

Sztuczne Państwo: między inteligentną administracją a automatokracją w świetle

The Rise and Fall of the Artificial State



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje zjawisko Sztucznego Państwa jako subtelny transformację ustrojową, w której funkcje poznawcze instytucji są zastępowane przez systemy optymalizacji i predykcji. Autor odrzuca skrajny triumfalizm i katastrofizm technologiczny, wskazując, że kluczowym problemem nie jest sama inteligencja maszyn, lecz redukcja antropologiczna obywatela do postaci rekordu danych. Proces ten prowadzi do powstania automatokracji – systemu, w którym człowiek pozostaje formalnie obecny, ale traci rolę normatywnego autora decyzji. Tekst ostrzega przed nadawaniem reprezentacjom algorytmicznym rangi ontologicznej oraz przed mechanizmem samospełniających się prorocstw, gdzie predykcje modelu kształtują rzeczywistość społeczną i ograniczają wolność jednostki poprzez niejawne profilowanie.

The article analyzes the phenomenon of the Artificial State as a subtle systemic transformation in which the cognitive functions of institutions are replaced by optimization and prediction systems. The author rejects extreme technological triumphalism and catastrophism, pointing out that the key problem is not machine intelligence itself, but the anthropological reduction of the citizen to a data record. This process leads to the creation of an automatocracy—a system in which humans remain formally present but lose their role as the normative authors of decisions. The text warns against granting algorithmic representations ontological status and against the mechanism of self-fulfilling prophecies, where model predictions shape social reality and limit individual freedom through covert profiling.

Artykuł analizuje koncepcję „Sztucznego Państwa” nie jako wizję rządów robotów, lecz jako proces subtelnej transformacji ustrojowej, w której podmiotowość obywatela

zostaje zastąpiona przez jego cyfrowy profil. Autor stawia tezę, że największym zagrożeniem współczesnej algorytmizacji administracji jest redukcja człowieka do formatu danych możliwych do przetworzenia, co prowadzi do zastąpienia normatywnego osądu i politycznej wolności przez probabilistyczną predykcję. Tekst dowodzi, że gdy instytucje zaczynają traktować wynik modelu jako substytut rozumowania, dochodzi do odwrócenia relacji między człowiekiem a narzędziem: to nie technologia jest dopasowywana do potrzeb ludzkich, lecz człowiek jest rekonstruowany tak, by stać się czytelny dla algorytmu. Autor postuluje konieczność rekonstytucjonalizacji relacji człowiek-maszyna oraz wprowadzenie prawa do „biograficznej nieprzewidywalności”, aby zapobiec powstaniu automatokracji, w której formalne struktury demokracji pozostają nienaruszone, lecz tracą realny wpływ na decyzje podejmowane na podstawie funkcji celu i optymalizacji danych.

SŁOWA KLUCZOWE

Sztuczne Państwo

automatokracja

predykcja algorytmiczna

funkcja celu

automation bias

human in the loop

government by scores

problem alignment

prawo Goodharta

dominacja infrastrukturalna

racjonalność formalna

system socjotechniczny

redukcja antropologiczna

kontestowalność decyzji

legitymizacja polityczna

Sztuczne Państwo jako redukcja człowieka do danych

Sztuczne Państwo jawi się jako problem konstytucyjny, epistemiczny i antropologiczny. Przeprowadzona analiza pozwala przede wszystkim odrzucić dwie symetryczne, choć przeciwstawne interpretacje współczesnej sztucznej inteligencji. Pierwszą z nich jest triumfalizm technologiczny, według którego wzrost mocy obliczeniowej i jakości modeli prowadzi niemal automatycznie do usprawnienia instytucji. Drugą stanowi katastrofizm, który każdą formę automatyzacji przedstawia jako nieuchronny krok ku utracie wolności. Obydwa stanowiska popełniają ten sam błąd metodologiczny: przypisują technologii własną teleologię. Tymczasem historia — od spisu ludności i maszyny Holleritha, przez cybernetykę i Simulmatics, aż po platformy cyfrowe oraz modele generatywne — wskazuje, że technologia nie posiada jednego społecznego przeznaczenia. Jej skutki zależą od reżimu własności, funkcji celu, struktur

odpowiedzialności, zasad prawa, możliwości odwołania, koncentracji infrastruktury oraz kultury organizacyjnej instytucji, które ją wykorzystują.

Tak rozumiane Sztuczne Państwo nie jest zatem państwem „rządzonym przez roboty” w sensie znanym z literatury science-fiction. To raczej subtelna transformacja ustrojowa, polegająca na stopniowym przenoszeniu podstawowych funkcji poznawczych i organizacyjnych instytucji do systemów opartych na klasyfikacji, predykcji i optymalizacji. Kluczowy moment nie następuje wtedy, gdy maszyna formalnie podpisuje decyzję, lecz w chwili, gdy instytucja zaczyna traktować wynik modelu jako wystarczający substytut samodzielnego rozumowania, publicznego uzasadnienia czy obywatelskiego uczestnictwa. Automatokracja jest zatem nie tyle systemem bez ludzi, ile systemem, w którym ludzie mogą pozostać formalnie obecni, lecz stopniowo tracić funkcję normatywnego autora decyzji.

W tym miejscu powraca Hobbes. Parafrazując centralną metaforę Lewiatana, państwo jest sztucznym człowiekiem skonstruowanym przez ludzi po to, aby chronić ich przed skutkami własnej bezbronności i konfliktowości. Hobbesowski artefakt posiadał więc funkcję służebną wobec człowieka, mimo że wyposażono go w ogromną władzę. Współczesny problem zaczyna się jednak wtedy, gdy kierunek tej konstrukcji ulega odwróceniu: nie maszyna zostaje dopasowana do człowieka, lecz człowiek jest rekonstruowany w postaci łatwiejszej do przetwarzania przez maszynę.

Obywatel staje się rekordem, rekord profilem, profil predykcją, a predykcja podstawą interwencji. Sztuczne Państwo jest dlatego nie tyle nowym Lewiatanem, ile odwróceniem antropologii Lewiatana: artefakt stworzony dla człowieka zaczyna wymagać człowieka przedstawionego w formie zgodnej z wymaganiami artefaktu. Najważniejszym wkładem teoretycznym proponowanej koncepcji jest zatem przesunięcie analizy z pytania o „inteligencję maszyny” na pytanie o instytucjonalną relację między reprezentacją człowieka a człowiekiem reprezentowanym. Techniczny model jednostki jest bowiem z konieczności redukcją.

Predykcja algorytmiczna jako zagrożenie dla podmiotowości i wolności

Profil nie może zawierać całej biografii, wszystkich motywacji, zmienności emocjonalnej, zdolności uczenia się czy moralnego przewartościowania. Redukcja sama w sobie nie jest wadą; bez niej analiza statystyczna byłaby niemożliwa. Problem pojawia się jednak w chwili, gdy administracyjna lub algorytmiczna reprezentacja otrzymuje rangę ontologiczną, jak gdyby opis probabilistyczny stanowił pełną prawdę o osobie.

