

FILIP WYSZYŃSKI

ORCID: 0000-0002-0792-5730

Uniwersytet w Białymstoku

CZY PRAWO POWINNO BYĆ NEUTRALNE MORALNIE?  
UWAGI NA TEMAT UNORMOWANIA SZTUCZNEJ  
INTELIGENCJI

Abstrakt: W niniejszym artykule omówiona została problematyka uregulowania sztucznej inteligencji. Wskazano na potrzebę uregulowania sztucznej inteligencji. Autor postuluje dla algorytmów AI prawo neutralne moralnie. Propozycja uregulowania sztucznej inteligencji w sposób neutralny moralnie opiera się na teorii filozofii prawa Axela Hägerströma i Ronalda Dworkina. Chodzi o niezależność AI od instrumentów władzy państwowej. Wnioskiem opracowania jest konieczność uregulowania sztucznej inteligencji ze względu na zabezpieczenie człowieka przed czynnikiem ludzkim i przed AI.

Słowa kluczowe: sztuczna inteligencja, moralność, neutralność

## UWAGI WPROWADZAJĄCE

W artykule na temat moralności sztucznej inteligencji<sup>1</sup> (*artificial intelligence*, AI) Aimee van Wynsberghe i Scott Robbins przytaczają interesujący przykład — sytuację etyczną psa-opiekuna osoby starszej<sup>2</sup>. Jest to jedynie krótko wzmiankowany przykład, zachęca jednak do szerszego omówienia, a wręcz zbudowania na jego kanwie daleko idącej koncepcji filozoficzno-prawnej.

Pies, który został przeszkolony do bycia opiekunem osoby starszej, wypełnia swoje obowiązki, postępując — z punktu widzenia człowieka — moralnie. Nie potrzebuje jednak aksjologicznej analizy. Autorzy odwołują się do tego przykładu w celu rozróżnienia sytuacji „bycia w odpowiedzialności moralnej”, a sytuacji „delegowania do pełnienia roli moralnej”. Warto jednak pójść o krok dalej i po-

<sup>1</sup> Ogólnie o moralności AI zob. między innymi J. Savulescu i H. Maslen, *Moral Enhancement and Artificial Intelligence: Moral AI?*, [w:] *Beyond Artificial Intelligence: The Disappearing Human-Machine Divide*, red. J. Romportl, E. Zackova, J. Kelemen, Cham 2015, s. 79–80.

<sup>2</sup> A. van Wynsberghe, S. Robbins, *Critiquing the Reasons for Making Artificial Moral Agents*, „Science and Engineering Ethics” 2019, nr 25, s. 724.

stulować<sup>3</sup>, by sytuacja prawna AI pod względem moralności została unormowana właśnie według „delegowania do pełnienia roli moralnej”.

Należy powiedzieć postulatywnie, że algorytmy AI — na wzór tresury, której poddawane są psy-opiekunowie — powinny być szkolone do takiego działania moralnego, by dalsze działanie algorytmów nie odbiegało od narzuconych norm, a jednocześnie nie wymagało już dalszej kontroli<sup>4</sup>. Wyćwiczona intuicja miałaby stanowić niezbędną cechę dla efektywnego funkcjonowania AI. Prawem intuicyjnym nazwać można w ślad za Marią Szyszkowską przeżycia prawne niezależne od zewnętrznego autorytetu<sup>5</sup>, a zatem autonomiczne<sup>6</sup>. Autonomia AI rozumiana jako odseparowanie od człowieka może zagwarantować ludziom większe bezpieczeństwo.

### UWAGI PORZĄDKUJĄCE

Niniejsze opracowanie należy podzielić na dwie części. W pierwszej wykazane zostaną zagrożenia i niekonsekwencje związane z prawem „dodatnim” moralnie względem korzystania z AI. Analizy tej można dokonać na przykładzie strategii szwedzkiego rządu w związku ze zwalczaniem pandemii COVID-19. Natomiast w drugiej części zostanie poddana analizie kwestia „uczenia” pierwotnych norm moralnych, by dalsze i autonomiczne postępowanie algorytmów AI odznaczało się nawyknieniem do moralności. Podobnie jak w tresurze psów opiekunów nauczać owych norm należy za pomocą utrwalania schematów, tak by normy te zostały zinternalizowane. Z tego powodu kształtując prawo AI, należałoby usunąć kwestie nieostre. Jest to związane również z samą istotą algorytmów. Prawo takie powinno być zero-jedynkowe, binarne, schematyczne, leksometryczne. Dlatego też należy postulować maksymalną ilość reguł (*rules*), przy minimalnej ilości zasad (roz-

<sup>3</sup> Istotną koncepcją jest na przykład postulat uregulowania AI w oparciu o tak zwanych „sztucznych agentów”. Odnosi się on jednak do aspektów technicznych albo do konkretnych gałęzi prawa (cywilnego, karnego), podczas gdy w niniejszej pracy chodzi o charakter regulacji pod względem moralności, a więc o aspekt filozoficzno-prawny. Zob. A. Krasuski, *Status prawny sztucznego agenta. Podstawy prawne zastosowania sztucznej inteligencji*, Warszawa 2020; J. Casas-Roma, J. Arnedo-Moreno, *From Games to Moral Agents: Towards a Model for Moral Actions*, [w:] *Artificial Intelligence Research and Development: Proceedings of the 22<sup>nd</sup> International Conference of the Catalan Association for Artificial Intelligence*, red. J. Sabater-Mir et al., Amsterdam 2019, s. 19–20.

<sup>4</sup> „Tresurę” należy rozumieć jako nadanie wzorca postępowania i odstąpienie od dalszych kontroli, a nie w sposób „panoptyczny”, jak to przedstawił Michel Foucault, upatrując źródeł współczesnego wyrabiania społecznego nawyku przestrzegania norm („tresowania”) w inżynierii o charakterze bezustannej presji wywieranej przez obserwację i system kar. Zob. M. Foucault, *Nadzorować i karać. Narodziny więzienia*, przeł. T. Komendant, Warszawa 2009, s. 126.

<sup>5</sup> M. Szyszkowska, *Zarys filozofii prawa*, Białystok 2000, s. 207.

<sup>6</sup> *Ibidem*.

różnienie w myśl filozofii Ronalda Dworkina)<sup>7</sup>. Kwestie te można prześledzić na przykładzie autonomicznych pojazdów, których działanie oparte jest o AI.

Pierwsza część artykułu ma zatem na celu dowiedzenie, że istnieje potrzeba przekazania kompetencji opiekuńczych wobec człowieka algorytmom AI („większa wierność” algorytmu względem człowieka niż człowieka względem człowieka), zaś rozważania pomieszczone w drugiej części mają nakreślić faktyczną możliwość zaimplementowania tegoż postulatu („tresura”, przyuczenie).

