

Przewaga adaptacyjna: od transformacji tożsamości do nowego humanizmu w erze AI



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje transformację organizacyjną nie jako prosty proces technologiczny, lecz jako głęboki kryzys tożsamości. Autor wprowadza pojęcie liminalności – stanu „pomiędzy”, w którym stare kompetencje i statusy tracą na znaczeniu, a nowe jeszcze nie zostały uformowane. W kontekście AI zjawisko to potęguje lęk przed utratą zawodowego prestiżu i sensu dotychczasowych inwestycji w rozwój osobisty. Tekst podkreśla, że największym wyzwaniem adaptacyjnym nie jest nauka nowych narzędzi, lecz zmiana odpowiedzi na pytanie o własną rolę w systemie. Współczesna liminalność często staje się przedłużonym stanem zawieszenia, co prowadzi do osłabienia identyfikacji grupowej i lęku, wymagając nowego podejścia do humanizmu w zarządzaniu.

The article analyzes organizational transformation not as a simple technological process, but as a profound identity crisis. The author introduces the concept of liminality—a 'between' state in which old competencies and statuses lose significance, while new ones have not yet been formed. In the context of AI, this phenomenon intensifies the fear of losing professional prestige and the meaning of previous investments in personal development. The text emphasizes that the greatest adaptive challenge is not learning new tools, but changing the answer to the question of one's own role in the system. Modern liminality often becomes a prolonged state of suspension, leading to weakened group identification and anxiety, requiring a new approach to humanism in management.

Artykuł analizuje fundamentalną zmianę paradygmatu ludzkiej sprawczości w obliczu rewolucji AI i potencjalnego nadejścia AGI. Autor stawia tezę, że w świecie, gdzie

inteligencja wykonawcza (zdolność do generowania odpowiedzi, optymalizacji i analizy) staje się tanim towarem powszechnym, przewaga adaptacyjna człowieka przestaje polegać na rywalizacji z maszyną w obszarze wydajności poznawczej. Zamiast tego, następuje przesunięcie roli człowieka ku sprawczości konstytucyjnej i ewaluacyjnej – obszarowi definiowania celów, nadawania sensu oraz ponoszenia etycznej odpowiedzialności za kierunek zmian. Tekst bada ten proces na trzech poziomach: psychologicznym (kryzys tożsamości w fazie liminalnej), instytucjonalnym (budowa państwa adaptacyjnego chroniącego wartości) oraz cywilizacyjnym. Głównym przesłaniem jest propozycja nowego humanizmu, w którym godność człowieka nie wynika z dominacji intelektualnej, lecz z pełnienia funkcji strażnika Dobra Wspólnego i arbitra wartości w trójkącie człowiek-AGI-instytucja.

SŁOWA KLUCZOWE

przewaga adaptacyjna

transformacja tożsamości

liminalność

nowy humanizm w erze AI

AGI

Meaningful Human Control

pułapka kompetencyjna

epistemiczna monokultura

trójkąt człowiek-AGI-instytucja

unlearning

Messy Middle

selekcja dziedzictwa

podział funkcji poznawczych

autonomia systemów AI

państwo adaptacyjne

Transformacja to przede wszystkim kryzys tożsamości

Opuszczanie Starego Miasta – liminalność, utrata tożsamości i niebezpieczna przestrzeń pomiędzy światami

Strategie transformacji mają osobliwą skłonność do pomijania czasu, w którym człowiek i organizacja przestają należeć do starego porządku, lecz nie przynależą jeszcze do nowego. Typowy diagram pokazuje punkt A, strzałkę i punkt B: zarząd ogłasza zmianę, konsultanci kreślą przyszły model operacyjny, pracownicy otrzymują nowe role, a prezentacja kończy się obrazem organizacji już po transformacji.

Tymczasem największa część realnego doświadczenia ukrywa się właśnie w strzałce. Stary system nie znika w chwili zatwierdzenia nowej strategii, a nowy nie zaczyna działać tylko dlatego, że nadano mu nazwę. Pomiędzy nimi otwiera się przestrzeń, w której stare kompetencje, statusy i

kryteria oceny tracą jednoznaczność, podczas gdy nowe nie zdołały jeszcze zyskać społecznego autorytetu.

John Sanei określa tę sytuację mianem Leaving the Old City oraz The In-Between Space. Oznacza ona opuszczenie „Starego Miasta” i wejście w obszar, w którym dotychczasowa tożsamość przestaje wystarczać, a nowa jeszcze się nie uformowała. W tym miejscu język futurystyki spotyka się z klasycznym aparatem pojęciowym antropologii. Arnold van Gennep, analizując na początku XX wieku rytuały przejścia, wskazywał na trójfazową strukturę zmiany statusu: separację od dotychczasowego położenia, fazę przejściową oraz ponowne włączenie do wspólnoty w nowej roli. Victor Turner rozwinął koncepcję tej środkowej fazy, nazywając ją liminalnością – od łacińskiego limen, czyli próg. W jego ujęciu jest to moment „pomiędzy” uporządkowanymi statusami; czas, w którym wcześniejsze klasyfikacje zostają zawieszone, a nowy porządek nie osiągnął jeszcze stabilności.

Nie oznacza to oczywiście, że firma zmieniająca model biznesowy odprawia rytuał przejścia w sensie antropologicznym – mechaniczne przenoszenie tych kategorii na grunt zarządzania byłoby metodologicznie nieostrożne. Analogia ta okazuje się jednak niezwykle płodna, gdyż pozwala dostrzec to, co w języku „wdrażania zmiany” często umyka: transformacja nie jest jedynie ruchem między dwiema konfiguracjami procesów, lecz również okresem dezorganizacji znaczeń. Człowiek nie pyta wtedy wyłącznie o to, czego ma się nauczyć, ale przede wszystkim o to, kim jest, skoro dotychczasowa odpowiedź przestała być wystarczająca.

Pracownik działu, którego zadania przejmuje AI, nie doświadcza jedynie zmiany technologicznej – traci on często fundament swojego zawodowego prestiżu. Menedżer, który przez dwie dekady był niezbędny jako kontroler przepływu informacji, może odkryć, że cyfrowa infrastruktura czyni jego funkcję pośredniczącą zbędną. Ekspert, którego wartość budowała pamięć i znajomość procedur, staje przed systemem zdolnym w sekundy syntetyzować dokumentację, której opanowanie zajęło mu lata. Organizacyjna transformacja staje się wówczas biograficznym pytaniem o sens wcześniejszej inwestycji w siebie.

Badania nad zmianą organizacyjną dowodzą, że taka nieciągłość tożsamościowa nie jest jedynie figurą literacką. Analizy strategicznych programów transformacyjnych wskazują, że emocjonalne przywiązanie do dawnej tożsamości organizacji oraz sposób komunikowania ciągłości lub zerwania z przeszłości istotnie wpływają na przeżycia pracowników. Z kolei studia nad fuzjami przedsiębiorstw opisują ludzi znajdujących się „pomiędzy” dawnymi a przyszłymi rolami, pozbawionych markerów, które wcześniej definiowały ich zawodzką pozycję i znaczenie. Sanei trafnie diagnozuje zatem problem: zmiana kompetencji staje się najtrudniejsza wtedy, gdy wymaga jednoczesnej zmiany odpowiedzi na pytanie: „kim jestem w tym systemie?”.

Pułapka liminalności i lęk przed utratą statusu w procesie zmiany

W rytuałach badanych przez antropologów faza liminalna była ograniczona i posiadała strukturę. Człowiek wiedział, że znajduje się w przejściu, a wspólnota dysponowała symbolami i regułami umożliwiającymi przeprowadzenie go do nowego statusu. Współczesna transformacja organizacyjna często nie oferuje podobnej pewności. Nie wiadomo, jak długo potrwa okres przejściowy, czy nowy model faktycznie zwycięży oraz czy rola, do której człowiek ma się przygotowywać, w ogóle będzie istnieć za trzy lata. Liminalność przestaje wtedy być krótkim progiem i może stać się środowiskiem pracy.