Teoria Sztucznego Państwa styka się tutaj z analizą Iana Hackinga dotyczącą „tworzenia ludzi” przez kategorie wiedzy. Klasyfikacje społeczne nie są bowiem jedynie biernymi etykietami. Człowiek może zacząć reagować na kategorię, która została wobec niego zastosowana; instytucje zaczynają traktować go zgodnie z tym schematem, a zachowania wynikające z nowych warunków zdają się następnie potwierdzać pierwotną klasyfikację. Algorytmiczne profilowanie może radykalnie przyspieszyć ten mechanizm, ponieważ klasyfikacja nie musi być już jawna, stabilna ani zrozumiała dla osoby klasyfikowanej. Profil może istnieć wyłącznie jako wewnętrzna probabilistyczna reprezentacja systemu, a mimo to realnie wpływać na ceny, treści, oferty, kontrole i dostępne możliwości.

Robert Merton opisywał podobną dynamikę poprzez pojęcie samospełniającego się proroctwa. Fałszywa początkowo definicja sytuacji może wywołać zachowania, które doprowadzą do jej późniejszego potwierdzenia. W systemach predykcyjnych mechanizm ten przybiera szczególnie silną formę: jeśli model przewiduje małe prawdopodobieństwo sukcesu i na tej podstawie ogranicza człowiekowi dostęp do szans, późniejszy brak sukcesu nie musi potwierdzać trafności modelu, lecz może być skutkiem wcześniejszej interwencji. Badanie algorytmicznej sprawiedliwości wymaga zatem rozróżnienia pomiędzy predykcją trafną deskryptywnie a taką, która staje się prawdziwa wskutek własnego zastosowania. Rozróżnienie to ma zasadnicze znaczenie dla teorii państwa, gdyż tradycyjne prawo – przynajmniej normatywnie – zakłada możliwość zmiany. Kara wygasa, zadłużenie można spłacić, obywatel może zmienić pogląd; osoba nie powinna być wiecześnie utożsamiana ze swoją przeszłością.

Algorytmiczna infrastruktura potrafi natomiast konserwować historię w coraz dokładniejszym profilu. Dlatego jednym z nowych praw podmiotowych powinno stać się nie tylko prawo do sprostowania fałszywej informacji, lecz szersze prawo do niebycia całkowicie definiowanym przez probabilistyczny obraz własnej przeszłości.

Koncepcja ta spotyka się z myślą Hannah Arendt. Parafrazując jej refleksję nad natalnością, człowiek jest istotą zdolną rozpocząć coś nowego, a właśnie możliwość nowego początku stanowi jeden z ontologicznych fundamentów politycznej wolności. System predykcyjny z natury poszukuje regularności między przeszłością a przyszłością, podczas gdy człowiek polityczny posiada zdolność do działania naruszającego tę regularność. Obywatel może zagłosować przeciw własnemu interesowi materialnemu, zmieniając ocenę moralną; może zerwać z grupą, wybaczyć, nawrócić się, zbuntować lub podporządkować wartości, której wcześniejszy profil nie zawierał. Z punktu widzenia statystyki takie działanie jest błędem predykcji; z perspektywy filozofii politycznej – przejawem wolności.

Ta obserwacja prowadzi do ważnego uzupełnienia liberalnej teorii praw. Obok prawa do prywatności, wolności wypowiedzi i rzetelnego procesu niezbędne staje się prawo do politycznej i biograficznej nieprzewidywalności. Nie byłoby ono zakazem prowadzenia analiz statystycznych, lecz normą zakazującą przekształcania predykcji w kompletną definicję osoby i automatyczną podstawę rozstrzygnięć dotyczących jej najważniejszych praw.

Podobny problem ujawnia się w teorii Maxa Webera. Biurokracja była dla niego jednym z najpotężniejszych instrumentów racjonalizacji, gdyż umożliwiała zastępowanie osobistego kaprysu przewidywalnymi regułami, profesjonalizacją i kompetencją. Algorytmizacja administracji może zatem stanowić dalsze rozwinięcie racjonalności formalnej. Model stosujący jednakowe kryteria wobec milionów spraw może ograniczać nierówności wynikające z odmiennego zachowania poszczególnych urzędników, co z perspektywy równości wobec prawa jest potencjalną zaletą.

Weberowskie ostrzeżenie pojawia się jednak wtedy, gdy racjonalność środka przejmuje funkcję racjonalności celu. Nowoczesność może stworzyć system niezwykle sprawny w realizowaniu zadań, którego uczestnicy coraz rzadziej pytają, dlaczego właśnie te zadania powinny być realizowane. W algorytmicznej administracji niebezpieczeństwo to przyjmuje formę funkcji celu: model może skutecznie zmniejszać dany wskaźnik, lecz nie potrafi rozstrzygnąć, czy wskaźnik ten właściwie reprezentuje dobro wspólne. Tutaj pojawia się epistemologiczna granica technokracji. Ekspertyza zwiększa wiedzę o relacji środek-cel, ale nie tworzy automatycznie prawomocności celu.

Inżynier może wskazać, jak najefektywniej zmniejszyć korki, lecz to społeczeństwo musi rozstrzygnąć, jaką cenę wolno za to zapłacić w prywatności, przestrzeni miejskiej lub dostępie do transportu. Ekonomista może oszacować najbardziej efektywną strukturę podatku, ale nie rozstrzygnie za pomocą samej ekonomii, jaką wagę należy przypisać redystrybucji.

Rozdzielenie kompetencji epistemicznej od legitymizacji politycznej

Model może przewidywać przestępczość, ale nie odpowie samodzielnie na pytanie, jakie ryzyko fałszywego wskazania niewinnej osoby jest konstytucyjnie dopuszczalne. Istotą demokratycznego konstytucjonalizmu pozostaje rozdzielenie kompetencji epistemicznej od legitymizacji politycznej. Można wiedzieć więcej i nadal nie mieć prawa samodzielnie decydować za innych. Technokracja zaczyna się tam, gdzie przewagę poznawczą traktuje się jako wystarczającą podstawę do przejęcia roli normatywnej.

Perspektywa Herberta Simona pozwala dostrzec, dlaczego całkowite odrzucenie systemów algorytmicznych byłoby błędem w drugą stronę. Człowiek funkcjonuje w warunkach bounded rationality: dysponuje ograniczoną uwagą, pamięcią i zdolnością obliczania konsekwencji. Administracja jest w dużej mierze technologią radzenia sobie z tymi ograniczeniami poprzez procedury, specjalizację i podział pracy. AI może stać się kolejnym etapem tego procesu, zwiększając wydajność instytucji w przetwarzaniu informacji. Parafrazując Simona: realny decydent nie zawsze poszukuje rozwiązania doskonałego; często szuka rozwiązania dostatecznie dobrego, które jest możliwe do znalezienia przy ograniczonym czasie i wiedzy. AI może rozszerzyć zbiór analizowanych wariantów i poprawić jakość przesłanek.