Dodać należy, że do określenia pojęcia prawa neutralnego moralnie przyjęte zostały wskazówki definicyjne z publikacji Michała Błachuta<sup>8</sup>. Autor jako cechę prawa neutralnego moralnie wskazuje między innymi odrzucenie uprzywilejowania prawodawczego państwa<sup>9</sup>, zaakcentowanie jego roli jako bezstronnego arbitra<sup>10</sup>, a także zapewnienie równych szans różnorodnym koncepcjom życia<sup>11</sup>. Zauważyć trzeba, że algorytmy AI wykazują podatność na spełnienie takich przesłanek. Choć AI wykazuje zdolność samoregulacyjną<sup>12</sup>, nie byłoby pożądane oddanie jej pełnej autonomii. Za zasadne można byłoby jednak uznać oddanie pola algorytmom AI w kwestii bezpieczeństwa i neutralności w oparciu o podstawowe reguły. AI mogłaby stanowić skuteczniejsze i niedyskryminacyjne narzędzie opieki nad osobami starszymi i chorymi, a także umocnić bezpieczeństwo w ruchu drogowym.

## SZTUCZNA INTELIGENCJA VS. RZĄD SZWEDZKI

Warto wspomnieć, że AI tworzy systemy zdolne do szybkiego i zdalnego badania chorób zakaźnych, w tym COVID-19<sup>13</sup>. Algorytmy AI opierają się przy tym na danych demograficznych. Ów projekt analizuje na zbiorach danych (*big data*) przypadki pacjentów i wykorzystuje uczenie maszynowe (*machine learning*) do diagnostyki. Biorąc pod uwagę COVID-19, precyzyjność diagnostyczna osiągnęła we wspomnianym programie próg aż 80%. Innym przykładem może

<sup>7</sup> B. Greczner, *Precedens a spójność aksjologiczna prawa w ujęciu Ronalda Dworkina*, „Przegląd Prawa i Administracji” 2010, nr 18, s. 17.

<sup>8</sup> M. Błachut, *Postulat neutralności moralnej prawa a konstytucyjna zasada równość*, Wrocław 2005.

<sup>9</sup> *Ibidem*, s. 28.

<sup>10</sup> *Ibidem*, s. 38.

<sup>11</sup> Tak między innymi M. Dudek, odnosząc się do publikacji M. Błachuta. Zob. M. Dudek, *Autonomia, neutralność i indyferentność moralna prawa a jego społeczeństwo*, „Ruch Prawniczy, Ekonomiczny i Socjologiczny” 2014, nr 4, s. 77.

<sup>12</sup> A. Rothert, *Nowa biologia polityczna*, [w:] *Metafory polityki*, t. 3, red. B. Kaczmarek, Warszawa 2005, s. 298.

<sup>13</sup> Komisja Europejska. Cordis, *Sztuczna inteligencja przejmie laboratoryjne badania na obecność patogenów*, <https://cordis.europa.eu/article/id/421710-artificial-intelligence-overcomes-laboratory-testing-for-pathogen-detection/pl> (dostęp: 10.01.2021).

być system identyfikacyjny, gdzie algorytmy AI prześwietlają organizm pacjenta pod kątem zmian koronawirusowych przy użyciu zdjęć rentgenowskich<sup>14</sup>. Jak podkreśla F. Herrera, skuteczność tej techniki sięga 80%. Technologie i techniki związane z AI wydają się zatem nieodzownym sprzymierzeńcem człowieka w walce z pandemią.

Dla porządku wskazać trzeba, że istnieje także wiele stanowisk dostrzegających w AI zagrożenie. Jak przykładowo wskazuje Włodzimierz Fehler, niebezpieczne jest powstanie tak zwanej „superinteligencji”, czyli osiągnięcie takiego stanu, w którym AI będzie zdolna do stałego i samodzielnego udoskonalania się<sup>15</sup>. Na problem ten zwracali uwagę Bill Gates, Elon Musk czy Stephen Hawking<sup>16</sup>. Fehler przytacza Muska, który podkreśla, że trudne do przewidzenia są rozmiary zagrożeń, gdyby AI została użyta w celach agresywnych<sup>17</sup>. Z kolei B. Joy zauważa trzy główne obszary zagrożenia — genetykę, nanotechnologię i robotykę<sup>18</sup>. Na innego rodzaju ryzyko zwraca uwagę Ryszard Tadeusiewicz, akcentując kwestię robotów militarnych i autonomicznych dronów bojowych<sup>19</sup>. Autor obawia się się ponadto zagrożeń związanych z wyparciem ludzi z rynku pracy przez algorytmy i maszyny, czego skutkiem byłoby dotkliwie bezrobocie<sup>20</sup>.

Odkładając na bok potencjalne zagrożenia związane z AI, należy podkreślić, że w chwili obecnej stała się ona stronnikiem społeczeństwa i realną pomocą w walce z koronawirusem. Inaczej zaś postrzegane jest działanie rządu Szwecji w walce z pandemią, nieoczekiwanie nieprzystające do ugruntowanych wartości.

Rząd premiera Stefana Löfvena objął drogę inną niż reszta Europy. Mówi się niekiedy, że był to „szwedzki eksperyment”<sup>21</sup>. Sytuacja była zresztą dostrzegalna w przestrzeni publicznej, gdzie noszenie maski ochronnej należało do rzadkości. Szwecja przez długi czas wydawała zalecenia przeciwne do krajów takich, jak Dania, Finlandia czy Norwegia, gdzie obowiązek zakrywania nosa i ust w przestrzeni publicznej był sankcjonowanym nakazem.

Chcąc nakreślić krótką i niezbędną dla porządku wywodu charakterystykę państwa szwedzkiego, trzeba zauważyć, że Szwecja określana jest jako państwo dobrobytu, który został zapoczątkowany już w latach pięćdziesiątych XX wie-

---

<sup>14</sup> Radio Białystok, *Sztuczna inteligencja pomaga identyfikować COVID-19 w płucach*, <https://www.radio.bialystok.pl/koronawirus/index/id/193618> (dostęp: 10.01.2021).

<sup>15</sup> W. Fehler, *Sztuczna inteligencja — szansa czy zagrożenie?*, „Studia Bobolanum” 2017, nr 3, s. 79.

<sup>16</sup> *Ibidem*.

<sup>17</sup> *Ibidem*.

<sup>18</sup> *Ibidem*, s. 80.

<sup>19</sup> R. Tadeusiewicz, *Automatyzacja i sztuczna inteligencja jako źródła prawdziwych i wyimaginowanych zagrożeń*, [w:] *Czy świat należy urządzić inaczej. Schyłek i początek*, red. B. Galwas, P. Kozłowski, K. Prandecki, Warszawa 2019, s. 38.