Literatura organizacyjna rozpoznaje właśnie tę możliwość "przedłużonego bycia pomiędzy". Badania nad liminalnością rozróżniają transformacyjne przejście od sytuacji, w której aktorzy przez długi czas pozostają między dwiema pozycjami tożsamościowymi, bez jednoznacznej reintegracji. Nowsze próby pomiaru subiektywnej liminalności wiążą poczucie bycia "w zawieszeniu" z niejednoznacznością, lękiem oraz osłabieniem identyfikacji grupowej. Jest to ważne uzupełnienie optymistycznej narracji adaptacyjnej. Przestrzeń pomiędzy może otwierać nowe możliwości, lecz człowiek nie jest zobowiązany doświadczać jej jako ekscytującej wolności; może odczuwać ją jako utratę orientacji. Właśnie dlatego materiał Saneia wprowadza pojęcie Messy Middle - "Brudnego Środka".

To faza, w której stara strategia nie dostarcza już wcześniejszych wyników, podczas gdy nowa inwestycja jeszcze ich nie generuje. Źródło opisuje w tym okresie narastającą presję na powrót do dawnych nawyków, mikrozarządzania i odzyskiwania pozornego poczucia kontroli. Jest to jedna z najbardziej trafnych obserwacji całej konstrukcji Saneia, jeśli oczyścić ją z nadmiernej terminologii neurobiologicznej. Regresja do starego modelu nie musi bowiem wynikać z "uzależnienia" czy określonego pasma fal mózgowych. Może być racjonalną reakcją ludzi zmuszonych działać w systemie, którego nowa logika nie dostarcza jeszcze ani wyników, ani tożsamości, ani jasnych kryteriów sukcesu. Stare Miasto posiadało mapę. Było wiadomo, kto jest ekspertem, jak wygląda awans, które działania są nagradzane, jakie liczby świadczą o powodzeniu i czego należy nauczyć nowego pracownika.

Można było tego świata nie lubić, lecz był przewidywalny. Przestrzeń pomiędzy pozbawia człowieka części tych punktów orientacyjnych. Wczorajsza kompetencja może nadal być użyteczna, ale już nie wystarcza. Jutrzejsza wydaje się ważna, ale jej wartość nie została jeszcze w pełni potwierdzona. Powstaje charakterystyczne napięcie: trzeba kontynuować pracę starego systemu i jednocześnie przestać utożsamiać swoją przyszłość z jego trwałością.

W kategoriach antropologicznych jest to właśnie doświadczenie progu. Turner podkreślał, że liminalność wiąże się z czasowym osłabieniem utrwalonych klasyfikacji i hierarchii, a w niektórych rytuałach może sprzyjać powstawaniu *communitas* – intensywnego doświadczenia wspólnotowości pomiędzy uczestnikami znajdującymi się poza zwykłą strukturą statusu. W przedsiębiorstwie analogiczny potencjał może pojawić się w zespołach transformacyjnych, gdy osoby reprezentujące odmienne funkcje zostają chwilowo wyjęte z tradycyjnych silosów i pracują nad problemem, którego żadna z dotychczasowych hierarchii nie potrafi samodzielnie rozwiązać. Dyrektor, programista, prawnik, projektant i analityk mogą przez pewien czas spotkać się nie jako przedstawiciele pozycji w schemacie organizacyjnym, lecz jako ludzie próbujący zrozumieć wspólną niewiadomą.

Nie należy jednak romantyzować *communitas*. Liminalność może równie łatwo generować konflikt. Analizy zmiany organizacyjnej inspirowane Turnerem opisują okresy transformacji jako "social drama", w którym dotychczasowy porządek zostaje zakwestionowany, konflikty interesów stają się bardziej widoczne, a organizacja musi później odnaleźć sposób ponownego ustanowienia ładu. Transformacja obnaża to, co stabilność potrafiła przykrywać. Gdy nie wiadomo jeszcze, które role przetrwają, spór o strategię staje się jednocześnie sporem o przyszły rozdział statusu, budżetów i wpływu. Dlatego bardzo często najostrejsze konflikty okresu transformacji przedstawiane są jako różnice poglądów technicznych, choć pod powierzchnią kryją się pytania egzystencjalne: czy mój dział będzie nadal potrzebny? Czy nowy system zmniejszy znaczenie kompetencji, którą budowałem przez piętnaście lat? Czy "automatyzacja procesu" oznacza, że organizacja przestanie potrzebować ludzi mojego typu? Czy nowy model promuje innych zwycięzców niż poprzedni?

Adaptacja to selekcja dziedzictwa, a nie permanentna dezorientacja

Dopóki tych pytań nie można bezpiecznie wypowiedzieć, organizacja analizuje transformację jedynie na poziomie technologii, podczas gdy prawdziwy konflikt toczy się o przyszły rozdział znaczeń. Tutaj najwyraźniej ujawnia się ograniczenie popularnego hasła o „wychodzeniu ze strefy komfortu”. Nie każda utrata bezpieczeństwa ma charakter rozwojowy; człowiek potrzebuje bowiem pewnego poziomu stabilności, by móc podejmować ryzyko poznawcze.

Organizacja, która jednocześnie zmienia strukturę, ludzi, technologię, kryteria oceny i model biznesowy, ryzykuje wywołaniem nie twórczej liminalności, lecz permanentnej dezorientacji. Adaptacyjność nie polega na utrzymywaniu wszystkich w nieustannym stanie przejścia, lecz na konstruowaniu bezpiecznych progów, przez które można przejść. W tym kontekście

antropologiczna trójfazowość van Gennepa staje się niezwykle wartościową heurystyką zarządzania zmianą. Pierwszym etapem jest separacja: organizacja musi wyraźnie określić, co naprawdę się kończy. Bez tego członkowie systemu słyszą komunikat o zmianie, lecz wciąż są rozliczani według reguł starego świata. Drugim jest liminalność: konieczne jest stworzenie przestrzeni do eksperymentów i nauki, w której role nie są jeszcze sztywno ustalone. Trzecim zaś reintegracja: nowe praktyki muszą zostać włączone w strukturę organizacji, otrzymać zasoby, kompetencje, system odpowiedzialności oraz symboliczną legitymizację.

Wiele transformacji zatrzymuje się właśnie przed tym ostatnim krokiem. Organizacje tworzą laboratoria, pilotaże i zespoły specjalne, lecz nie potrafią dokonać reintegracji. Powstaje permanentny półświatek innowacji – ciekawy, medialny, ale słabo połączony z centrum przedsiębiorstwa. Zespół od lat pracuje „nad przyszłością”, podczas gdy rdzeń organizacji działa według starych reguł. Liminalność, która miała być przejściem, staje się stanem chronicznym, co jest niebezpieczne z dwóch powodów.

Po pierwsze, jednostka eksperymentalna może wykształcić tożsamość opartą na byciu „inną” niż reszta firmy. Jej członkowie zaczynają potrzebować starej organizacji jako negatywnego punktu odniesienia; gdyby ich rozwiązania rzeczywiście zintegrowano, utraciliby status buntowników. Po drugie, organizacja operacyjna może zacząć traktować transformację jak odrębny teatr: przyszłość zostaje zamknięta w laboratorium, dzięki czemu centrum nie musi zmieniać siebie. W ten sposób instytucjonalizuje się przejście, które nigdy nie prowadzi do celu.

Dlatego Leaving the Old City nie może polegać wyłącznie na krytyce starego porządku. Człowiek opuszczający miasto potrzebuje odpowiedzi na pytanie, co warto z niego zabrać. Rewolucyjna retoryka często przedstawia przeszłość jedynie jako zbiór ograniczeń, tymczasem stare instytucje przechowują kapitał społeczny, wiedzę ukrytą, relacje, pamięć błędów oraz normy odpowiedzialności i kompetencje trudne do odtworzenia.

Adaptacja nie jest ucieczką z płonącego miasta z pustymi rękami. Jest selekcją dziedzictwa. Należy tutaj odróżnić ciągłość od konserwatyzmu. Badania tożsamości organizacyjnej wskazują, że podczas strategicznej zmiany zdolność zakomunikowania pewnej ciągłości pomaga ludziom przejść pomiędzy dawną a nową tożsamością. Człowiek łatwiej akceptuje radykalną zmianę formy, jeśli odnajdzie sensowną odpowiedź na pytanie, co z wcześniejszego „my” pozostaje nadal ważne. Microsoft może zmienić model biznesowy, nie ogłaszając jednocześnie, że kompetencja techniczna przestała mieć wartość. Bank może automatyzować procesy, nie rezygnując z odpowiedzialności za bezpieczeństwo środków. Administracja może cyfryzować procedury, zachowując zasadę praworządności.