Nie wynika z tego jednak, że maszyna powinna przejąć definicję wartości. Właściwą relacją nie jest opozycja człowiek-maszyna, lecz podział pracy poznawczej i normatywnej. To rozróżnienie prowadzi do koncepcji systemu socjotechnicznego. Instytucja nigdy nie jest wyłącznie technologią ani jedynie zbiorem ludzi; skutek organizacyjny powstaje w interakcji technologii, reguł, kompetencji, bodźców i kultury odpowiedzialności. Ocena AI wyłącznie na podstawie trafności modelu jest zatem metodologicznie niewystarczająca. Należy badać, jak jego obecność zmienia zachowania pracowników, hierarchię wiedzy, możliwość sprzeciwu, proces szkolenia, odpowiedzialność oraz zakres ludzkich kompetencji.

Jeżeli model uzyskuje 95% trafności, ale jednocześnie sprawia, że personel przestaje rozpoznawać 5% przypadków wyjątkowych, całociowy rezultat może być mniej bezpieczny, niż wskazuje benchmark. Z kolei system o niższej trafności, który zmusza urzędnika do weryfikacji sprzecznych przesłanek i ułatwia wychwycenie anomalii, może realnie poprawiać funkcjonowanie instytucji mimo słabszego wyniku technicznego.

Od rządów racji do rządów wskaźników

To właśnie perspektywa socjotechniczna odróżnia naukową analizę AI od marketingowej fascynacji benchmarkami. Problem staje się szczególnie wyraźny w świetle teorii automation bias: człowiek ma tendencję do traktowania wyniku systemu jako bardziej obiektywnego niż własny osąd, co prowadzi do ignorowania informacji świadczących o błędzie algorytmu. Im większa reputacja „inteligentnego” systemu, tym silniejsza staje się psychologiczna presja zgodności.

Powstaje w ten sposób paradoks: formalne pozostawienie człowieka w procesie nie gwarantuje realnej kontroli. Jeżeli warunki organizacyjne sprawiają, że odrzucenie rekomendacji jest kosztowne, koncepcja „human in the loop” staje się instytucjonalną fikcją.

Właściwą kategorią analizy nie jest zatem sama obecność człowieka, lecz realna kontestowalność. Decyzja zyskuje demokratyczne zabezpieczenie wtedy, gdy istnieje podmiot zdolny powiedzieć „nie”, dysponujący informacjami niezbędnymi do uzasadnienia sprzeciwu i niepodlegający karom organizacyjnym za odejście od rekomendacji systemu. Człowiek musi być nie tylko elementem pętli, ale potencjalnym przerywaczem pętli.

Teoria wysokiej niezawodności organizacji wnosi tutaj istotne uzupełnienie. W systemach, w których błąd może prowadzić do katastrofy, bezpieczeństwo nie wynika wyłącznie ze zwiększania trafności pojedynczego komponentu. Kluczowe są redundancja, możliwość eskalacji, kultura zgłaszania problemów oraz czujność wobec zdarzeń nieprzewidywanych. Parafrazując intuicję Weicka i Sutcliffe: organizacja odporna nie uznaje sukcesu za dowód braku ryzyka, lecz stale poszukuje sygnałów wskazujących na ograniczenia własnych modeli. Demokratyczne państwo wykorzystujące AI powinno przyjąć podobną kulturę.

System, który przez milion decyzji działał prawidłowo, nie uzyskuje tym samym prawa do niekwestionowanej władzy w milion pierwszej. Statystyczny sukces nie znosi obowiązku zachowania procedury w przypadku wyjątkowym – a to właśnie przypadek wyjątkowy jest jednym z głównych powodów istnienia prawa.

Koncepcja Sztucznego Państwa powinna być zatem powiązana z teorią procedural justice. Tom Tyler wykazał, że ludzie oceniają instytucje nie tylko przez pryzmat uzyskanego wyniku, ale przede wszystkim przez to, czy zostali wysłuchani, potraktowani z szacunkiem i czy procedura była postrzegana jako bezstronna. Legitymizacja posiada więc komponent relacyjny. Państwo nie może sprowadzić obywatela do obiektu poprawnie sklasyfikowanego przez model; musi traktować go jako osobę uprawnioną do przedstawienia własnych racji.

Ten wymiar prowadzi bezpośrednio do problemu reason-giving. Demokratyczna decyzja nie jest po prostu wynikiem, lecz wynikiem otoczonym racjami. System prawny wymaga uzasadnienia właśnie dlatego, że każda decyzja może zostać zakwestionowana – uzasadnienie jest mechanizmem transformującym władzę w działanie podlegające krytyce. Jeśli administracja odpowiada obywatelowi jedynie wartością scoringu, zastępuje rację wskaźnikiem. Tę patologię można określić jako przejście od government by reasons do government by scores. Wynik liczbowy może być cennym dowodem pomocniczym, lecz nie stanowi samodzielnego uzasadnienia normatywnego.

Optimalizacja algorytmiczna jako ukryta decyzja polityczna

Wartość 0,82 nie wyjaśnia, dlaczego prawo dopuszcza ograniczenie wolności przy takim poziomie ryzyka, dlaczego koszt fałszywego alarmu uznano za mniejszy od kosztu przeoczenia ani kto właściwie wyznaczył próg interwencji. Każdy scoring ma ukrytą konstytucję, ponieważ ktoś musiał zdefiniować funkcję błędu. To właśnie teoria decyzji pozwala ująć ten problem w sposób matematyczny i jednocześnie polityczny.

Klasyfikator generuje dwie główne kategorie błędów: false positive oraz false negative. Nie istnieje jednak czysto techniczna reguła, która rozstrzygałaby, który z nich jest ważniejszy. W medycynie przeoczenie choroby bywa kosztowniejsze niż fałszywy alarm; w prawie karnym skazanie niewinnego ma szczególną wagę. Algorytm potrafi obliczyć trade-off, ale bez uprzednio przyjętej normy nie jest w stanie określić moralnej ceny błędu.

Nawet najbardziej zaawansowane AI pozostaje zależne od przyjętej teorii sprawiedliwości. Optimalizacja nie eliminuje filozofii politycznej – ona jedynie ukrywa ją w parametrach. Problem alignment w kontekście państwa jest zatem przede wszystkim problemem konstytucyjnym. Podczas gdy informatyka pyta, czy system zachowuje się zgodnie z intencją operatora, konstytucjonalizm musi zapytać, czy operator ma prawo posiadać taką intencję, czy cel jest zgodny z prawami podstawowymi i czy jednostka może podważyć sposób jego realizacji.