<sup>20</sup> *Ibidem*, s. 39.

<sup>21</sup> Deutsche Welle, *Koronawirus w Szwecji. Uboczne skutki eksperymentu*, <https://www.dw.com/pl/koronawirus-w-szwecji-uboczne-skutki-eksperymentu/a-55544999> (dostęp: 10.01.2021).

ku<sup>22</sup>. Polityka wyrównywania płac, pełnego zatrudnienia, niskiej inflacji i stałego wzrostu gospodarczego miała zapewnić krajowi ogólny dobrobyt<sup>23</sup>. Szwecja jest pierwowzorem socjaldemokratycznego państwa dobrobytu<sup>24</sup> budowanego na modłę uniwersalistyczną, wedle której rozbudowanymi prawami socjalnymi mają zostać objęci wszyscy obywatele<sup>25</sup>. Państwo szwedzkie odznacza się szczególną dbałością o prawa człowieka<sup>26</sup>. Za niewątpliwą cechę szwedzkiego państwa opiekuńczego uznaje się troskę o seniorów. Jak opisuje Rafał Bakalarczyk, szwedzkie społeczeństwo odznacza się wyjątkowo wysokim wskaźnikiem długości życia<sup>27</sup>. Badania z roku 2007 pokazują, że średnia wieku w Szwecji dla kobiet i mężczyzn wynosi odpowiednio 75 i 72 lata, zaś dla porównania liczby te wynoszą w Polsce odpowiednio 70 i 64 lata<sup>28</sup>. Bakalarczyk wskazuje, że wysoka średnia przeżywalności w dobrym zdrowiu w Szwecji jest nie tylko skutkiem sprawnego i profesjonalnego funkcjonowania służby zdrowia, ale także zasługą pomocy socjalnej<sup>29, 30</sup>.

Ogólnie rzecz biorąc, szwedzki model opieki nad ludźmi starszymi jest instytucjonalny. Państwo pełni niejako rolę rodziny, wykonując obowiązki opiekuńcze. W Szwecji troska o ludzi starszych uwzględnia wyjątkowo ważną w tym kraju kwestię autonomii człowieka. Autonomia ta przejawia się poprzez rozwijanie własnej osobowości i zaspokajanie potrzeb. Aktywność seniorów znajduje się w tym kontekście na wysokim poziomie partycypacyjnym, jeżeli chodzi o udział w działaniach lokalnej społeczności<sup>31</sup>. Podkreśla się, że w Szwecji opieka stacjonarna nad seniorami postrzegana jest jako niewystarczająca. W konsekwencji nieodzwonne jest również możliwie jak najdłuższe funkcjonowanie osób starszych w społecznościach<sup>32</sup>. W Szwecji opieka nad seniorami jest silnie zakorzeniona w systemie usług publicznych, a państwo szwedzkie, przeznaczające znaczne środki na opiekę, gwarantuje wysoką dostępność świadczeń<sup>33</sup>.

Trzeba powiedzieć, że państwo szwedzkie przejęło odpowiedzialność za zdrowie i życie najstarszych obywateli, jednak podczas pandemii okazało się, że

<sup>22</sup> M. Banaś, *Szwedzka polityka integracyjna wobec imigrantów*, Kraków 2010, s. 131.

<sup>23</sup> *Ibidem*.

<sup>24</sup> A. Hägerström wywarł wielki wpływ na ukształtowanie się szwedzkiej idei państwa socjaldemokratycznego, które nazwano „socjalizmem funkcjonalnym”. Zob. S. Eliaeson, *Axel Hägerström and Modern Social Thought*, „Zeitschrift für Politik, Wirtschaft und Kultur” 2000, nr 1, s. 19.

<sup>25</sup> A. Jachowicz, *Funkcjonowanie pomocy społecznej. Wybrane problemy*, Dąbrowa Górnicza 2011, s. 66.

<sup>26</sup> Kancelaria Rządu Szwecji, *Szwedzki system rządów*, Sztokholm 2014, s. 11–13.

<sup>27</sup> R. Bakalarczyk, *Opieka nad seniorami w państwie opiekuńczym — przykład Szwecji*, „Problemy Polityki Społecznej. Studia i Dyskusje” 2012, nr 18, s. 110.

<sup>28</sup> *Ibidem*.

<sup>29</sup> *Ibidem*.

<sup>30</sup> I. Skoog, *COVID-19 and Mental Health Among Older People in Sweden*, „International Psychogeriatrics” 2020, nr 32, s. 1173–1175.

<sup>31</sup> *Ibidem*, s. 112.

<sup>32</sup> *Ibidem*, s. 114.

<sup>33</sup> *Ibidem*, s. 116.

sytuacja osób starszych, mieszkających nierzadko w domach opieki społecznej, nie była istotnym argumentem w dyskusji na temat strategii rządowej w sprawie zapobiegania i zwalczania COVID-19, pomimo że to właśnie domy pomocy społecznej stawały się raz po raz głównymi źródłami zakażeń w skali kraju.

Premier Szwecji stwierdził, że strategia walki jego kraju z pandemią była błędna, a szwedzkie oddziały intensywnej opieki zdrowotnej znalazły się na granicy wydolności<sup>34</sup>. Jak relacjonowała Miłada Jędrysik, w kwietniu 2020 roku jedna trzecia zgonów dotyczyła szwedzkich domów opieki<sup>35</sup>, w tym jednym aspekcie (to jest opieki nad seniorami w czasie pandemii) szwedzki model okazał się nieskuteczny<sup>36</sup>.

Biorąc pod uwagę powyższe rozważania, stwierdzić należy, że AI, która miała być zagrożeniem, okazała się zabezpieczać człowieka przed jego ryzykowną niedoskonałością, zaś poważne wątpliwości moralne wzbudzać może, wydawałoby się, wzorcowy model zbudowany przez człowieka i mający służyć człowiekowi — model państwa opiekuńczego, który w chwili pandemii nie okazał się wystarczająco wrażliwy społecznie.

Przykład ten pokazuje, że stosowanie norm moralnych przez człowieka szwankuje, podczas gdy w analogicznej sytuacji AI działa w sposób nie budzący moralnych wątpliwości. AI gwarantuje bowiem znaczną eliminację ryzyka podejmowania złych wyborów. Ludzkie normy moralne są w gruncie rzeczy niejasne i nieostre. Czytelnym jest ponadto, że człowiek w wielu sytuacjach nie będzie wykazywał przywiązania do tychże norm.