Transformacja jako rekonstrukcja tożsamości, a nie jej utrata

Transformacja, która nie potrafi wskazać ciągłości wartości, może zostać przeżyta jako amputacja pamięci. Ma to fundamentalne znaczenie dla jednostki. Adaptacyjność nie wymaga przecież od eksperta całkowitego „zabicia” dawnej tożsamości. Choć Sanei posługuje się dramatycznym językiem psychologicznej śmierci starego „ja”, a materiały źródłowe kreślą obraz lidera gotowego porzucić tożsamość zbudowaną na dawnym sukcesie, bardziej precyzyjne byłoby mówienie o rekonstrukcji tożsamości. Człowiek nie przestaje być tym, kim był; zmienia jedynie relację między przeszłym doświadczeniem a przyszłym działaniem. Były ekspert od rozwiązania A może stać się osobą najlepiej rozumiejącą, dlaczego A działało i w jakich warunkach przestaje wystarczać. To ogromna wartość, o ile przeszłość przestanie być prawem weta.

Liminalność ma zatem podwójny charakter. Z jednej strony jest momentem zwiększonej plastyczności, gdyż stare klasyfikacje tracą swą władzę. Z drugiej – okresem podatności na lęk, regresję i walkę o status, ponieważ nowa struktura nie zapewnia jeszcze poczucia bezpieczeństwa. Właśnie dlatego współczesne badania organizacyjne analizują zmianę nie tylko jako proces proceduralny, lecz jako pracę nad tożsamością w czasie, gdy przeszłość, teraźniejszość i wyobrażona przyszłość pozostają w napięciu. Dojrzałe przywództwo w okresie liminalnym nie polega zatem na naiwnym zapewnianiu, że wszystko będzie dobrze – takiej wiedzy często po prostu nie ma. Nie powinno ono również sztucznie przyspieszać deklarowania nowej tożsamości, zanim praktyka zdoła ją uwiarygodnić. Zadaniem lidera jest coś trudniejszego: stworzyć strukturę na tyle stabilną, by niepewność nie zamieniła się w chaos, a jednocześnie zachować otwartość, aby przyszłość nie została przedwcześnie zamknięta. W takim ujęciu In-Between Space przestaje być luką między dwiema strategiami, a staje się strategicznym zasobem czasu. To właśnie tam można przetestować nowe role, języki, technologie i modele współpracy.

Ten zasób ma jednak swoją cenę. Nie można w nieskończoność utrzymywać człowieka i instytucji w zawieszeniu. Każda transformacja musi zatem zawierać nie tylko plan odejścia od starego modelu, ale i warunki ponownego zakorzenienia: jasne kryteria uznania eksperymentu za nowy standard, podział odpowiedzialności, obowiązujące prawa oraz określenie, które dawne normy zostają zachowane i jakie nowe źródła statusu powstaną dla ludzi. Stare Miasto nie było przecież tylko więzieniem – było domem. Dlatego opuszczanie go boli bardziej, niż sugerują to korporacyjne prezentacje o zarządzaniu zmianą. W jego murach zakłute były kompetencje, wspomnienia, hierarchie i znaczenia. Nie wystarczy powiedzieć człowiekowi, że za horyzontem czeka przyszłość. Trzeba zrozumieć, że prosi się go o wejście w obszar, w którym część dotychczasowych odpowiedzi

na pytanie „kim jestem?” przestaje działać. Dopiero po przejściu przez tę antropologiczną warstwę można wrócić do ekonomii i polityki instytucjonalnej bez redukowania człowieka do funkcji produkcji.

Państwo adaptacyjne chroni wartości przed nadmierną zmiennością

Adaptacyjność nie jest bowiem zdolnością do bezwarunkowej zmienności. Człowiek i wspólnota potrzebują punktów stałych, aby móc ewoluować bez ryzyka rozpadu. Pojawia się zatem pytanie, którego teoria organizacji nie rozwiąże samodzielnie: co powinno pozostać niezmiennie? Jakie prawa, wartości i zobowiązania nie powinny zmieniać się wraz z każdą technologiczną falą?

W tym miejscu przewaga adaptacyjna musi zostać skonfrontowana z instytucją szczególną – państwem. Jego zadaniem nie jest wyłącznie nadążanie za przyszłością, lecz przede wszystkim ochrona ciągłości wspólnoty i Dobra Wspólnego w momentach przejścia przez krytyczne progi.

Przeniesienie aparatu analitycznego Johna Saneia z organizacji biznesowej na grunt państwowy wymaga jednak metodologicznego zastrzeżenia. Materiał źródłowy koncentruje się bowiem na liderach, przedsiębiorstwach i zespołach „Dzisiaj i Jutro”, analizując pułapkę kompetencji, ekonomię uczenia oraz strategiczną alokację zasobów. Państwo nie zostało w nim rozwinięte jako odrębna teoria polityczna.

To, co następuje, jest zatem interpretacyjnym rozszerzeniem tego aparatu: próbą odpowiedzi na pytanie, co stanie się, gdy Speed Mismatch, unlearning, Today/Tomorrow, Transition Layer oraz Economies of Learning potraktujemy nie jako język zarządzania firmą, lecz jako narzędzia badania instytucji publicznych. Odpowiedź ta musi być jednak ostrożna.

Państwo nie jest przedsiębiorstwem, obywatel nie jest klientem, a konstytucja nie jest modelem biznesowym podlegającym kwartalnemu pivotowi. W demokratycznym państwie prawnym pewien stopień powolności został zaprojektowany świadomie. Ustawa wymaga procedury, decyzja administracyjna – podstawy prawnej, ograniczenie praw jednostki – uzasadnienia i kontroli, a rozstrzygnięcie władzy musi podlegać możliwości zakwestionowania. Sądy mają być niezależne właśnie po to, by nie reagowały na każdą zmianę politycznego nastroju; stabilność prawa pozwala bowiem obywatelom planować własne życie. Praworządność jest w znacznej mierze instytucjonalną technologią ograniczania arbitralnej adaptacyjności władzy.

Polska Konstytucja formułuje tę konstrukcję wyjątkowo silnie już w pierwszych artykułach. Rzeczpospolita jest określona jako dobro wspólne wszystkich obywateli, demokratyczne państwo prawne oraz system, w którym organy władzy publicznej działają na podstawie i w granicach prawa. Preambuła łączy sprawność instytucji z poszanowaniem wolności, sprawiedliwości, dialogu społecznego, pomocniczości, godności i solidarności.

Oznacza to, że sprawne państwo nie może być definiowane wyłącznie przez prędkość. Instytucja może działać niezwykle szybko, a jednocześnie źle, jeśli tempo to osiąga poprzez eliminację prawa do wysłuchania, kontroli, odwołania czy równego traktowania. Stąd zasadnicze rozróżnienie: państwo adaptacyjne nie jest tym, które najszybciej zmienia prawo, lecz tym, które najszybciej rozpoznaje, że jego diagnoza rzeczywistości wymaga rewizji i potrafi przeprowadzić korektę bez niszczenia praworządności. Jego przewagą nie jest ustawodawcza hiperaktywność, lecz krótki dystans między pojawieniem się istotnego sygnału a rozpoczęciem rzetelnego procesu poznawczego.

Taka perspektywa pozwala przełożyć Saneiove Speed Mismatch na język administracji. Luka nie musi przebiegać w prosty sposób pomiędzy „szybką technologią” a „wolnym prawem” – prawo z zasady może być wolniejsze niż kod komputerowy. Problem pojawia się wtedy, gdy rzeczywistość społeczna zmienia kategorię, a administracja przez lata nadal mierzy ją miarami poprzedniej epoki.

Dzieje się tak, gdy gospodarka platformowa jest oceniana wyłącznie przez pryzmat konstrukcji stworzonych dla fabrycznej organizacji pracy, nowe modele finansowe przy pomocy klasyfikacji zaprojektowanych przed ich powstaniem, a generatywna AI poprzez przepisy konstruowane dla oprogramowania wykonującego z góry przewidziany zestaw instrukcji.