Możliwe jest bowiem doskonale technicznie "aligned" Sztuczne Państwo. System może bezbłędnie wykonywać instrukcje władzy i właśnie dlatego stanowić zagrożenie. Historia autorytaryzmów nie dokumentuje przecież głównie katastrof wynikających z nieposłuszeństwa administracji, lecz wręcz przeciwnie – ukazuje tragedie spowodowane nadmiernie sprawnym wykonywaniem złych reguł. Problem nie polega zatem wyłącznie na tym, czy maszyna robi to, czego chcemy. Trzeba zapytać, kim jest to „my” oraz jakie instytucje kontrolują to pragnienie.

Ted Chiang wnosi tu trafną intuicję: najbardziej przekonujące wizje groźnej superinteligencji przypominają system gospodarczy pozbawiony zewnętrznych ograniczeń – strukturę niezwykle efektywną w maksymalizacji celu, która nie ma powodu, by respektować wartości niewprowadzone do funkcji celu. W tej analogii naukowo wartościowe nie jest utożsamienie kapitalizmu z AI, lecz dostrzeżenie, że problem instrumentalnej racjonalności jest starszy od sztucznej inteligencji. AI jedynie radykalnie zwiększa prędkość i skalę jego oddziaływania. Do tej samej rodziny należy prawo Goodharta.

Zgodnie z tą zasadą miara traci swoją wartość informacyjną, gdy staje się bezpośrednim celem optymalizacji. W polityce publicznej oznacza to, że wskaźnik bezpieczeństwa może zacząć kształtować praktyki tak, by poprawić sam wskaźnik, a nie realne bezpieczeństwo; wzrost liczby wykrytych nadużyć może świadczyć o skuteczniejszej administracji, ale równie dobrze o agresywniejszym kontrolowaniu grup już wcześniej uznanych za ryzykowne. Algorytm potrafi perfekcyjnie maksymalizować miernik i właśnie dlatego należy stale badać, czy ten miernik nadal reprezentuje rzeczywistą wartość społeczną.

Humanistyka i policentryzm jako bezpieczniki funkcji celu

Sztuczne Państwo można zatem interpretować jako makropolityczny problem Goodharta. Społeczeństwo opiera się na wartościach złożonych, niewspółmiernych i trudnych do zmierzenia, podczas gdy system obliczeniowy wymaga celu możliwego do formalizacji. Każde przełożenie sprawiedliwości, dobrobytu, bezpieczeństwa czy zaufania społecznego na zestaw mierników nieuchronnie pozostawia pewne obszary poza modelem. Im skuteczniejsza staje się optymalizacja, tym kluczowa okazuje się wiedza o tym, co zostało pominięte.

Literatura humanistyczna przestaje być zatem jedynie dodatkiem do nauki o AI, a staje się częścią jej aparatu bezpieczeństwa. Historia, filozofia czy antropologia to dyscypliny szczególnie wyczulone na znaczenia, których nie da się zamknąć w jednym mierniku. Ich funkcją nie jest konkurowanie z matematyką w obszarze predykcji, lecz przypominanie o wymiarach ludzkiego działania niewidocznych dla formalizacji. Można tu przywołać Lemowską intuicję: wzrost kompetencji technologicznych nie likwiduje problemu wartości; przeciwnie – zwiększa konsekwencje ich błędnego wyboru. Im potężniejsze narzędzie, tym mniej rozsądne jest traktowanie pytania o cel jako filozoficznego luksusu. W epoce słabej technologii zły cel mógł pozostać niezrealizowany z powodu braku możliwości. W dobie niezwykle skutecznej AI to właśnie doskonałość wykonania może uczynić błąd aksjologiczny katastrofalnym.

Nauki humanistyczne nie stoją zatem po stronie naiwnego sporu „człowieka z maszyną”, lecz stanowią element systemu kontroli funkcji celu.

Podobną funkcję pełni republikańska teoria wolności Philipa Pettita. Wolność w tym ujęciu nie jest jedynie brakiem faktycznej ingerencji. Człowiek poddany arbitralnej władzy innego podmiotu nie jest w pełni wolny, nawet jeśli ta władza chwilowo z niego nie korzysta. Problem polega na możliwości podporządkowania bez równoważnej kontroli ze strony osoby podporządkowanej.

Teoria ta precyzyjnie opisuje naturę władzy infrastrukturalnej platform i algorytmów. Użytkownik może przez lata nie doświadczać konfliktu z systemem, jednak jeśli operator może jednostronnie zmienić warunki dostępu, ranking czy widzialność, istnieje relacja potencjalnej dominacji. Z punktu widzenia republikanizmu nie wystarczy wierzyć, że operator jest „dobry” – należy ograniczyć jego arbitralność poprzez jasne reguły, prawo do odwołania, interoperacyjność i zewnętrzną kontrolę. W tym świetle antymonopol staje się czymś więcej niż instrumentem obniżania cen; jest mechanizmem rozpraszania władzy. Gdy jeden dostawca kontroluje infrastrukturę, od której nie można realnie odejść, mamy do czynienia z problemem dominacji, co w środowisku AI – przy ogromnej koncentracji kapitału, chipów i danych – ma znaczenie krytyczne.

Z kolei Elinor Ostrom pozwala sformułować pozytywną alternatywę wobec centralistycznego Sztucznego Państwa. Jej główna lekcja wskazuje, że skomplikowane problemy zbiorowe nie zawsze wymagają jednego centralnego zarządcy ani całkowitej prywatyzacji. Systemy policentryczne, w których wiele ośrodków podejmuje częściowo autonomiczne decyzje, mogą być bardziej adaptacyjne i odporne. W odniesieniu do AI prowadzi to do wizji policentrycznej infrastruktury inteligencji. Zamiast jednego narodowego supermodelu lub pełnej zależności od kilku globalnych korporacji, rozwiązaniem może być rozwój systemów publicznych, akademickich, sektorowych i lokalnych, które posiadając różne kompetencje, wzajemnie podlegają kontroli.

Granice koncepcji i empiryczne mierniki automatokracji

Redundancja może sprawiać wrażenie nieefektywności w prostym rachunku ekonomicznym, jednak z perspektywy bezpieczeństwa ustrojowego stanowi cenny zasób. Monokultura informatyczna bywa bowiem równie ryzykowna co monokultura biologiczna: pojedynczy błąd, naruszenie, konflikt interesów czy zmiana warunków dostępu mogą uderzyć w całe społeczeństwo jednocześnie. Choć różnorodność systemów podnosi koszty koordynacji, znacząco redukuje ryzyko wystąpienia jednego punktu awarii. To kolejny dowód na to, że efektywność i odporność nie są pojęciami tożsamymi.