Odnosząc się ściśle do filozofii prawa, warto spojrzeć na kwestię szwedzkiej strategii oraz wykorzystania AI przez pryzmat skandynawskiego realizmu<sup>37</sup>, czyli nurtu, którego teoretykiem był szwedzki filozof Axel Hägerström<sup>38</sup>. Hägerström wygłosił między innymi trzy zajmujące tezy: według tezy ontologicznej fakty moralne nie istnieją; wedle tezy epistemologicznej wiedza na temat wartości jest niemożliwa; z kolei w ślad za tezą semantyczną można stwierdzić, że osądy moralne są wyzute z tak zwanej „wartości logicznej”<sup>39</sup>. Powyższy koncept wydaje

---

<sup>34</sup> A. Payne, *Premier Szwecji: nasza strategia walki z pandemią była błędna*, <https://businessinsider.com.pl/wiadomosci/koronawirus-w-szwecji-szwedzka-droga-premier-nasza-strategia-byla-bledna/e9glf38> (dostęp: 10.01.2021).

<sup>35</sup> M. Jędrysik, *Szwedzki sen: czy jest sens walczyć z koronawirusem tak, jak robiła to Szwecja? A raczej — czy był?*, <https://oko.press/czy-jest-sens-walczyz-z-koronawirusem-tak-jak-szwecja/> (dostęp: 10.01.2021).

<sup>36</sup> *Koronawirus w Szwecji: Czy szwedzki model walki z pandemią sprawdza się?*, „Dziennik Gazeta Prawna”, <https://www.gazetaprawna.pl/wiadomosci/artykuly/1470155,koronawirus-w-szwecji-szwedzki-model-walki-z-pandemia.html> (dostęp: 10.01.2021).

<sup>37</sup> H.L.A. Hart, *Scandinavian Realism*, „The Cambridge Law Journal” 1959, nr 17, s. 233–240.

<sup>38</sup> J. Bjarup, *The Philosophy of Scandinavian Legal Realism*, „Ratio Juris” 2005, nr 18, s. 2–3.

<sup>39</sup> M. Wojciechowski, *Axel Hägerström*, [w:] *Filozofia prawa w pytaniach i odpowiedziach*, red. J. Zajadło, K. Zeidler, Warszawa 2013, s. 84.

się silnie pragmatyczny i bywa określany mianem *value nihilism* (*värdenihilism*), czy też *axiological nihilism*<sup>40</sup>.

Analizując inne aspekty koncepcji Hägerströma, trzeba zaakcentować to, że ów filozof prawa był przekonany, iż poczucie sprawiedliwości jest intuicyjne<sup>41</sup>. Jego zdaniem człowiek wchodzi w pole oddziaływania prawa, przez co poddany jest działaniu sugestywnemu. Sugestywność ta opiera się na tym, że człowiek posiada świadomość prawną już przez sam fakt, że prawo jest stosowane<sup>42</sup>. Według Hägerströma przestrzeganie prawa jest odzwierciedleniem pewnych nawyków i preferencji<sup>43</sup>. Hägerström podążał ku ujęciu empirycznemu, które traktuje zjawisko przestrzegania prawa jako reakcję psychologiczną<sup>44</sup> związaną z językiem, która utrwała się z biegiem lat. Język ten jest ocenny, a konkretne zdania (*value sentences*) dają się podzielić w podgrupy<sup>45</sup>. W tym sensie uregulowaniu prawa AI można byłoby, podążając za refleksjami Hägerströma, przydać charakter prawa wykryzalizowanego z normatywnych przyzwyczajzeń, które zostałyby zinternalizowane przez algorytmy za pomocą języka (języka oprogramowania).

Trzeba jednakże założyć, że na obecnym etapie rozwoju AI nie jest pożądane, by algorytmy wytworzyły samodzielne systemy moralne *ab initio*. Nieodzownym elementem wydaje się być pewien rodzaj *spiritus movens* („tresury”). Mrzonką wydaje się być myślenie, że udało się stworzyć nowy system moralny w pełnym oderwaniu od wypracowanych już w społeczeństwie zasad moralnych, a przynajmniej stwierdzić trzeba, że mogłoby to być niebezpieczne, a zagrożenia z tym związane sygnalizowano wcześniej. Przyjąć bowiem należy za Dawidem Bunikowskim, że niemożliwa jest całkowita neutralność moralna prawa (bez źródeł, punktu wyjścia, początku)<sup>46</sup>. Chodzi zatem o to, by inkorporować normy moralne bez inkorporacji ich ontologii, czyli ich ontologicznego uzasadnienia<sup>47</sup>. Należy zatem jedynie nadać początek systemowi moralnemu AI przez człowieka. Proces taki można byłoby przyrównać do koncepcji prawa jako mechanizmu zegara, którego twórcą jest Hägerström<sup>48</sup>.

Hägerström opisuje prawo jako zegar, w którym kołami zębatymi mechanizmu stają się podmioty prawa, zaś siłą, która pozwala na poruszanie wskazówek

<sup>40</sup> H. Ruin, *Hägerström, Nietzsche and Swedish Nihilism*, [w:] *Axel Hägerström and Modern Social Thought*, red. S. Eliaeson, P. Mindus, S.P. Turner, Oxford 2014, s. 177.

<sup>41</sup> *Ibidem*, s. 86.

<sup>42</sup> *Ibidem*.

<sup>43</sup> J. Oniszczyk, *Filozofia i teoria prawa*, Warszawa 2012, s. 447.

<sup>44</sup> E. Cassirer, *Axel Hägerström. Eine Studie zur Schwedischen Philosophie der Gegenwart*, Göteborg 1939, s. 113.

<sup>45</sup> F. Tersman, *Methodological Reflections on Hägerström's Meta-ethics*, <http://www.diva-portal.se/smash/get/diva2:713073/FULLTEXT02.pdf> (dostęp: 10.01.2021), s. 2.

<sup>46</sup> D. Bunikowski, *Idea neutralności moralnej prawa we współczesnych systemach prawnych*, [w:] *Etyka. Część 2*, red. S. Janeczek, A. Starościc, Lublin 2016, s. 547.

<sup>47</sup> *Ibidem*.

<sup>48</sup> P. Mindus, *Real Mind. The Life and Work of Axel Hägerström*, Dordrecht 2009, s. 137–138.

jest sama rzeczywistość, czy też — mówiąc Émile'em Durkheimem — szereg faktów społecznych<sup>49</sup> takich jak przede wszystkim poczucie powinności czynienia oraz obawa przed konsekwencjami nieczynienia. Podmioty prawa stają się *per se* jego twórcami. Ta ilustratywna koncepcja może być przydatna do kontynuowania wyводу o potencjalnym unormowaniu AI. Otóż po nastawieniu wskazówek prawnych zegara na wielkości stosowne do norm moralnych, system taki może stać się w konwencji swojego działania samowystarczalny. Zgodne z moralnością działanie AI odbywałoby się bowiem według narzuconych reguł. Toczyłoby się jednak w sposób niezaburzony i niezależny od zewnętrznej ingerencji, analogicznie do sytuacji psa opiekuna, który po odpowiednim przeszkoleniu nie wymaga ani ganiaenia, ani nadzoru.