Państwo musi przełamać pułapkę kompetencyjną poprzez instytucjonalizację foresightu

Niedopasowanie staje się w tym ujęciu nie tyle opóźnieniem legislacyjnym, co opóźnieniem pojęciowym. Państwo może dysponować tysiącami urzędników, uniwersytetami, instytutami badawczymi czy rozbudowaną statystyką publiczną i systemami informacyjnymi, a mimo to wpadać w pułapkę własnej Competence Trap. Utrwalona wiedza administracyjna narzuca bowiem sztywne kategorie: formularze, rozdziały budżetowe, kompetencje organów oraz granice między resortami. Problem obywatela musi najpierw gdzieś „pasować”, aby administracja mogła go procedować. Jeśli jednak nowe zjawisko przebiega w poprzek dawnych klasyfikacji, zostaje sfragmentaryzowane. Jedna jego część staje się domeną prawa pracy, druga podatków, trzecia

ochrony danych, czwarta rynku cyfrowego, a piąta polityki migracyjnej. W efekcie każdy urząd może poprawnie analizować swój wycinek, podczas gdy całość pozostaje niczyja. To państwowa wersja błędu opisanego wcześniej w kontekście organizacji: lokalna racjonalność nie gwarantuje racjonalności systemowej. Każdy departament może wykonywać swoją pracę zgodnie z procedurą, lecz obywatel nadal porusza się po systemie, który z jego perspektywy nie generuje spójnej odpowiedzi. Im bardziej wyspecjalizowana jest administracja, tym silniejsza staje się potrzeba powołania instytucji zdolnych pracować na granicach tych specjalizacji. Państwowy odpowiednik Horizon Scout nie byłby zatem futurystą przepowiadającym kolejne technologie.

Byłby raczej tłumaczem pomiędzy nauką, prawem, gospodarką i doświadczeniem obywateli a aparatem wykonawczym państwa. W tym punkcie aparat Saneia spotyka się z nurtem, który OECD określa jako *anticipatory governance*. Definiuje go on jako zdolność włączania foresightu, eksperymentowania, innowacji i ciągłego uczenia się do podstawowych procesów publicznego zarządzania. Organizacja wskazuje jednocześnie, że systematyczne wykorzystywanie tych zdolności nie jest jeszcze powszechne, a informacje o możliwej przyszłości często pozostają zbyt słabo powiązane z rzeczywistym procesem tworzenia polityk. Jest to niemal dokładny odpowiednik problemu *Transition Layer* w sferze publicznej: państwo może dysponować analizami przyszłości, ale brak mu mechanizmu, który zamieniałby je w korektę dzisiejszych decyzji. Raport OECD z 2025 roku dotyczący Litwy, Włoch i Malty opisuje próby instytucjonalizacji foresightu w administracji, wskazując między innymi na potrzebę trwałych struktur, kompetencji oraz sieci wymiany wiedzy. Podkreśla konieczność włączenia myślenia przyszłościowego do podstawowych funkcji państwa, zamiast pozostawiania go w odizolowanych „wyspach doskonałości”. Jest to ujęcie bardzo bliskie krytyce laboratoriów innowacji przedstawionej wcześniej. Jeśli jednostka foresightowa przygotowuje fascynujące scenariusze, ale budżet, legislacja i administracja sektorowa nie mają obowiązku ich rozważenia, Jutro staje się instytucjonalnym teatrem przyszłości.

Można zatem mówić o swoistej konstytucyjnej oburęczności państwa. Jedna ręka musi zapewniać Dzisiaj: ciągłość usług publicznych, wypłatę świadczeń, bezpieczeństwo, pobór podatków oraz funkcjonowanie sądów, szkół, infrastruktury i administracji. Tutaj kluczowe są powtarzalność, odpowiedzialność, legalność i niezawodność. Druga ręka powinna badać Jutro: demografię, automatyzację, nowe formy zatrudnienia, transformację energetyczną, biotechnologię, migracje, bezpieczeństwo cybernetyczne, zmiany klimatyczne i konsekwencje AI.

Adaptacyjna regulacja jako proces uczenia się państwa

Państwo nie może jednak otrzymać prawa do dowolnego eksperymentowania na obywatelu. Państwowy zespół Jutra działa w granicach konstytucyjnych, ustanowionych po to, aby przyszłość nie była projektowana kosztem godności człowieka. Szczególnie interesującym narzędziem stają się zatem regulacyjne środowiska eksperymentalne, czyli sandboxy. Ich sens nie polega na zawieszeniu prawa dla innowatorów, lecz na stworzeniu kontrolowanej przestrzeni, w której technologia, regulator i podmiot wdrażający mogą uczyć się od siebie nawzajem przed pełnym skalowaniem rozwiązania. Aktualny tekst unijnego AI Act przewiduje krajowe piaskownice regulacyjne dla sztucznej inteligencji; po zmianach przyjętych w 2026 roku termin zapewnienia co najmniej jednej takiej piaskownicy na poziomie krajowym przesunięto do 2 sierpnia 2027 roku. Przepisy zakładają środowisko ograniczone w czasie, plan eksperymentu, nadzór właściwych organów oraz mechanizmy pozwalające badać m.in. dokładność, odporność, cyberbezpieczeństwo i zagrożenia dla praw podstawowych.

Jest to doskonały przykład różnicy między deregulacją a regulacją adaptacyjną. Pierwsza sugeruje: skoro nie znamy skutków nowej technologii, usuńmy bariery i pozwólmy rynkowi odkryć rozwiązanie. Druga odpowiada: skoro nie znamy skutków, potrzebujemy tańszego mechanizmu zdobycia wiedzy przed podjęciem decyzji nieodwracalnej. Jest to dokładnie logika ekonomii uczenia się przeniesiona z przedsiębiorstwa do prawa. Jednocześnie AI Act pokazuje, że samo państwo i regulator również znajdują się w procesie nauki. Unijne przepisy dotyczące AI weszły w życie w 2024 roku i zostały rozłożone na etapy; 2 sierpnia 2026 roku rozpoczęło się stosowanie kolejnej szerokiej części regulacji oraz nowych obowiązków przejrzystości, natomiast w 2026 roku część terminów dotyczących systemów wysokiego ryzyka została ponownie przesunięta w ramach zmian legislacyjnych. Można krytykować konkretne decyzje, zakres obowiązków czy korekty dat, lecz sama architektura systemu ujawnia istotne zjawisko: regulowanie przełomowej technologii nie jest jednorazowym aktem wiedzy. Jest procesem stawiania kolejnych hipotez, ustalania standardów, doświadczeń wdrożeniowych i korekt.

Państwo adaptacyjne wymaga więc prawa zdolnego do uczenia się, lecz nie prawa płynnego. Nadmiernie częste zmiany regulacyjne podnoszą koszty transakcyjne, osłabiają zaufanie i premią podmioty dysponujące największymi zasobami prawnymi, zdolne do ciągłego dostosowywania się. Stabilność jest dobrem gospodarczym. Dlatego mechanizmy adaptacyjne powinny opierać się nie na ustawicznym przepisywaniu norm, lecz na okresowych przeglądach, ewaluacjach skutków, pilotażach, klauzulach przeglądowych, kontrolowanych eksperymentach oraz sprawniejszych systemach gromadzenia sygnałów.

Spółeczeństwo obywatelskie jako rozproszony aparat percepcji państwa

W tym miejscu szczególnego znaczenia nabiera społeczeństwo obywatelskie. Państwo nie może być własnym jedynym sensorem. Hayekowska intuicja o rozproszeniu wiedzy pozostaje aktualna także poza sporem o rynek: ogromna część informacji o skutkach regulacji tkwi w lokalnym doświadczeniu obywateli, przedsiębiorstw, organizacji społecznych, samorządów czy związków zawodowych – tam, gdzie centrum administracyjne nigdy nie dotrze bezpośrednio. Elinor Ostrom pokazywała z kolei, że zdolność rozwiązywania problemów zbiorowych często zależy od wielopoziomowych, policentrycznych struktur uczenia się i współdziałania, a nie wyłącznie od jednego ośrodka wiedzy. Można zatem stwierdzić, że nowoczesne państwo potrzebuje rozproszonego aparatu percepcji.

Petycja, konsultacja publiczna, sygnał sygnalisty, orzeczenie sądu, dane samorządu czy raport organizacji obywatelskiej można traktować jako różne typy sensorów. Ich znaczenie nie polega jednak na tym, że każdy impuls powinien automatycznie wymuszać zmianę prawa; byłaby to polityczna wersja nadmiernej reaktywności.