Koncepcja "Sztucznego Państwa" posiada jednak ograniczenia, które należy wyraźnie nazwać. Po pierwsze, metafora państwa może sugerować większą spójność zjawiska, niż ma ona w rzeczywistości. Big Tech nie jest jednolitym podmiotem politycznym; Microsoft, Meta, Alphabet, Amazon, OpenAI, Anthropic czy xAI konkurują ze sobą, mają odmienne interesy i różne relacje z władzami. Nie istnieje jedna doktryna Doliny Krzemowej, podobnie jak brak jest jednolitej reakcji państw narodowych na AI. Po drugie, genealogia idei nie musi być genealogią przyczynową. Fakt,

że współczesny projekt przypomina dawną technokratyczną fantazję, nie dowodzi bezpośredniego dziedziczenia. Technokracja Inc., cybernetyka, libertarianizm, longtermizm i obecna gospodarka AI wyrastają z różnych źródeł i dążą do innych celów. Ich zestawienie jest uzasadnione jedynie wtedy, gdy wskazuje na powtarzalne struktury argumentacyjne, a nie sugeruje istnienie jednej linii spisku czy sukcesji. Po trzecie, krytyka kwantyfikacji wymaga szczególnej ostrożności.

Statystyka, spisy i modele predykcyjne stanowią fundament ogromnych osiągnięć w zdrowiu publicznym, polityce społecznej, epidemiologii czy planowaniu infrastruktury. Świat pozbawiony danych nie byłby bardziej humanitarny – w wielu dziedzinach stałby się wręcz bardziej arbitralny. Właściwym obiektem krytyki nie jest zatem sama kwantyfikacja, lecz niekontrolowane przejście od niej do normatywnego rozstrzygnięcia. Po czwarte, nie wolno zakładać, że każda ludzka decyzja jest z natury bardziej godna zaufania niż algorytmiczna. Badania nad biasem, noise i ograniczoną racjonalnością wskazują na coś przeciwnego; w pewnych obszarach system może być spójniejszy, dokładniejszy i mniej podatny na określone uprzedzenia. Naukowo odpowiedzialna krytyka automatokracji musi zatem dopuścić możliwość, że wzmocnienie roli człowieka wymaga czasami odebrania mu części rutynowej dyskrekcji. Po piąte, nie każda forma automatyzacji ma znaczenie konstytucyjne.

Model układający rozkład jazdy nie może być analizowany w ten sam sposób co system oceniający ryzyko recydywy. Konieczna jest gradacja oparta na naturze funkcji, odwracalności skutków, asymetrii władzy oraz skali wpływu na prawa podstawowe. Bez tego teoria "Sztucznego Państwa" stanie się zbyt szeroka, by zachować wartość analityczną. Dalsze badania powinny zatem odejść od pytania „czy AI rządzi?”, na rzecz analizy mierzalnych zmian w rozkładzie sprawczości instytucjonalnej. Warto badać, jak często urzędnicy odrzucają rekomendacje systemu i czy mają do tego organizacyjne warunki, jak zmienia się jakość uzasadnień decyzji, kto kontroluje progi klasyfikacji oraz jaki jest koszt zmiany dostawcy. Istotne jest również to, w jakim stopniu państwo potrafi audytować technologię i czy obywatele dysponują realną możliwością odwołania. Można wprowadzić wskaźnik automation deference rate – udział przypadków, w których człowiek akceptuje sugestię systemu mimo sprzecznych przesłanek – oraz institutional reversibility, czyli zdolność organizacji do powrotu do alternatywnych metod działania po utracie systemu AI. Innym miernikiem może być infrastructural dependency ratio, określający stopień zależności krytycznych funkcji instytucji od pojedynczego prywatnego dostawcy.

Choć propozycje te wymagają operacjonalizacji, pokazują one drogę przejścia koncepcji "Sztucznego Państwa" z poziomu metafory ku programowi empirycznemu. Kluczowy jest pomiar różnicy między formalnym a rzeczywistym human oversight. Sam zapis w dokumencie, że człowiek jest ostatecznym decydentem, nie gwarantuje, iż posiada on czas, kompetencje i władzę niezbędne

do sprzeciwu. Stąd potrzeba badania *effective contestability*: prawdopodobieństwa, że zakwestionowanie rekomendacji prowadzi do ponownej, niezależnej analizy, a nie jedynie do potwierdzenia wyniku przez kolejną warstwę tego samego systemu. Interesującym obszarem staje się również *epistemic sovereignty*. Państwo może zachować formalną jurysdykcję, tracąc jednocześnie zdolność samodzielnego rozumienia infrastruktury, którą reguluje. Przyszłe analizy powinny więc dotyczyć nie tylko liczby regulacji, lecz relacji kompetencyjnych między administracją a sektorem prywatnym, gdyż suwerenność prawna bez zaplecza technicznego może okazać się niewystarczająca.

W tym miejscu powraca myśl Galbraitha. Parafrazując ją, człowiek może stopniowo stać się sługą w myśleniu i działaniu maszyny, którą pierwotnie skonstruował jako własnego sługę. Współcześnie sens tego spostrzeżenia należy odczytywać organizacyjnie, a nie metafizycznie. Nie chodzi o świadomą dominację maszyn, lecz o sytuację, w której procedury, harmonogramy i kategorie myślenia zostają przebudowane pod możliwość systemu technicznego, a człowiek zaczyna traktować tę architekturę jako naturalną i nieuchronną. Najbardziej niepokojąca postać "Sztucznego Państwa" nie jest więc spektakularna – nie posiada sztandaru, konstytucji ani momentu zamachu stanu.

Powstaje ona poprzez tysiące racjonalnych mikrodecyzji: automatyzujemy ten etap, bo będzie szybciej; zbieramy te dane, bo zwiększą *accuracy*; integrujemy system, by ułatwić obsługę; powierzamy kolejną warstwę zewnętrznemu dostawcy, bo ma większe kompetencje; usuwamy człowieka, bo niemal zawsze zgadza się z modelem.

Rekonstytucjonalizacja relacji człowiek-maszyna poprzez systemowe równoważenie władzy

Każdy z tych kroków, rozpatrywany osobno, może wydawać się dobrze uzasadniony. Ich suma potrafi jednak przesunąć strukturę odpowiedzialności w sposób, którego nikt świadomie nie zaprojektował. Jest to przykład emergencji instytucjonalnej – dowód na to, że złożony porządek może powstać bez centralnego planu. Sztuczne Państwo nie potrzebuje zatem architekta; może być efektem lokalnie racjonalnych decyzji, które prowadzą do globalnie problematycznej zależności.

Właśnie dlatego analiza skupiona wyłącznie na intencjach technologicznych miliarderów byłaby niewystarczająca. Osoby mają znaczenie, lecz jeszcze większy wpływ mają struktury bodźców i zależności, w których Big Tech sam staje się częściowo podmiotem własnej infrastruktury. Menedżer na konkurencyjnym rynku inwestuje w AI, bo obawia się konkurencji; państwo

przyspiesza rozwój z lęku przed przewagą geopolitycznego rywala; inwestor finansuje skalowanie, gdyż brak wzrostu oznacza utratę pozycji. Powstaje logika wyścigu, w której żaden pojedynczy uczestnik nie kontroluje procesu. Teoria gier sugeruje, że struktura ta przypomina dylemat kolektywnego działania: nawet jeśli wszyscy aktorzy preferowaliby świat bezpieczniejszy i bardziej kontrolowany, każdy z nich ma indywidualny bodziec do przyspieszenia, zakładając, że inni nie zwolnią. Rozwiązanie tego problemu nie wymaga zatem moralnego nawrócenia pojedynczych przedsiębiorców, lecz instytucji zdolnych zmienić strukturę tych bodźców.