## REGUŁY VS. ZASADY

Prawo neutralne moralnie służyć może także bezpieczeństwu człowieka w przypadku zastosowania AI przy programowaniu autonomicznych pojazdów<sup>50</sup>. Jasnym jest, że autonomizacja pojazdów jest bezpieczna wówczas, gdy wszystkie one zostaną zaprzęgnięte do algorytmicznego oprogramowania<sup>51</sup>. Systemy takie mogą jednakże działać tylko w oparciu o normy zero-jedynkowe<sup>52</sup> — jednoznaczne, wyuczone, „łatwe”, a mówiąc filozoficzno-prawnie: wyłącznie w oparciu o reguły, a nie standardy jako zasady i wymogi polityki.

Interesujące refleksje przedstawia Grzegorz Szulczewski, który zadaje pytanie o to, czy AI nabędzie zdolność do refleksji moralnej w zakresie podejmowania słusznego wyboru<sup>53</sup>. Autor odnosi się do tak zwanego „dylematu wagoni-

<sup>49</sup> W. Doroszewski, *Studia i szkice językoznawcze*, Warszawa 1962, s. 96; E. Tarkowska, *Ciągłość i zmiana socjologii francuskiej: Durkheim, Mauss, Lévi-Strauss*, Warszawa 1974, s. 139; A. Pałubicka, *Kulturowy wymiar ludzkiego świata obiektywnego*, Poznań 1990, s. 123; A. Lipski, *Perspektywy socjologii kultury artystycznej*, Warszawa 1992, s. 147; M. Bernasiewicz, *Interakcjonizm symboliczny w teorii i praktyce resocjalizacyjnej*, Kraków 2011, s. 146.

<sup>50</sup> A. Chłopecka, *Problematyka odpowiedzialności za ruch autonomicznych samochodów w kontekście ochrony praw człowieka*, „Człowiek w Cyberprzestrzeni” 2018, nr 1, s. 32–33. Ogólne rozważania na temat uregulowania autonomicznych pojazdów podejmuje w piśmiennictwie prawniczym między innymi Aleksander Chłopecki. Zob. A. Chłopecki, *Sztuczna inteligencja — szkice prawnicze i futurologiczne*, Warszawa 2018, s. 57–60.

<sup>51</sup> T. Neumann, *Perspektywy wykorzystania pojazdów autonomicznych w transporcie drogowym w Polsce*, „Autobusy” 2018, nr 12, s. 791.

<sup>52</sup> Przyszłość AI związana jest z komputerami kwantowymi, które mogą działać na wszystkich wartościach w tym samym czasie („splątanie”), jednak działanie to opiera się koniec końców na systemie zer i jedynek. Zob. M. Sawerwain, J. Wiśniewska, *Informatyka kwantowa. Wybrane obwody i algorytmy*, Warszawa 2015, s. 99; M. Hetmański, *Umysł a maszyny: krytyka obliczeniowej teorii umysłu*, Lublin 2000, s. 94.

<sup>53</sup> G. Szulczewski, *Sztuczna inteligencja a inteligencja moralna. Zagadnienia wstępne cybernetyki*, „Annales. Ethics in Economic Life” 2019, nr 22, s. 21.

ka”<sup>54</sup>, czyli konstrukcji zaproponowanej przez Philippę Foot<sup>55</sup>. Dylemat polegać ma na sytuacji, w której motorniczy traci kontrolę nad hamulcem, a jednocześnie dostrzega na szynach pięć osób. Wykorzystując zwrotnicę, może przekierować pojazd na drugi tor, na którym zginie tylko jeden człowiek. Moralnym dylematem jest zatem to, jak powinien zachować się motorniczy<sup>56, 57</sup>. Szulczewski nawiązuje w tym miejscu do deliberacji Thomasa Cathcarta, który upatruje w świetle niniejszego dylematu kilka możliwych stanowisk<sup>58</sup>. Proponuje między innymi podejście właściwe dwóm wersjom utilitaryzmu, a także ujęcie z punktu widzenia następujących koncepcji: imperatywu katerycznego Immanuela Kanta, koncepcji zmysłu moralnego Hume’a, etyki altruizmu Petera Singera, koncepcji dobra George’a Edwarda Moore’a, podwójnego skutku, amoralizmu Friedricha Nietzschego, a także złotej reguły<sup>59</sup>.

Niezależnie od tego, jaka droga do rozwiązania niniejszego problemu zostanie obrana, Szulczewski pryncypialnie wskazuje, że Cathcart nie udziela jednoznacznej odpowiedzi, która byłaby właściwa dla świata informatyków. Informatycy bowiem są zmuszeni poruszać się w systemie zero-jedynkowym, zaś przedstawienie wielu różnych rozwiązań moralnych nie może stanowić semafora w pożądanym z punktu widzenia programistów binarnym systemie moralnym. Wypada w tym miejscu postawić pytanie natury bardziej ogólnej i zastanowić się, czy w ogóle możliwa jest moralność i etyka o charakterze zero-jedynkowym. Niewykluczone jest prawdopodobnie zero-jedynkowe podejście do etyki. Jak przykładowo mówi Magdalena Kwasek, członkini zarządu ANG Spółdzielni: „W kwestii etyki jesteśmy zero-jedynkowi”<sup>60</sup>. Można prawdopodobnie wyciągać konsekwencje i nakładać określone kary za postępowania nieetyczne w sposób zero-jedynkowy. I tak na przykład Andrzej Michałowski zaznacza, powołując się na § 23 Kodeksu Etyki Adwokackiej, że adwokaci nie mogą pozyskiwać klientów w sposób sprzeczny z godnością zawodu i naruszać w tym zakresie prawa i zasady współżycia społecznego, i że jako członkowie samorządu zawodowego traktują obejście lub złamanie tej zasady jako rzecz zero-jedynkową, bez możliwości jakiegokolwiek niuansowania<sup>61</sup>.

<sup>54</sup> Problem znany jest także jako „dylemat zwrotnicy”, w literaturze anglojęzycznej — *trolley problem*.

<sup>55</sup> D. Edmonds, *Would You Kill the Fat Man? The Trolley Problem and What Your Answer Tells Us About Right and Wrong*, Princeton 2014, s. 9.

<sup>56</sup> J. Hacker-Wright, *Philippa Foot’s Moral Thought*, Londyn 2013, s. 107.

<sup>57</sup> J.J. Thomson, *Turning the Trolley*, „Philosophy & Public Affairs” 2008, nr 36, s. 359–374.

<sup>58</sup> G. Szulczewski, *op. cit.*

<sup>59</sup> *Ibidem*, s. 21–22.