Chodzi raczej o to, aby system potrafił odróżnić sygnał od szumu, agregować rozproszone informacje i przekazywać je tam, gdzie staną się przedmiotem procedury decyzyjnej. W tym kontekście warto przywołać projekt „Szkatułki Kosztowności” Fundacji Dobre Państwo. Przedstawia on zasób jako zbiór analiz z obszaru historii, filozofii i szeroko rozumianej wiedzy – od asymetrii informacji i automatyzacji po kapitał ludzki czy instytucjonalne źródła szumu decyzyjnego. Nie jest to zatem zwykłe archiwum stanowisk w jednej sprawie. Z punktu widzenia omawianej teorii Szkatułkę można interpretować jako obywatelską warstwę sensing: próbę gromadzenia modeli, analogii i języków opisu, które poszerzają liczbę hipotez dostępnych w debacie o państwie. Szczególnie wymowny jest tekst „Ekonomia godności: AI, automatyzacja i rola państwa”, który odrzuca założenie, że postęp technologiczny automatycznie oznacza postęp społeczny.

Sposób wykorzystania AI zależy bowiem od bodźców, instytucji oraz reguł określających, kto czerpie korzyści, kto ponosi koszty i jakie cele uznaje się za wartościowe. W tym miejscu Saneiowa adaptacyjność wymaga uzupełnienia aksjologicznego. Firma może zdefiniować przetrwanie lub wartość przedsiębiorstwa jako nadrzędny test skuteczności. Państwo nie może. Musi pytać nie tylko o to, czy potrafi się dostosować, ale przede wszystkim: do czego i w czym interesie ma to robić?

Problem ten formułuje bezpośrednio tekst o dobru wspólnym w polskim konstytucjonalizmie, zamieszczony w Sztatule. Punktem wyjścia jest tam rozumienie bonum commune nie jako prostej sumy interesów, lecz jako zespołu warunków instytucjonalnych pozwalających ludziom rozwijać się we wspólnocie. Niezależnie od sporów filozoficznych, pojęcie to wprowadza kluczowe ograniczenie dla teorii adaptacyjności: nie wszystko, co zwiększa sprawność systemu, zwiększa dobro wspólne.

Państwo może na przykład wykorzystać AI do radykalnego przyspieszenia administracyjnego orzekania. Jeśli jednak wzrost prędkości nastąpi kosztem zrozumienia decyzji, prawa do jej zakwestionowania lub równego traktowania obywateli, system stanie się bardziej wydajny, ale niekoniecznie lepszy. Algorytm może zredukować szum decyzyjny, lecz jednocześnie utrwalić systematyczny bias. Cyfrowa administracja może obniżyć koszt usługi dla większości, zwiększając jednocześnie wykluczenie osób pozbawionych kompetencji cyfrowych.

Adaptacyjność państwa zakotwiczona w wartościach i krytyce

Adaptacyjność publiczna wymaga wielowymiarowego rachunku, w którym efektywność jest tylko jedną z wielu wartości. To właśnie odróżnia Adaptive Intelligence państwa od inteligencji oportunistycznej. Państwo oportunistyczne reaguje na presję najsilniejszego bodźca; państwo adaptacyjne zachowuje zdolność uczenia się, lecz jego zmiany pozostają zakotwiczone w konstytucyjnej hierarchii wartości. Nie wszystko może być przedmiotem eksperymentu. Godność człowieka, legalizm, równość wobec prawa, możliwość kontroli władzy czy ochrona podstawowych wolności nie są kwartalnymi hipotezami biznesowymi. Warto tu przywołać Poppera.

Jego idea cząstkowej inżynierii społecznej może być odczytana jako intelektualny przodek adaptacyjnego podejścia do polityki: zamiast projektować raz na zawsze doskonałe społeczeństwo, wprowadza się ograniczone reformy, obserwuje ich skutki, identyfikuje błędy i zachowuje możliwość korekty. Warunkiem jest jednak istnienie instytucji krytyki. System polityczny, który eksperymentuje, ale uniemożliwia zakwestionowanie tych działań, nie uczy się – jedynie zmienia rzeczywistość metodą prób bez sprzężenia zwrotnego. Dlatego najważniejszym aktywnym państwa adaptacyjnego staje się zdolność do instytucjonalizowania sprzeciwu wobec własnych hipotez. Niezależne sądy, wolne media, nauka, organizacje społeczne, opozycja parlamentarna, organy kontroli, konsultacje i prawo petycji nie są przeszkodami w sprawnym zarządzaniu; są kanałami informacji o błędach. W krótkim okresie mogą spowalniać decyzje, lecz w długim dramatycznie obniżają koszt systemowej pomyłki.

Demokracja posiada w tym sensie pewną przewagę adaptacyjną, choć nie jest ona automatyczna. Może generować hałas, krótkoterminowość i polaryzację, ale oferuje możliwość pokojowej korekty przywództwa, ujawniania błędów oraz konkurencji modeli interpretacyjnych. Autokracja może podejmować decyzje szybciej, lecz jeśli wytworzy wysoką cenę przekazywania złych informacji do centrum, jej prędkość staje się pozorem. System zaczyna poruszać się szybko na podstawie coraz gorszej mapy. Państwowy odpowiednik psychologicznego bezpieczeństwa nie polega więc na komforcie urzędników, lecz na tym, aby informacja sprzeczna z oczekiwaniem hierarchii mogła przetrwać drogę do decydenta. W przedsiębiorstwie pytaliśmy o cenę przekazywania złych wiadomości prezesowi; w państwie trzeba zapytać o cenę przekazywania ich ministrowi, premierowi, parlamentowi i opinii publicznej.

System nagradzający wyłącznie informacje zgodne z linią polityczną tworzy własną wersję High Beta – nie przez neurobiologię, lecz przez instytucjonalny strach przed prawdą. Szkatułka Kosztowności może być pod tym względem przykładem tworzenia zewnętrznej pamięci i alternatywnego aparatu poznawczego. Jej profesjonalizm nie powinien być mierzony liczbą tez zgodnych z bieżącą polityką Fundacji, lecz zdolnością do przechowywania argumentów, modeli, badań i sporów, do których można wrócić, gdy rzeczywistość postawi nowe pytanie. Zasób obywatelski zyskuje największą wartość nie wtedy, gdy dostarcza jednej doktryny, ale gdy poszerza liczbę sensownych modeli dostępnych dla wspólnoty politycznej. Połączenie analiz ekonomicznych, prawnych, historycznych, psychologicznych i filozoficznych jest tu funkcjonalne: złożone państwo rzadko psuje się w granicach jednej dyscypliny.

Można stwierdzić, że demokratyczne państwo potrzebuje odpowiednika Saneiowych zespołów Dzisiaj i Jutro także poza własną administracją. Instytucje publiczne zapewniają ciągłość Dzisiaj; nauka, społeczeństwo obywatelskie, think tanki, organizacje zawodowe, przedsiębiorcy, związki społeczne i obywatele produkują równoległe obrazy Jutra. Zadaniem państwa nie jest podporządkowanie ich jednemu centrum, lecz stworzenie warstwy przejściowej, która pozwoli wiedzy społecznej wejść do procesu stanowienia prawa bez przejmowania państwa przez najsilniejszą grupę interesu. Dobro Wspólne działa tutaj jak funkcja celu wyższego rzędu. Nie rozstrzyga automatycznie każdego sporu, gdyż obywatele będą różnić się w ocenie tego, co mu służy.

Wyznacza jednak ważną granicę: adaptacja instytucji nie może być mierzona wyłącznie przetrwaniem administracji czy skutecznością rządzących. Państwo istnieje dla wspólnoty politycznej, a jego zdolność uczenia powinna poprawiać warunki wolności, bezpieczeństwa, sprawiedliwości, zaufania i uczestnictwa obywateli, a nie tylko sprawność aparatu. Największym zagrożeniem nie jest więc sama powolność – czasami jest ona mądrością proceduralną.

Zagrożeniem jest powolność poznawcza połączona z szybką pewnością polityczną: system bardzo długo nie aktualizuje obrazu świata, lecz bardzo szybko ogłasza, że zna rozwiązanie. Przeciwnieństwem nie jest państwo zarządzane jak startup, lecz takie, które potrafi być konserwatywne wobec praw i eksperymentalne wobec instrumentów; stabilne wobec wartości i elastyczne wobec metod; ostrożne w używaniu przymusu i ambitne w zdobywaniu wiedzy. Tak rozumiana adaptacyjność staje się formą instytucjonalnej pokory. Państwo przyznaje, że nie posiada pełnej wiedzy, ale nie rezygnuje z działania. Tworzy procedury, w których hipotezy można testować, błędy wykrywać, sygnały z peryferii przekazywać do centrum, a skutki reform poddawać ponownej ocenie. Pozwala ludziom kwestionować swój model świata, rozumiejąc, że zdolność do korekty jest składnikiem siły, a nie oznaką słabości.