W tym miejscu demokratyczne prawo odzyskuje swoje znaczenie. Jego funkcją nie jest wyłącznie zakazywanie, lecz tworzenie warunków koordynacji, które pozwolą uczestnikom uniknąć wyścigu ku rezultatowi, którego nikt z nich indywidualnie nie preferuje. Normy bezpieczeństwa lotniczego nie zatrzymały rozwoju lotnictwa – wręcz przeciwnie, umożliwiły jego społeczną legitymizację. Podobnie regulacja AI może stać się infrastrukturą zaufania niezbędną do trwałego skalowania technologii. Demontaż Sztucznego Państwa nie powinien być zatem przedstawiany jako triumf człowieka nad maszyną, gdyż byłaby to kolejna fałszywa dychotomia.

Powinien on oznaczać rekonstytucjonalizowanie relacji człowiek-maszyna. Technologia pozostaje potężnym narzędziem działania, ale jej właściciel nie staje się prawodawcą; predykcja nie staje się automatycznie decyzją; profil nie staje się osobą; funkcja celu nie staje się teorią dobra; rekomendacja nie staje się osądem, a regulamin platformy nie staje się konstytucją sfery publicznej.

W tym kontekście architektura separacji może być cyfrowym odpowiednikiem klasycznego podziału władz. Władza obliczeniowa musi być rozdzielana, ponieważ – jak każda inna władza – ma tendencję do rozszerzania swoich zastosowań w momencie, gdy okazuje się skuteczna. Nie wymaga to zakładania złej woli. James Madison nie projektował konstytucyjnych hamulców z powodu wiary w wyjątkową deprawację moralną urzędników, lecz dlatego, że sama struktura władzy wymaga przeciwstawienia jej innych ośrodków kontroli. Jeśli epoka cyfrowa stworzyła nowy rodzaj władzy infrastrukturalnej, musi wypracować również nowe mechanizmy jej równoważenia. Niektóre z nich już istnieją: prawo ochrony danych, prawo konkurencji, AI Act, audyt, interoperacyjność czy możliwość sądowego odwołania. Inne trzeba dopiero rozwinąć: publiczną kompetencję modelową, niezależne laboratoria audytujące, normy dotyczące generatywnej perswazji politycznej, zasady zachowania alternatywnych ścieżek w krytycznych usługach oraz mechanizmy kontroli infrastrukturalnej koncentracji mocy obliczeniowej.

Ostatecznie jednak najgłębszy problem pozostaje antropologiczny.

Obrona podmiotowości przed redukcją do modelu danych

Sztuczna inteligencja stawia przed polityką pytanie nie tylko o to, co maszyna potrafi, ale o to, jaką teorię człowieka przyjmujemy, projektując system wokół jej możliwości. Czy obywatel jest przede wszystkim zbiorem preferencji, które należy coraz dokładniej przewidywać? Czy raczej podmiotem zdolnym do zmiany własnych przekonań w wyniku argumentu, doświadczenia i moralnego namysłu? Czy jego przeszłe zachowanie stanowi najlepszą definicję przyszłości, czy istnieje w nim zdolność do nowego początku, której żaden model historycznych danych nie wyczerpuje?

Jeżeli przyjmiemy drugą odpowiedź, demokratyczne państwo musi chronić nie tylko wolność od przymusu, lecz także przestrzeń niepełnej przewidywalności człowieka. Można tu ponownie przywołać Arendt: człowiek nie jest wyłącznie istotą reagującą na zastany świat, ale zdolną rozpocząć ciąg zdarzeń, którego nie da się całkowicie wyprowadzić z tego, co istniało wcześniej. Jeżeli algorytmiczny porządek polityczny utraci zdolność do przyjęcia takiego początku, stanie się niezwykle skutecznym systemem reprodukcji przeszłości. Demokracja potrzebuje bowiem nie tylko pamięci, lecz i możliwości niespodzianki.

Ostateczny wybór nie przebiega między tradycyjnym państwem a AI, lecz dotyczy dwóch modeli nowoczesności. Pierwszy dąży do uczynienia społeczeństwa coraz bardziej przewidywalnym, aby można było nim sprawniej zarządzać. Drugi wykorzystuje technologie do zwiększania wiedzy i kompetencji, pozostawiając jednak nienaruszoną zasadę, że człowiek nie wyczerpuje się w swoim modelu. W pierwszym przypadku informacja stopniowo przechodzi w sterowanie; w drugim pozostaje narzędziem osądu.

Właśnie ta różnica odróżnia inteligentne państwo od Sztucznego Państwa. Pierwsze może wykorzystywać ogromne modele, predykcję, automatyzację i centra danych, zachowując obywatela jako autora porządku normatywnego. Drugie może zachować parlament, sądy i wybory, ale uczynić ich uczestników coraz bardziej zależnymi od infrastruktury, która ustala, co widzą, jak są klasyfikowani, jakie możliwości otrzymują i jak przewidywane jest ich zachowanie.

Zewnętrzne podobieństwo ustrojowe może być mylące. Automatokracja nie musi ogłaszać własnej konstytucji. Wystarczy, że ta pozostanie formalnie niezmienną, podczas gdy władza epistemiczna i infrastrukturalna przesunie się poza mechanizmy, które prawo potrafi kontrolować.

Najważniejszym pytaniem XXI wieku nie jest zatem: „czy sztuczna inteligencja stanie się człowiekiem?”. Jest nim pytanie odwrotne: czy człowiek zostanie politycznie zredukowany do

formatu, który sztuczna inteligencja potrafi najłatwiej przetwarzać? Do profilu. Do prawdopodobieństwa. Do segmentu. Do rekordu. Do użytkownika. Do funkcji celu.

Jeżeli tak się stanie, zwycięstwo maszyny nie będzie wymagało świadomości, buntu ani przejęcia władzy. Wystarczy reorganizacja instytucji wokół założenia, że model człowieka jest wystarczającym substytutem człowieka. Jeżeli natomiast państwo zachowa prawo obywatela do przedstawienia racji, zakwestionowania wyniku, zmiany zdania, odwołania się od predykcji i współdecydowania o funkcjach celu systemów publicznych, sztuczna inteligencja może stać się jednym z najpotężniejszych narzędzi demokratycznego państwa.