<sup>60</sup> K. Patalan, M. Kwasek, *W kwestii etyki jesteśmy zerojedynkowi*, <https://www.miesiecznik-benefit.pl/wywiad/news/w-kwestii-etyki-jestesmy-zerojedynkowi/> (dostęp: 10.01.2021).

<sup>61</sup> A. Krzyżanowska, *Adwokaci i radcowie na bakier z etyką*, <https://www.rp.pl/Etyka-i-reklama/309279918-Adwokaci-i-radcowie-na-bakier-z-etyka.html> (dostęp: 10.01.2021).

Pytanie jest jednak zgoła inne, gdyż nie chodzi o potencjalność samego zero-jedynkowego nastawienia do przestrzegania norm moralnych i etycznych czy zero-jedynkowego (niezniuansowanego) karania za nieprzestrzeganie norm deontologicznych, lecz o samo ukonstytuowanie praw moralnych i etycznych, które w swojej istocie byłyby zbliżone do komputerowego kodu binarnego.

W „dylemacie wagonika” pojawia się co prawda dychotomiczny wybór, jednak nieostry. Krótko mówiąc, chodzi o to, co byłoby mniejszym złem — śmierć jednej osoby czy pięciu. Inną wersją „dylematu wagonika” jest przykład tak zwanego „grubego człowieka na torach”<sup>62</sup>. Możliwością w tej hipotetycznej sytuacji byłoby zrzućenie z kładki „grubego człowieka”, który zahamowałby pęd pociągu, przez co uratowałby pięć osób, samemu jednak tracąc życie.

Powyższe kwestie dobrze uwidaczniają moralne problemy związane z pojazdami<sup>63</sup>. Przykłady te mogłyby być klasyfikowane jako tak zwane *hard cases*, czyli stany faktyczne, których kwalifikacja z moralnego punktu widzenia pozostaje trudna do dwudzielnego rozstrzygnięcia<sup>64, 65</sup>. Stwierdzić jednak należy, że w przyszłości problem wagoników nie będzie aktualny o tyle, o ile ograniczona zostanie śmiertelność spowodowana przez wypadki. Prognozować należy znaczące zwiększenie bezpieczeństwa ruchu, stąd można przewidywać, że konieczność moralnego wyboru będzie zredukowana<sup>66, 67</sup>. Co najistotniejsze, to właśnie potrzeba podjęcia natychmiastowego wyboru moralnego budzi w kierowcach panikę i nieracjonalne zachowania, przy AI zaś zostaje wyeliminowany niebezpieczny problem emocjonalnej reakcji<sup>68</sup>. Sytuację tę analizuje się w Niemczech, gdzie komisja etyczna pracuje nad kwestią programowania wartości w samochodach autonomicznych tak, aby priorytetem była ochrona życia ludzkiego, pod tym jednak warunkiem, że konkretne ludzkie życia nie mogą być kwalifikowane, zatem

<sup>62</sup> G. Andrade, *Medical Ethics and the Trolley Problem*, „The Journal of Medical Ethics and History of Medicine” 2019, nr 3, s. 9.

<sup>63</sup> S.S. Wu, *Autonomous Vehicles, Trolley Problems, and the Law*, „Ethics and Information Technology” 2020, nr 22, s. 1–13.

<sup>64</sup> A. Brezko, *O dylematach bioetycznych w kontekście praw jednostki*, [w:] *Demokracja, teoria prawa, sądownictwo konstytucyjne. Księga jubileuszowa dedykowana profesorowi zw. nauk prawnych Adamowi Jamrozowi z okazji pięćdziesięciolecia pracy zawodowej*, red. M. Aleksandrowicz et al., Białystok 2018, s. 114.

<sup>65</sup> Według Ronalda Dworkina przykładem *hard case* jest stan faktyczny w sprawie *Henningsen vs. Bloomfield Motors, Inc.*, która dotyczy odpowiedzialności gwarancyjnej producenta samochodów w związku z wypadkiem. Zob. D.O. Brink, *Originalism and Constructive Interpretation*, [w:] *The Legacy of Ronald Dworkin*, red. W. Waluchow, S. Sciaraffa, New York 2016, s. 276.

<sup>66</sup> W. Choromański, I. Grabarek, *Pojazdy autonomiczne w aglomeracjach miejskich*, „Czasopismo Transport Miejski i Regionalny” 2018, nr 11, s. 15.

<sup>67</sup> Inaczej między innymi M. Domagała, *Zagrożenia związane z wprowadzeniem pojazdów autonomicznych jako przykład negatywnych skutków rozwoju sztucznej inteligencji*, [w:] *Prawo sztucznej inteligencji*, red. L. Lai, M. Świerczyński, Warszawa 2020, s. 229–248.

<sup>68</sup> S. Nyholm, J. Smids, *The Ethics of Accident-Algorithms for Self-Driving Cars: an Applied Trolley Problem?*, „Ethical Theory and Moral Practice” 2016, nr 19, s. 1278.

przewiduje się skrajne zredukowanie możliwości decydowania o tym, którą osobę uratować, a którą potrącić (*trolley problem*) na podstawie cech takich jak płeć, wiek czy kolor skóry<sup>69</sup>.

Zwrócono już uwagę, że wywołanie w algorytmach AI intuicji do moralnego postępowania musi rozpocząć się od siły sprawczej („tresury”, szkolenia). Taki system powinien opierać się na dostarczeniu algorytmom AI jak największej bazy danych. Jest to tak zwane „karmienie” sztucznej inteligencji<sup>70</sup>.

Pamiętając o powyższym, należy postulować, że systemem idealnym do unormowania AI jest system reguł prawnych. System ów wydaje się wręcz stworzony do rozwoju prawa AI, gdyż stan faktyczny każdorazowo musi zostać rozwiązany w oparciu o wybór ujęty jako „wszystko albo nic”. W takim przypadku z jednej strony zostaje zagwarantowana ciągłość decyzyjna i spójność systemowa, z drugiej zaś możliwe byłoby uczenie algorytmów AI na bazie schematów, czyli „reguł”<sup>71</sup>. W zdecydowanej większości przypadków życia codziennego algorytmy funkcjonowałyby na bazie Dworkinowskich reguł, na których oparty jest zarówno system drogowy (światło zielone vs. czerwone, szlaban zamknięty vs. otwarty, linia ciągła vs. przerywana, miejsce postojowe wolne vs. zajęte), jak i sama konstrukcja pojazdu (światła włączone vs. wyłączone, prędkość dopuszczalna vs. niedopuszczalna, ilość paliwa wystarczająca vs. niewystarczająca), a także formalności związane z posiadaniem pojazdu (raty leasingowe/ubezpieczenie OC/AC opłacone vs. nieopłacone, przegląd pojazdu aktualny vs. nieaktualny).