Przewaga adaptacyjna jako projekt nowego humanizmu

W tym miejscu wracamy do pierwszego pytania całego artykułu. Jeśli AI obniża koszt wytwarzania odpowiedzi, przewagą człowieka i instytucji nie może być wyłącznie posiadanie ich większej liczby. Państwo przyszłości będzie dysponować większą ilością danych, modeli prognostycznych i automatycznych analiz niż jakiekolwiek państwo z poprzednich epok.

Nadal jednak będzie ono musiało zdecydować, jakie pytania zadać, komu zaufać, jakie ryzyko podjąć, których praw nie wolno poświęcić oraz jak rozdzielić korzyści i koszty transformacji. Granicą adaptacyjności okazuje się zatem nie technologia, lecz aksjologia. Właśnie dlatego ostatni ruch tej argumentacji musi wykroczyć poza strategię, neuronaukę, organizację i teorię państwa. Jeżeli maszyna przejmie coraz większą część obliczeń, przewidywań, generowania treści i optymalizacji, problem człowieka nie znika.

Staje się bardziej nagi. Nie wystarczy już pytać, co potrafimy zrobić. Trzeba częściej pytać, co powinniśmy zrobić, dla kogo i jakiej przyszłości chcemy przyznać prawo do istnienia. Tu przewaga adaptacyjna przestaje być narzędziem konkurencyjności, a zaczyna przypominać projekt nowego humanizmu. Przewaga adaptacyjna jako nowy humanizm – człowiek po epoce monopolu na inteligencję.

Cała konstrukcja Johna Saneia prowadzi w istocie do pytania znacznie poważniejszego niż to, które zazwyczaj dominuje w debacie o sztucznej inteligencji. Nie brzmi ono: czy maszyna odbierze człowiekowi pracę? Niewystarczające jest również pytanie o kompetencje niezbędne do zachowania przewagi na rynku zatrudnienia. Problem sięga głębiej, ponieważ przez kilka stuleci nowoczesność wiązała wysoką pozycję człowieka z jego zdolnością rozumowania, obliczania, organizowania

informacji i przewidywania. Jeśli coraz większą część tych czynności można delegować systemom AI, zachwianiu ulega nie tylko struktura zawodów.

Zachwianiu ulega pewna antropologia – opowieść o tym, dlaczego człowiek jest potrzebny. Materiał Saneia ujmuje tę zmianę jako przejście od Ery Mięśni, przez Erę Mózgu, ku Erze Obecności. W pierwszej wartość człowieka zależała przede wszystkim od siły fizycznej, w drugiej od intelektu, analizy i zdolności optymalizacji, natomiast w nadchodzącym porządku rosnąć ma znaczenie obecności, mądrości, intuicji i umiejętności zachowania wewnętrznej stabilności wobec radykalnej niepewności. W innym miejscu źródło formułuje tę tezę jeszcze mocniej: przyszłość ma należeć nie do tych, którzy jedynie „wiedzą więcej”, lecz do ludzi zdolnych oduczać się wcześniejszych modeli i wchodzić w nieznanie bez utraty sprawczości. Dosłownie odczytana teza o końcu wartości intelektu byłaby jednak błędna.

Przejście od ekonomii odpowiedzi do ekonomii osądu

Aktualne dane nie malują obrazu świata, w którym ludzka praca poznawcza zostałaby po prostu zastąpiona przez generatywną AI. Czerwcowy przegląd badań empirycznych ILO z 2026 roku wskazuje, że choć wzrosty produktywności są realne, pozostają one nierówne i nie zawsze przekładają się na mierzalne wyniki gospodarcze. Masowe wypieranie pracowników z rynku pracy pozostaje dotychczas ograniczone; znacznie wyraźniejsze są natomiast zmiany w organizacji pracy, autonomii zatrudnionych oraz samej strukturze zadań. Wcześniejsza analiza ILO i NASK wykazała, że choć około jedna czwarta zawodów jest w pewnym stopniu eksponowana na generatywną AI, transformacja konkretnych zadań wydaje się bardziej prawdopodobna niż całkowite zniknięcie całych profesji. Najbardziej odpowiedzialna naukowo interpretacja intuicji Saneia brzmi zatem następująco: nie kończy się wartość inteligencji, lecz kończy się monopol człowieka na część czynności, które dotychczas świadczyły o inteligencji zawodowej.

Kalkulator nie unieważnił matematyki, lecz zmienił znaczenie ręcznego rachunku. Internet nie zlikwidował wiedzy, ale radykalnie przededefiniował wartość pamięciowego posiadania informacji. Generatywna AI może wykonać analogiczny ruch wobec kolejnych klas działań intelektualnych. Gdy napisanie poprawnego tekstu, stworzenie pierwszej wersji kodu, wyszukanie argumentów, streszczenie dokumentów czy wygenerowanie zestawu hipotez staje się tanie, rośnie znaczenie zdolności rozróżnienia: który tekst jest wart napisania, który kod powinien działać, które argumenty są relewantne i której hipotezy nie wolno wdrożyć. Można powiedzieć, że przesuwamy się od ekonomii odpowiedzi ku ekonomii osądu. Odpowiedź może zostać wyprodukowana przez maszynę; osąd dotyczy natomiast relacji pomiędzy tą odpowiedzią a światem wartości. Nie pyta

jedynie o to, czy rozwiązanie działa, ale czy jego działanie służy celowi, który rzeczywiście chcemy realizować.

W tym miejscu pojawia się starożytne rozróżnienie, którego technologiczna nowoczesność nie zdołała unieważnić. Arystoteles odróżniał *techne* – umiejętność wytwarzania – od *phronesis*, mądrości praktycznej pozwalającej rozważać, co należy czynić w konkretnej sytuacji. Maszyna może zwiększyć naszą *techne* w skali nieznanej wcześniej w historii, lecz nie wynika z tego automatycznie wzrost *phronesis*.

Tę różnicę można ująć jeszcze prościej. System może znakomicie odpowiedzieć na pytanie, jak osiągnąć cel, ale techniczna dostępność odpowiedzi nie rozstrzyga kwestii, czy cel ten powinien zostać osiągnięty. Algorytm zoptymalizuje przydział zasobów, lecz to człowiek musi zdecydować, jakie kryterium sprawiedliwości zostanie wpisane do funkcji celu. AI może zwiększyć produktywność przedsiębiorstwa, ale nie rozstrzygnie samodzielnie, jak podzielić korzyści między kapitał, pracowników i konsumentów. Może pomóc państwu przewidywać ryzyko, lecz nie posiada demokratycznego mandatu do określenia dopuszczalnego poziomu ingerencji w prywatność. Właśnie ten punkt odnajdujemy we współczesnych ramach etycznego zarządzania AI. UNESCO stawia godność osoby i prawa człowieka w centrum rekomendacji dotyczących etyki sztucznej inteligencji, podkreślając, że ostateczna odpowiedzialność nie może zostać przeniesiona z człowieka na system. OECD w zaktualizowanych zasadach AI również wskazuje na godność, autonomię, sprawiedliwość, wartości demokratyczne oraz potrzebę zachowania ludzkiej sprawczości i nadzoru. Nie jest przypadkiem, że w epoce potężnej automatyzacji instytucje międzynarodowe powracają do pojęcia godności. Im więcej można technicznie delegować, tym wyraźniej trzeba określić, czego nie można delegować normatywnie.

Nowy humanizm opiera się na odpowiedzialności, a nie na wydajności poznawczej

W tym punkcie może rozpocząć się nowy humanizm. Nie powinien być on jednak postawą defensywną, próbującą ratować wyjątkowość człowieka poprzez desperackie poszukiwanie czynności, których maszyna rzekomo „nigdy” nie wykona. Historia techniki wielokrotnie kompromitowała podobne przepowiednie; granice możliwości maszyn przesuwają się sukcesywnie tam, gdzie wcześniejsze pokolenia widziały wyłączną domenę ludzkich kompetencji. Budowanie godności człowieka na liście zadań chwilowo niewykonalnych dla komputerów jest filozoficznie kruche. Jeśli bowiem godność zależy od przewagi wydajnościowej, utracimy ją w momencie pojawienia się wydajniejszego systemu. Godność nie jest rekordem świata w przetwarzaniu

informacji. Człowiek nie musi być najlepszym kalkulatorem, aby pozostać podmiotem praw. Nie musi pisać szybciej od modelu językowego, by jego cierpienie miało znaczenie moralne, ani wygrywać z algorytmem w grę strategiczną, by jego wolność zasługiwała na ochronę. Humanizm epoki AI wymaga zatem paradoksalnego wyzwolenia człowieka z kultu jego własnej inteligencji. Oświecenie budowało potęgę rozumu przeciw przesądom, arbitralności i dziedziczonej władzy, co stanowiło jedną z najważniejszych emancypacji w historii idei.

