stylometrycznych. W ich przypadku zakłada się, że świadomie bądź nie, przez całe życie kształtujemy styl pisania charakterystyczny tylko dla nas. Jako czytelnicy czasem również jesteśmy zdolni do rozpoznania konkretnej osoby po danym tekście: jeśli wcześniej mieliśmy do czynienia z odpowiednią ilością tekstu, który wyszedł spod ręki tej osoby, to rozpoznamy często stosowane zwroty, szyk zdania czy sposób używania interpunkcji. W ramach uczenia maszynowego dostarczamy programowi pewne dane wejściowe, na podstawie których porównuje on te elementy w czasie. Musi mieć przestrzeń cech, w której tworzy sobie pewne wektory (czyli uporządkowane zbiory cech). Opisują one profil danej osoby, a kolejne teksty są z tym profilem porównywane. Nazywa się to stylometrycznym profilowaniem behawioralnym. Chodzi o to, że jedni wolą używać krótkich zdań (wówczas siłą rzeczy w tekście pojawia się więcej kropek i wielkich liter, co łatwo wykryć), inni nieco dłuższych. Piszących można też podzielić na tych, którzy wolą rzeczowniki, i na tych, którzy częściej stosują czasowniki itd.

Wnioski

Zastosowanie AI do badań szeroko rozumianej humanistyki jest stosunkowo nowym zagadnieniem. Budzi ono wiele kontrowersji i nieporozumień. Wynika to stąd, że – jak twierdzi czołowy teoretyk literatury i kultury Ryszard Nycz (2017, s. 174) – praca nad tekstem to koronna konkurencja nie tylko literaturoznawczej profesji. Humanisci zawsze pracowali z tekstami i nad tekstami, próbując dekonstruować, rozumieć czy dekonstruować komunikaty, znaczenia, sensy, symbole lub intencje. Tę tendencję umocnił tylko zwrot postmodernistyczny, który skonceptualizował całą rzeczywistość kulturową jako tekstualny świat. Stąd usiłowania

wychodzenia poza tekst czy źródła, a zwłaszcza próby adaptowania metod nauk przyrodniczych czy ścisłych do badań humanistycznych są/były zawsze postrzegane jako matematyzacja, która z reguły budziła sporo nieufności czy wręcz sprzeciwu.

Jednakże w ostatnich kilkudziesięciu latach, a zwłaszcza w XXI w., powstały nowe praktyki związane z upowszechnieniem się technologii cyfrowych. Świat nowych mediów – na czele z Internetem – stał się również przedmiotem zainteresowań humanistów. W dodatku dygitalizacja tzw. wytworów analogowych doprowadziła do transformacji klasycznego przedmiotu badań humanistycznych, czyli tekstów. Jak zauważa Rens Bod, informatyczne, jak go określa, podejście do nauk humanistycznych doprowadziło do odkrycia wielu nowych zjawisk. Bez tego podejścia, twierdzi, niemożliwe byłoby porównywanie setek tekstów naraz albo znajdowanie historycznych schematów w tysiącach zdigitalizowanych źródeł. Cyfrowe nauki humanistyczne przynoszą nie tylko nowe ustalenia, ale także pozwalają na zadawanie nowych pytań, które wcześniej nie byłyby możliwe do postawienia (Bod 2013, s. 465).

Stosowanie algorytmów sztucznej inteligencji wpisuje się w tę tendencję, czyli w podejście informatyczne w humanistyce. Jest ono charakterystyczne dla tzw. humanistyki cyfrowej, która szeroko korzysta z narzędzi ICT⁵, a ostatnio również zaczyna eksperymentować z AI właśnie. Jak widzieliśmy chociażby z zaprezentowanych tu analiz, to w badaniach języka (tradycyjnego przedmiotu humanistyki) sztuczna inteligencja ma największe chyba osiągnięcia. Analizując teksty czy mowę, AI jest w stanie dostrzec różnego typu relacje, korelacje i cechy, których nie odnajdzie nawet badacz z długoletnim doświadczeniem. Jest to tym bardziej istotne, kiedy próbujemy badać duże korpusy tekstualne.

Algorytmy AI i uczenie maszynowe można, jak widzieliśmy, owocnie zastosować do świata obrazów i zjawisk ze sfery medialnej. Metody AI mogą pomóc w analizie zarówno określonych cech fotografii (i innej grafiki), jak i tego, co znajduje się na poszczególnych obrazach czy klatkach filmowych. Jest to analiza zautomatyzowana i nie trzeba dodawać, jakie to ma znaczenie w przypadku badania dużych kolekcji medialnych, określanych mianem *big data*.

Jednym z najbardziej obiecujących kierunków rozwoju AI jest uczenie maszynowe nienadzorowane. W przypadku badań humanistycznych czy z zakresu nauk społecznych algorytmy AI nie będą znały wyniku (na wejściu), lecz w kolekcji danych zmuszone będą szukać i klasyfikować szeroko rozumiane obiekty. Przykładem tego typu analiz jest program InfraNodus, który służy do interpretacji różnych korpusów tekstualnych. Jego algorytmy, oparte na grafach, potrafią skonceptualizować (jak widzieliśmy) każdy tekst jako sieć, wykryć dominujące tematy

⁵ Skrót od *information and communication technology*.

oraz główne kategorie zbiorów cyfrowych bądź analogowych przekonwertowanych do postaci cyfrowej.

I uwaga na koniec. Uczni chcący w przyszłości w pełni wykorzystywać potencjał AI do nauk humanistycznych zmuszeni będą do nabycia zupełnie nowych kompetencji z dziedziny IT – przede wszystkim metod Data Science: podstaw statystyki, algebry, programowania (zwłaszcza w Python i R), analizy danych oraz ich wizualizacji, i oczywiście do opanowania bibliotek programistycznych przeznaczonych dla AI.

Bibliografia

- Bod R., 2013, *Historia humanistyki. Zapomniane nauki*, Warszawa: Aletheia.
- Chollet F., 2018, *Deep learning*, London: Manning Publications.
- Flasiński M., 2018, *Wstęp do sztucznej inteligencji*, Warszawa: PWN.
- Aurelian G., 2018, *Uczenie maszynowe z użyciem Scikit-Learn i TensorFlow*, Gliwice: Helion.
- Kaplan J., 2019, *Sztuczna inteligencja*, Warszawa: PWN.
- Lee H., Sohn I., 2016, *Big data w przemyśle*, tłum. M. Barabowski, Warszawa: PWN.
- Kisielewicz A., 2017, *Sztuczna inteligencja i logika*, Warszawa: WNT, wyd. II.
- Mayer-Schonberger V., Cukier K., 2014, *Big data. Rewolucja, która zmieni nasze myślenie, pracę i życie*, Warszawa: Mt Biznes.
- Nycz R., 2017, *Kultura jako czasownik. Sondowanie nowej humanistyki*, Warszawa: IBL.
- Osigna D., 2019, *Deep learning. Receptury*, Gliwice: Helion.
- Patrterson J., Gibson A., *Deep learning. A Practitioner's Approach*, O'Reilly Media, <http://csis.pace.edu/~ctappert/cs855-18fall/DeepLearningPractitionersApproach.pdf> (dostęp: 10.01.2020).
- Szpunar M., 2018, „Kultura algorytmów”, *Zarządzanie w Kulturze* 1: 1–10.
- Stephenson D., 2018, *Big Data Demystified. How to use Big Data, Data Science and AI to Make Better Business and Gain Competitive Advantage*, London: FT Press.