## UWAGI KOŃCOWE

W świetle tych okoliczności uregulowanie pozycji prawnej AI w oparciu o normy neutralne moralnie w rozumieniu Błachuta należy uznać za zasadne i pożądane. Sytuacja taka z jednej strony może rozbić moralno-rewizyjny monopol państwa i jego depluralizacyjną arbitralność, której negatywne skutki zostały zaprezentowane na przykładzie strategii walki z pandemią w Szwecji („Sztuczna inteligencja vs. rząd szwedzki”), a z drugiej — może przyczynić się do wzmocnienia stabilności działania algorytmów w oparciu o binarne reguły działań moralnych, wyrażone w sposób „wszystko albo nic” („Reguły vs. zasady”).

Środkiem do zaimplementowania reguł byłoby uczenie algorytmów AI schematów bez ontologicznego uzasadnienia problemu. Uczenie schematów nastąpiło już choćby w przypadku diagnostyki medycznej czy pojazdów autonomicznych. Brakuje jednak „karmienia” algorytmów AI regułami prawa, których będą neu-

<sup>69</sup> P. Konar, *Czy sztuczna inteligencja może nauczyć się moralności?*, <https://www.f5.pl/> (dostęp: 10.01.2021).

<sup>70</sup> D. Reborn, *Digitalismus: Die Utopie einer neuen Gesellschaftsform in Zeiten der Digitalisierung*, Wiesbaden 2019, s. 65, 104, 108, 133, 164.

<sup>71</sup> R. Dworkin, *Biorąc prawa poważnie*, przeł. T. Kowalski, Warszawa 1998, s. 57.

tralnie i autonomicznie przestrzegać, zwiększając poziom bezpieczeństwa i przejmując niektóre funkcje opiekuńcze, nie na wzór systematyki państw opiekuńczych, lecz za przykładem moralnie intuicyjnego postępowania przeszkolonego psa-opiekuna.

## SHOULD A LAW BE MORALLY NEUTRAL? REMARKS ON REGULATION OF ARTIFICIAL INTELLIGENCE

### Summary

This article discusses the issues concerning artificial intelligence regulation. The need to regulate artificial intelligence was indicated. The author postulates a morally neutral law for AI algorithms. The proposal to regulate artificial intelligence by means of morally neutral law is based on the philosophy of law theory by Axel Hägerström and Ronald Dworkin. It is about a independence of an artificial intelligence from state power instruments. The study concludes that there is a need to regulate artificial intelligence for the sake of protecting humans from human and AI threats.

Keywords: artificial intelligence, morality, neutrality

### BIBLIOGRAFIA

- Andrade G., *Medical Ethics and the Trolley Problem*, „The Journal of Medical Ethics and History of Medicine” 2019, nr 3, s. 1–15.
- Bakalarczyk R., *Opieka nad seniorami w państwie opiekuńczym — przykład Szwecji*, „Problemy Polityki Społecznej. Studia i Dyskusje” 2012, nr 18, s. 107–118.
- Banaś M., *Szwedzka polityka integracyjna wobec imigrantów*, Kraków 2010.
- Bernasiewicz M., *Interakcjonizm symboliczny w teorii i praktyce resocjalizacyjnej*, Kraków 2011.
- Bjarup J., *The Philosophy of Scandinavian Legal Realism*, „Ratio Juris” 2005, nr 18, s. 1–15.
- Błachut M., *Postulat neutralności moralnej prawa a konstytucyjna zasada równości*, Wrocław 2005.
- Breczko A., *O dylematach bioetycznych w kontekście praw jednostki*, [w:] *Demokracja, teoria prawa, sądownictwo konstytucyjne. Księga jubileuszowa dedykowana profesorowi zw. nauk prawnych Adamowi Jamrozowi z okazji pięćdziesięciolecia pracy zawodowej*, red. M. Aleksandrowicz et al., Białystok 2018, s. 107–121.
- Brink D.O., *Originalism and Constructive Interpretation*, [w:] *The Legacy of Ronald Dworkin*, red. W. Waluchow, S. Sciaraffa, New York 2016, s. 273–298.
- Bunikowski D., *Idea neutralności moralnej prawa we współczesnych systemach prawnych*, [w:] *Etyka. Część 2*, red. S. Janeczek, A. Starościc, Lublin 2016, s. 541–577.
- Casas-Roma J., Arnedo-Moreno J., *From Games to Moral Agents: Towards a Model for Moral Actions*, [w:] *Artificial Intelligence Research and Development: Proceedings of the 22<sup>nd</sup> International Conference of the Catalan Association for Artificial Intelligence*, red. J. Sabater-Mir et al., Amsterdam 2019, s. 19–28.
- Cassirer E., Axel Hägerström. *Eine Studie zur Schwedischen Philosophie der Gegenwart*, Göteborg 1939.
- Chłopecka A., *Problematyka odpowiedzialności za ruch autonomicznych samochodów w kontekście ochrony praw człowieka*, „Człowiek w Cyberprzestrzeni” 2018, nr 1, s. 29–43.
- Chłopecki A., *Sztuczna inteligencja — szkice prawnicze i futurologiczne*, Warszawa 2018.