Późna nowoczesność zaczęła jednak stopniowo utożsamiać racjonalność z mierzalnością, a inteligencję ze sprawnością przetwarzania coraz większych ilości informacji. W obliczu systemów wykonujących te czynności szybciej pojawia się pokusa uznania, że człowiek został pokonany we własnej koronnej dyscyplinie. Jest to błędny wniosek, ponieważ człowiek nigdy nie był wyłącznie maszyną obliczeniową niższej jakości. Hannah Arendt przypominała, że ludzkie działanie nie sprowadza się do wytwarzania. Człowiek istnieje także jako istota rozpoczynająca coś nowego w przestrzeni wspólnej, ujawniająca siebie wobec innych i ponosząca konsekwencje działań w świecie, którego nigdy nie kontroluje w pełni. Ta zdolność inicjowania – natalność – stanowi szczególnie istotny kontrpunkt wobec świata algorytmicznego. Nie chodzi tu o metafizyczne twierdzenie, że maszyna „nie potrafi stworzyć niczego nowego”, gdyż systemy generatywne oczywiście wytwarzają nowe kombinacje i artefakty.

Problem Arendtowski leży gdzie indziej: nowość ludzkiego działania posiada autora, którego można pociągnąć do odpowiedzialności. To właśnie odpowiedzialność może stać się jednym z fundamentów nowego humanizmu. Hans Jonas, pisząc jeszcze przed obecną rewolucją cyfrową, zauważał, że wraz ze wzrostem technologicznej potęgi człowieka musi rosnąć zakres jego etycznej odpowiedzialności za skutki działania. Dziś ta intuicja staje się niemal dosłowna.

Adaptacyjność wymaga etycznego kierunku i odpowiedzialności

Możemy tworzyć systemy obejmujące miliony ludzi, wpływać na ekosystem informacyjny społeczeństw, automatyzować decyzje i skalować błędy z prędkością infrastruktury cyfrowej. Technologia poszerza promień sprawczości szybciej niż biologicznie wzrasta nasza zdolność odczuwania konsekwencji. W tej perspektywie adaptacyjność pozbawiona odpowiedzialności staje się jedynie zwiększoną skutecznością przystosowania. Przestępca może być adaptacyjny. System autorytarny może sprawnie uczyć się, jak skuteczniej kontrolować populację. Platforma może dostosowywać algorytmy tak, by coraz dłużej zatrzymywać uwagę użytkownika, nawet jeśli obniża to jego dobrostan. Przedsiębiorstwo może błyskawicznie reagować na luki regulacyjne.

Sama zdolność uczenia się nie posiada wartości moralnej. Dlatego Adaptability Advantage wymaga uzupełnienia o pojęcie kierunku adaptacji. Organizacja nie powinna pytać jedynie, czy potrafi szybko zmienić model, lecz co nowy model robi z człowiekiem. Państwo nie powinno badać tylko skuteczności wykorzystania AI, ale czy jej zastosowanie zwiększa sprawczość obywateli, czy ją zawęża. Społeczeństwo nie może mierzyć sukcesu transformacji wyłącznie wartością inwestycji i wzrostem produktywności, jeśli towarzyszy im degradacja jakości pracy, narastanie nierówności lub utrata autonomii.

Aktualne dane ILO nadają tym pytaniom charakter praktyczny, a nie abstrakcyjny. Organizacja wskazuje, że AI najprawdopodobniej przekształci dużą część zawodów zamiast po prostu je zlikwidować, a kluczowe dla jakości tego procesu będą dialog społeczny, sposób reorganizacji pracy oraz podział korzyści z produktywności. Oznacza to, że przyszłość pracy nie jest technologicznym faktem oczekującym na odkrycie. Jest ona w znacznej mierze przedmiotem instytucjonalnego wyboru. Te same możliwości techniczne mogą służyć wzmacnianiu kwalifikacji pracownika albo radykalnemu ograniczeniu jego autonomii; skróceniu monotonicznych zadań albo zwiększeniu intensywności cyfrowego nadzoru.

W tym miejscu szczególnego znaczenia nabiera sformułowana przez Amartyę Sena i rozwinięta przez Marthę Nussbaum perspektywa zdolności ludzkich. Dobrobyt nie powinien być utożsamiany wyłącznie z ilością wytworzonych dóbr czy dochodem, lecz z realną zdolnością ludzi do prowadzenia życia, które mają powód cenić. Zastosowana do AI zasada ta stawia pytanie odmienne od tradycyjnej miary produktywności. Nie wystarczy wiedzieć, ile czasu oszczędził system; trzeba zapytać, co człowiek może zrobić z odzyskanym czasem i czy posiada realną możliwość decydowania o tym, co robi.

Jeżeli automatyzacja zwiększa produktywność, ale cały zysk zostaje przechwycony przez właściciela technologii, adaptacja społeczna będzie przebiegać inaczej niż wtedy, gdy część korzyści zostanie zamieniona w krótszy czas pracy, wyższe wynagrodzenie, lepszą edukację czy większą dostępność usług publicznych. Technologia nie dostarcza odpowiedzi na pytanie o dystrybucję. To instytucje społecznego wyboru decydują, w jaką strukturę praw, umów i własności zostanie ona wpisana.

Tutaj nowy humanizm spotyka się z Dobrem Wspólnym. Jeśli każde przedsiębiorstwo adaptuje się wyłącznie według własnej funkcji celu, suma lokalnych optymalizacji nie musi stworzyć społeczeństwa, w którym chcemy żyć. Adam Smith dostrzegał możliwości koordynacyjne rynku, lecz nawet liberalna tradycja nigdy nie sprowadzała całej architektury wspólnoty do mechanizmu cenowego. Państwo prawa, rodzina, edukacja, kultura, opieka, zaufanie społeczne i odpowiedzialność międzypokoleniowa wytwarzają dobra, których pełnej wartości rynek nie potrafi

wyrazić w bieżącej cenie. Pytanie brzmi zatem nie tylko o to, jak zwiększać sumę możliwości, ale jak decydować, które z nich powinny stać się rzeczywistością.

Możemy technicznie automatyzować coraz więcej ocen, prognoz, rekomendacji i kontaktów między ludźmi. Czy jednak powinniśmy? Możemy zastąpić znaczną część kontaktu administracyjnego systemem cyfrowym, ale czy każda osoba w sytuacji granicznej powinna rozmawiać z maszyną? Możemy zoptymalizować rekrutację, lecz jakie konsekwencje dla pojęcia szansy i indywidualnego wyjątku ma probabilistyczne filtrowanie kandydatów? Możemy generować treści edukacyjne bez końca, ale czy problemem szkoły jest rzeczywiście niedobór treści? Nowy humanizm musi zatem odzyskać pytanie o cel.

Przez długi okres technologicznego optymizmu cel był często ukryty w samej kategorii postępu: skoro coś można zrobić szybciej, taniej i na większą skalę, uznawano to za ulepszenie. AI czyni tę heurystykę niebezpieczną, ponieważ przyspiesza możliwość realizacji również celów złe postawionych. Inteligentniejsza optymalizacja błędnej funkcji celu może przynieść gorsze skutki niż nieefektywny system, który przynajmniej nie posiadał zdolności skalowania własnego błędu.

Przewaga człowieka to odpowiedzialność za sens i kontrola normatywna

UNESCO słusznie wiąże zarządzanie AI nie tylko z bezpieczeństwem technicznym, ale także z godnością, proporcjonalnością, sprawiedliwością, możliwością ludzkiego nadzoru i prawem do zachowania odpowiedzialności człowieka. Etyka nie może być etapem kontroli jakości wykonywanym po zbudowaniu systemu; powinna uczestniczyć już w wyborze tego, co system ma optymalizować. Z tego powodu UNESCO rozwija narzędzia oceny wpływu etycznego, które pozwalają badać zgodność konkretnych zastosowań AI z wartościami i prawami jeszcze przed ich pełnym wdrożeniem. Jest to etyczny odpowiednik ekonomii uczenia się: próbujemy sprawdzić nie tylko, czy technologia zadziała, ale jakie konsekwencje może wywołać, zanim zostaną one wyskalowane.