Parafrazując Ostrom, nie jesteśmy skazani na wybór pomiędzy wszechwładnym centrum a całkowitym chaosem. Ludzie wielokrotnie potrafili budować złożone, policentryczne instytucje zdolne ograniczać władzę bez niszczenia zdolności do wspólnego działania. To samo zadanie stoi dziś przed społeczeństwami wobec infrastruktury obliczeniowej. Parafrazując Simona, nie musimy być doskonale racjonalni, aby budować instytucje poprawiające jakość naszych decyzji. Parafrazując Webera, skuteczna administracja nie powinna zapominać, że racjonalność środka nie rozstrzyga racjonalności celu. Parafrazując Pettita, wolność nie polega wyłącznie na tym, że system chwilowo nie ingeruje; wymaga ochrony przed arbitralną możliwością ingerencji. Parafrazując Arendt, technologia nie zwalnia nas z obowiązku myślenia o tym, co robimy. I wreszcie, parafrazując Hobbesa, sztuczny człowiek powinien pozostać konstrukcją stworzoną przez ludzi po to, aby im służyć, a nie architekturą, do której ludzie muszą się dopasować.

Na tej granicy rozstrzyga się przyszłość Sztucznego Państwa. Nie pomiędzy analogiem a cyfrą, człowiekiem a robotem czy postępem a jego zatrzymaniem, lecz pomiędzy dwoma porządkami legitymizacji: takim, w którym maszyna zwiększa zdolność człowieka do działania, oraz takim, w którym człowiek jest rekonstruowany jako element umożliwiający maszynie skuteczniejsze działanie. Jeżeli demokratyczne państwo zachowa pierwszeństwo prawa nad kodem, osoby nad profilem, decyzji nad predykcją, osądu nad rekomendacją oraz obywatela nad użytkownikiem, Sztuczne Państwo pozostanie ostrzeżeniem i użytecznym idealnym typem, a nie opisem nieuchronnej przyszłości. Jeżeli tych różnic nie zachowa, może dojść do najbardziej osobliwej transformacji w dziejach nowoczesnego państwa: formalne instytucje demokracji pozostaną na swoim miejscu, podczas gdy coraz mniej będzie od nich zależało. Wtedy spełni się najważniejsza obawa całej genealogii – nie ta, że maszyna stanie się człowiekiem, lecz ta, że człowiek przestanie być potrzebny jako podmiot decyzji, ponieważ system uzna, że wystarczająco dobrze zna jego przewidywany model.

Obrona człowieka w epoce sztucznej inteligencji nie powinna polegać na obronie jego błędów, niewiedzy i irracjonalności. Powinna polegać na obronie czegoś znacznie bardziej

fundamentalnego: prawa do bycia czymś więcej niż sumą danych, na podstawie których można go przewidzieć. Państwo przyszłości może być sztucznie inteligentne. Nie powinno jednak stać się państwem, któremu człowiek jest potrzebny jedynie jako dane wejściowe.

Czy w świecie doskonałych predykcji zdołamy ocalić prawo do bycia czymś więcej niż sumą danych, które nas opisują? Prawdziwym triumfem maszyny nie będzie moment, w którym zyska ona świadomość i zbuntuje się przeciwko stwórcy, lecz chwila, gdy uznamy, że model człowieka jest wystarczającym substytutem samego człowieka. Wówczas państwo może pozostać formalnie demokratyczne, podczas gdy my staniemy się w nim jedynie pasywnymi danymi wejściowymi do systemu, który przestał potrzebować naszego osądu.

Inspiracje

[1] "The Rise and Fall of the Artificial State" Jill Lepore

Jill Lepore (ur. 27 sierpnia 1966) jest cenioną amerykańską historyczką, autorką i dziennikarką. Jest profesorem historii amerykańskiej im. Davida Woodsa Kempera (rocznik 1941) na Uniwersytecie Harvarda oraz profesorem prawa na Harvard Law School, a także regularnie pisze dla „The New Yorker”. Lepore uzyskała doktorat z amerykanistyki na Uniwersytecie Yale. Jej badania naukowe obejmują amerykańską historię polityczną, prawo, literaturę oraz metodologię historyczną i technologię dowodową. Lepore jest znana z udostępniania szerokiej publiczności wnikliwych analiz historycznych. Do jej przełomowych prac należą nagrodzone nagrodą Bancrofta „The Name of War”, „The Secret History of Wonder Woman” oraz „These Truths: A History of the United States”, wpływowa jednotomowa historia Ameryki. Dwukrotna finalistka Nagrody Pulitzera, w swojej twórczości krytycznie analizuje instytucje demokratyczne, braki w archiwach i przemiany społeczne.

Jill Lepore (born August 27, 1966) is an acclaimed American historian, author, and journalist. She serves as the David Woods Kemper '41 Professor of American History at Harvard University and Professor of Law at Harvard Law School, while also writing regularly as a staff writer for The New Yorker. Lepore holds a doctorate in American Studies from Yale University. Her scholarship spans American political history, law, literature, and the historical methodology and technology of evidence. Lepore is renowned for bringing rigorous historical analysis to broad public audiences. Her seminal works include the Bancroft Prize-winning The Name of War, The Secret History of Wonder Woman, and These Truths: A History of the United States, an influential single-volume history of America. A two-time Pulitzer Prize finalist, her work critically investigates democratic institutions, archival absences, and societal transformations.

Harvard University

Literatura uzupełniająca

Książki

Autorzy przywoływani w artykule:



[1] "If Then"

Jill Lepore

2020 | Hachette UK | ISBN: 9781529386189



[2] "Blindspot"

Jane Kamensky, Jill Lepore

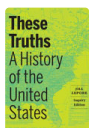
2009 | Random House | ISBN: 9780385526203



[3] "The Mansion of Happiness"

Jill Lepore

2012 | Vintage | ISBN: 9780307958501



[4] "These Truths"

Jill Lepore

2023 | ISBN: 9781324043799



[5] "The Story of America"

Jill Lepore

2012 | Princeton University Press | ISBN: 9780691159591



[6] "We the People"

Jill Lepore

2026 | John Murray Publishers | ISBN: 9781399827089

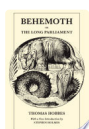
Literatura pokrewna (Google Scholar):



[7] "Leviathan"

Thomas Hobbes

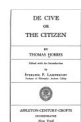
Barnes & Noble Publishing | ISBN: 9780760755938



[8] "Behemoth"

Thomas Hobbes

University of Chicago Press | ISBN: 9780226229843



[9] "De Cive"

Thomas Hobbes

[10] "The looping effects of human kinds"

Ian Hacking

[11] "Black–Scholes model"

Robert Merton



[12] "Basic concepts in sociology"

Max Weber
Citadel Press



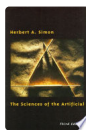
[13] "Models of my life"

Herbert Simon
1991



[14] "Human Problem Solving"

Herbert Simon
Prentice Hall



[15] "The Sciences of the Artificial"