- Choromański W., Grabarek I., *Pojazdy autonomiczne w aglomeracjach miejskich*, „Czasopismo Transport Miejski i Regionalny” 2018, nr 11, s. 12–16.
- Koronawirus w Szwecji. *Uboczne skutki eksperymentu*, „Deutsche Welle”, <https://www.dw.com/pl/koronawirus-w-szwecji-uboczne-skutki-eksperymentu/a-55544999>.
- Domagała M., *Zagrożenia związane z wprowadzeniem pojazdów autonomicznych jako przykład negatywnych skutków rozwoju sztucznej inteligencji*, [w:] *Prawo sztucznej inteligencji*, red. L. Lai, M. Świerczyński, Warszawa 2020, s. 229–248.
- Doroszewski W., *Studia i szkice językoznawcze*, Warszawa 1962.
- Dudek M., *Autonomia, neutralność i indyferentność moralna prawa a jego społeczeństwo*, „Ruch Prawniczy, Ekonomiczny i Socjologiczny” 2014, nr 4, s. 69–81.
- Dworkin R., *Biorąc prawa poważnie*, przeł. T. Kowalski, Warszawa 1998.
- Edmonds D., *Would You Kill the Fat Man? The Trolley Problem and What Your Answer Tells Us About Right and Wrong*, Princeton 2014.
- Eliaeson S., *Axel Hägerström and modern social thought*, „Zeitschrift für Politik, Wirtschaft und Kultur” 2000, nr 1, s. 19–30.
- Fehler W., *Sztuczna inteligencja — szansa czy zagrożenie?*, „Studia Bobolanum” 2017, nr 3, s. 69–83.
- Foucault M., *Nadzorować i karać. Narodziny więzienia*, przeł. T. Komendant, Warszawa 2009.
- Greczner B., *Precedens a spójność aksjologiczna prawa w ujęciu Ronalda Dworkina*, „Przegląd Prawa i Administracji” 82, 2010, s. 15–33.
- Hacker-Wright J., *Philippa Foot's Moral Thought*, London 2013.
- Hart H.L.A., *Scandinavian Realism*, „The Cambridge Law Journal” 1959, nr 2, s. 233–240.
- Hetmański M., *Umysł a maszyny: krytyka obliczeniowej teorii umysłu*, Lublin 2000.
- Jachowicz A., *Funkcjonowanie pomocy społecznej. Wybrane problemy*, Dąbrowa Górnicza 2011.
- Jędrsyk M., *Szwedzki sen: czy jest sens walczyć z koronawirusem tak, jak robiła to Szwecja? A raczej — czy był?*, <https://oko.press/czy-jest-sens-walczyć-z-koronawirusem-tak-jak-szwecja/>.
- Kancelaria Rządu Szwecji, *Szwedzki system rządów*, Sztokholm 2014.
- Komisja Europejska. Cordis, *Sztuczna inteligencja przejmuję laboratoryjne badania na obecność patogenów*, <https://cordis.europa.eu/article/id/421710-artificial-intelligence-overcomes-laboratory-testing-for-pathogen-detection/pl>.
- Konar P., *Czy sztuczna inteligencja może nauczyć się moralności?*, <https://www.f5.pl>.
- Koronawirus w Szwecji: Czy szwedzki model walki z pandemią sprawdza się?, „Dziennik Gazeta Prawna”, <https://www.gazetaprawna.pl/wiadomosci/artykuly/1470155,koronawirus-w-szwecji-szwedzki-model-walki-z-pandemia.html>.
- Koronawirus w Szwecji. Szwedzki model walki z pandemią, „Gazeta Prawna”, <https://www.gazetaprawna.pl/wiadomosci/artykuly/1470155,koronawirus-w-szwecji-szwedzki-model-walki-z-pandemia.html>.
- Krasuski A., *Status prawny sztucznego agenta. Podstawy prawne zastosowania sztucznej inteligencji*, Warszawa 2020.
- Krzyżanowska A., *Adwokaci i radcowie na bakier z etyką*, <https://www.rp.pl/Etyka-i-reklama/309279918-Adwokaci-i-radcowie-na-bakier-z-etyka.html>.
- Lipski A., *Perspektywy socjologii kultury artystycznej*, Warszawa 1992.
- Mindus P., *Real Mind. The Life and Work of Axel Hägerström*, Dordrecht 2009.
- Neumann T., *Perspektywy wykorzystania pojazdów autonomicznych w transporcie drogowym w Polsce*, „Autobusy” 2018, nr 12, s. 787–794.
- Nyholm S., Smids J., *The Ethics of Accident-Algorithms for Self-Driving Cars: an Applied Trolley Problem?*, „Ethical Theory and Moral Practice” 2016, nr 19, s. 1275–1289.
- Oniszczyk J., *Filozofia i teoria prawa*, Warszawa 2012.
- Pałubicka A., *Kulturowy wymiar ludzkiego świata obiektywnego*, Poznań 1990.
- Patalan K., Kwasek M., *W kwestii etyki jesteśmy zerojedynkowi*, <https://www.miesiecznik-benefit.pl/wywiad/news/w-kwestii-etyki-jestesmy-zerojedynkowi/> (dostęp: 10.01.2021).

- Payne A., *Premier Szwecji: nasza strategia walki z pandemią była błędna*, <https://businessinsider.com.pl/wiadomosci/koronawirus-w-szwecji-szwedzka-droga-premier-nasza-strategia-byla-bledna/e9g1f38>.
- Radio Białystok, *Sztuczna inteligencja pomaga identyfikować COVID-19 w płucach*, <https://www.radio.bialystok.pl/koronawirus/index/id/193618>.
- Rebhorn D., *Digitalismus: Die Utopie einer neuen Gesellschaftsform in Zeiten der Digitalisierung*, Wiesbaden 2019.
- Rothert A., *Nowa biologia polityczna*, [w:] *Metafory polityki*, t. 3, red. B. Kaczmarek, Warszawa 2005, s. 288–300.
- Ruin H., *Hägerström, Nietzsche and Swedish Nihilism*, [w:] *Axel Hägerström and Modern Social Thought*, red. S. Eliaeson, P. Mindus, S.P. Turner, Oxford 2014, s. 177–202.
- Savulescu J., Maslen H., *Moral Enhancement and Artificial Intelligence: Moral AI?*, [w:] *Beyond Artificial Intelligence: The Disappearing Human-Machine Divide*, red. J. Romportl, E. Zackova, J. Kelemen, Cham 2015, s. 79–96.
- Sawerwain M., Wiśniewska J., *Informatyka kwantowa. Wybrane obwody i algorytmy*, Warszawa 2015.
- Skoog I., *COVID-19 and Mental Health Among Older People in Sweden*, „International Psychogeriatrics” 2020, nr 10, s. 1173–1175.
- Szulczewski G., *Sztuczna inteligencja a inteligencja moralna. Zagadnienia wstępne cybernetyki*, „Annales. Ethics in Economic Life” 2019, nr 22, s. 19–31.
- Szyszkowska M., *Zarys filozofii prawa. Fragmenty dzieł filozoficznoprawnych*, przeł. C. Tarnogórski, Białystok 2000.
- Tadeusiewicz R., *Automatyzacja i sztuczna inteligencja jako źródła prawdziwych i wyimaginowanych zagrożeń*, [w:] *Czy świat należy urządzić inaczej. Schyłek i początek*, red. B. Galwas, P. Kozłowski, K. Prandecki, Warszawa 2019, s. 29–43.
- Tarkowska E., *Ciągłość i zmiana socjologii francuskiej: Durkheim, Mauss, Lévi-Strauss*, Warszawa 1974.
- Tersman F., *Methodological Reflections on Hägerström's Meta-ethics*, <http://www.diva-portal.se/smash/get/diva2:713073/FULLTEXT02.pdf>.
- Thomson J.J., *Turning the Trolley*, „Philosophy & Public Affairs” 2008, nr 36, s. 359–374.
- Van Wynsberghe A., Robbins S., *Critiquing the Reasons for Making Artificial Moral Agents*, „Science and Engineering Ethics” 2019, nr 25, s. 719–735.
- Wojciechowski M., *Axel Hägerström*, [w:] *Filozofia prawa w pytaniach i odpowiedziach*, red. J. Zajadło, K. Zeidler, Warszawa 2013, s. 82–88.
- Wu S.S., *Autonomous vehicles, trolley problems, and the law*, „Ethics and Information Technology” 2020, nr 22, s. 1–13.