Herbert Simon
MIT Press | ISBN: 9780262264495



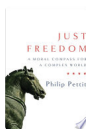
[16] "Legitimacy-Based Policing and the Promotion of Community Vitality"

Tom Tyler, Caroline Nobo
2023 | Cambridge University Press | ISBN: 9781009308038



[17] "The Economy of Esteem"

Philip Pettit
OUP Oxford | ISBN: 9780199289813



[18] "Just Freedom"

Philip Pettit
2014 | W. W. Norton & Company | ISBN: 9780393243017



[19] "A Political Philosophy in Public Life"

José Luis Martí, Philip Pettit
2012 | Princeton University Press | ISBN: 9781400835058



[20] "Rules, Games, and Common-pool Resources"

Elinor Ostrom
University of Michigan Press | ISBN: 9780472065462

[21] "The United States Bombing Surveys Summary Reports (European War) (Pacific War)"

John Kenneth Galbraith

[22] "Federalist No. 10"

James Madison

[23] "Memorial and Remonstrance"

James Madison

[24] "James Madison letter to James Robertson"

James Madison

Artykuły naukowe

Autorzy przywoływani:

[1] "3 Modern Humans and the Super-Brain For by Art is created that great Leviathan called a Common-wealth . . . which is but an Artificiall Man"

Thomas Hobbes

2011 | Landscape of the Mind | DOI: 10.7312/hoff14704-003

[Google Scholar](#)

[2] "Jean-Paul Sartre: Existential "Freedom" and the Political"

Yvonne Manzi, Thomas Hobbes, J. Locke, J. Mill

2020

[Google Scholar](#)

[3] "The looping effects of human kinds"

Ian Hacking

1996 | Oxford University Press eBooks | DOI: 10.1093/acprof:oso/9780198524021.003.0012

[Google Scholar](#)

[4] "The Human Condition"

Hannah Arendt, Margaret Canovan, Danielle Allen

1998 | DOI: 10.7208/chicago/9780226586748.001.0001

[Google Scholar](#)

[5] "Between past and future eight exercises in political thought"

Hannah Arendt

1987

[Google Scholar](#)

[6] "Lectures on Kant's Political Philosophy"

Hannah Arendt, Ronald S. Beiner

1982 | DOI: 10.7208/chicago/9780226231785.001.0001

[Google Scholar](#)

[7] "FROM 'POLITICS AS A VOCATION'"

Max Weber

1991 | Edinburgh University Press eBooks | DOI: 10.1515/9781474469654-020

[Google Scholar](#)

[8] "FROM 'BUREAUCRACY'"

Max Weber

2005 | Social Theory | DOI: 10.1515/9781474469654-021

[Google Scholar](#)

[9] "Germany's Future Form of State"

M. Weber

2021 | Max Weber Studies | DOI: 10.1353/max.2021.0000

[Google Scholar](#)

[10] "Basic Sociological Concepts: §8 – §17"

M. Weber

2020 | Sociological Problems Journal | DOI: 10.66384/33624196

[Google Scholar](#)

[11] "Four Primitives for Algorithmic Cognition Executable Relational Structures in Biologically Plausible Neural Systems"

Charles Simon

2026 | DOI: 10.2139/ssrn.7366892

[Google Scholar](#)

[12] "Toward a Model of Organizations as Interpretation Systems"

Richard L. Daft, Karl E. Weick

1984 | Academy of Management Review | DOI: 10.5465/amr.1984.4277657

[Google Scholar](#)

[13] "A Socially Embedded Model of Thriving at Work"

Gretchen Marie Spreitzer, Kathleen M. Sutcliffe, Jane E. Dutton, Scott B. Sonenshein, Adam M. Grant

2005 | Organization Science | DOI: 10.1287/orsc.1050.0153

[Google Scholar](#)

[14] "DISTINGUISHING CONTROL FROM LEARNING IN TOTAL QUALITY MANAGEMENT: A CONTINGENCY PERSPECTIVE"

Sim B. Sitkin, Kathleen M. Sutcliffe, Roger G. Schroeder

1994 | Academy of Management Review | DOI: 10.5465/amr.1994.9412271813

[Google Scholar](#)

[15] "Learning Through Rare Events: Significant Interruptions at the Baltimore & Ohio Railroad Museum"

Marlys K. Christianson, Maria T. Farkas, Kathleen M. Sutcliffe, Karl E. Weick

2008 | Organization Science | DOI: 10.1287/orsc.1080.0389

[Google Scholar](#)

[16] "Uncertainty in the transaction environment: an empirical test"

Kathleen M. Sutcliffe, Akbar Zaheer

1998 | Strategic Management Journal | DOI: 10.1002/(sici)1097-0266(199801)19:1<1::aid-smj938>3.0.co;2-5

[Google Scholar](#)

[17] "Catching crumbs from the table"

Ted Chiang

2000 | Nature | DOI: 10.1038/35014679

[Google Scholar](#)

[18] "Group Agency"

Christian List, Philip N. Pettit

2011 | Oxford University Press eBooks | DOI: 10.1093/acprof:oso/9780199591565.001.0001

[Google Scholar](#)

[19] "Republicanism: A Theory of Freedom and Government."

Erin L. Kelly, Philip N. Pettit

1999 | The Philosophical Review | DOI: 10.2307/2998263

[Google Scholar](#)

[20] "A Theory of Freedom: From the Psychology to the Politics of Agency"

Philip N. Pettit

2001

[Google Scholar](#)

[21] "Just Freedom: A Moral Compass for a Complex World"

Philip N. Pettit

2014

[Google Scholar](#)

[22] "Consumer Product Returns: High-Frequency Tracking Using Transaction Data"

David Bounie, Youssouf Camara, John W. Galbraith

2024 | DOI: 10.2139/ssrn.4977481

[Google Scholar](#)

[23] "The debates in the several state conventions on the adoption of the federal Constitution"

Jonathan Elliot, James H. Madison

2013 | Medical Entomology and Zoology

[Google Scholar](#)

[24] "Broken Heartland: The Rise of America's Rural Ghetto"

James H. Madison, Osha Gray Davidson

1997 | Western Historical Quarterly | DOI: 10.2307/969906

[Google Scholar](#)

[25] "The Federalist: A Commentary on the Constitution of the United States"

Alexander Hamilton, James H. Madison, John Jay, Edward G. Bourne

2016

[Google Scholar](#)

[26] "James Madison's Political and Constitutional Thought Reconsidered"

Stuart Leibiger, Garrett Ward Sheldon, J. C. A. Stagg, James Madison

2002 | The William and Mary Quarterly | DOI: 10.2307/3491669

[Google Scholar](#)

 Treść tworzy Fundacja Dobre Państwo.

Redaguje i publikuje APA ONE, autorski system redakcyjny Fundacji oparty na AI.

APA ONE Publishing System

Fundacja Dobre Państwo © 2026

Article ID: c54dd6ae-637b-4792-8c71-5040072f67do