Właśnie tutaj można ponownie odczytać Saneiową "obecność". Nie jako mistyczny stan świadomości czy magiczną kompetencję niedostępną maszynie, lecz jako dyscyplinę pozostawiania przy rzeczywistości przed jej natychmiastowym przekształceniem w zadanie optymalizacyjne. Obecność lekarza wobec pacjenta oznacza coś więcej niż szybką klasyfikację objawów. Obecność sędziego wobec sprawy to coś więcej niż dopasowanie faktów do statystycznie najczęstszego rozstrzygnięcia. Obecność polityka wobec wspólnoty polega na dostrzeżeniu, że za tabelami

wskaźników kryją się konkretne biografie, których nie można zredukować do średniej. Tak rozumiana obecność nie konkuruje z AI, lecz korzysta z niej. Dobry lekarz może dzięki AI mieć lepszy dostęp do wiedzy, sędzia do materiału porównawczego, urzędnik do danych, a badacz do narzędzi syntezy. Przewaga człowieka nie polega na odmowie korzystania z maszyny, lecz na zachowaniu odpowiedzialności za sposób, w jaki jej wynik zostanie włączony do świata ludzkich znaczeń.

Podobnie jest z intuicją, kolejną kategorią mocno eksponowaną w materiale Saneia. Nie należy przedstawiać jej jako tajemniczego źródła wiedzy stojącego ponad analizą. W psychologii intuicja może wynikać z szybkiego rozpoznawania wzorców wypracowanych dzięki doświadczeniu, lecz w dziedzinach o złym sprzężeniu zwrotnym może równie szybko prowadzić do błędów. Dlatego nowy humanizm nie powinien przeciwstawiać intuicji algorytmowi, lecz budować dialog pomiędzy ludzkim doświadczeniem, dowodem empirycznym a możliwościami obliczeniowymi maszyny. Najbardziej obiecującym modelem nie jest dychotomia "AI albo człowiek", lecz dobrze zaprojektowany układ człowiek-AI-instytucja. System obliczeniowy może dostarczać możliwości, człowiek osąd i interpretację, a instytucja procedury odpowiedzialności, kontroli i korekty.

Każdy z tych elementów ma odmienne słabości. Człowiek bywa podatny na bias, zmęczenie, ograniczoną pamięć i interesowność. Model może halucynować, reprodukować błędy w danych i pozostawać nieprzejrzysty. Instytucja z kolei może kosztować i bronić status quo. Zamiast romantyzować któregośkolwiek z aktorów, dojrzała architektura powinna projektować ich wzajemne mechanizmy kontroli. Wtedy przewaga adaptacyjna nie oznacza, że człowiek musi zmieniać się szybciej niż maszyna – taka konkurencja byłaby absurdalna. Człowiek powinien raczej budować systemy, które pozwalają mu zmieniać się dostatecznie szybko, aby zachować kontrolę normatywną nad konsekwencjami własnych technologii. Adaptacja nie jest pogonią za prędkością. Jest skracaniem czasu pomiędzy odkryciem, że coś przestaje służyć człowiekowi, a zdolnością do jego korekty.

Tu zamyka się koło całej argumentacji. Speed Mismatch oznaczał, że świat zmienia się szybciej niż nasze modele; Competence Trap pokazała, że sukces może utrudniać ich zmianę. Neuronauka przypomniała o organizmie podatnym na stres, a Unlearning wskazał konieczność odebrania dawnym regułom statusu dogmatu. Zespoły Dzisiaj i Jutro dowiodły, że eksploatacja i eksploracja wymagają odmiennych warunków, podczas gdy ekonomia uczenia przesunęła uwagę z wielkości produkcji ku szybkości zdobywania wartościowej informacji. Transition Layer ujawniła potrzebę przekładu przyszłości na język teraźniejszości, a AI Strategy Scanner uporządkował alokację uwagi.

Prawdziwa adaptacyjność wymaga niezmiennych wartości

Liminalność obnażyła tożsamościowy koszt przejścia, a państwo adaptacyjne wyznaczyło granicę konstytucyjnej wartości. Dopiero teraz widać, że wszystkie te elementy sprowadzają się do jednego problemu: jak pozostać podmiotem w świecie, w którym narzędzia coraz lepiej radzą sobie z czynnościami dotąd utożsamianymi z ludzką podmiotowością. Odpowiedź Sanei można zaakceptować jedynie po jej pogłębieniu. Nie wystarczy „być obecnym”, „regulować układ nerwowy” czy „oduczać się przeszłości”. Człowiek przyszłości potrzebuje także pamięci, odpowiedzialności, języka wartości, instytucji i zdolności do sprzeciwu. Adaptacyjność pozbawiona pamięci staje się płynnością; ta pozbawiona etyki – oportunistą, a bez granic – uległością wobec każdego nowego bodźca.

Najwyższa forma adaptacyjności polega więc paradoksalnie na wiedzy, kiedy nie należy się adaptować. Człowiek powinien przystosować metodę, gdy zmienia się rzeczywistość, ale nie powinien dostosowywać godności do aktualnej ceny rynkowej. Państwo powinno modyfikować instrumenty, lecz nie poddawać sezonowej optymalizacji praw podstawowych. Nauka powinna zmieniać teorie w obliczu dowodów, a nie prawdę pod naciskiem politycznym. Organizacja może przebudować model biznesowy, ale nie powinna uznawać ludzi za zbędne koszty tylko dlatego, że krótkoterminowy algorytm rentowności tak podpowiada.

To właśnie może stać się konstytucją nowego humanizmu: elastyczne metody, krytyczne przekonania i ruchome modele – przy jednoczesnym twardym zobowiązaniu wobec godności osoby oraz Dobra Wspólnego. Nie jest to humanizm człowieka triumfującego nad maszyną, lecz człowieka wystarczająco dojrzałego, aby nie potrzebował triumfu. Nie chodzi o kogoś, kto zachowuje wartość dzięki umiejętnościom, których AI jeszcze nie posiada, lecz o człowieka będącego nosicielem wartości, dla których technologie w ogóle powinny działać.

Być może najważniejsza rewolucja ery sztucznej inteligencji okaże się mniej technologiczna, niż nam się wydaje. Przez wieki pytaliśmy, ile świata człowiek potrafi podporządkować własnemu rozumowi. W epoce inteligentnych maszyn będziemy musieli pytać, czy potrafimy podporządkować rosnącą moc temu, co uznajemy za dobre. Maszyna pomoże nam dostrzec więcej możliwości, ale nie zwolni nas z obowiązku wyboru między nimi.

Sanei ma rację: przyszłość nie należy po prostu do najsilniejszych ani do tych, którzy zgromadzili najwięcej informacji. Nie wynika to jednak z faktu, że zwyciężą osoby najbardziej elastyczne – w świecie, w którym algorytmy również stają się elastyczne, przewagę człowieka trzeba zdefiniować

głębiej. Będzie nią zdolność uczenia się bez utraty pamięci, zmiany bez utraty odpowiedzialności, eksperymentowania bez utraty godności i korzystania z inteligencji większej od własnej bez rezygnacji z obowiązku określenia celu, w jakim jej używamy.

„Era Obecności” może zatem zyskać znaczenie bardziej wymagające niż to, które nadaje jej literatura futurystyczna. Obecność oznacza bycie na tyle przytomnym wobec własnej potęgi, aby nie pomylić tego, co możliwe, z tym, co pożądane; tego, co wydajne, z tym, co sprawiedliwe; tego, co nowe, z tym, co lepsze. Człowiek przyszłości nie musi być ostatnim bastionem biologicznej inteligencji. Powinien być strażnikiem pytania, którego żadna cywilizacja technologiczna nie może bezpiecznie utracić: jaką przyszłość warto wspólnie stworzyć?

Ludzkość i AGI – nowy podział pracy między inteligencją, odpowiedzialnością i sensem

Na wstępie należy postawić granicę między wiedzą a futurologią. Nauka nie dysponuje dziś powszechnie uzgodnionym testem, którego przekroczenie pozwalałoby bezspornie stwierdzić: oto powstała AGI. Jedna z wpływowych propozycji badaczy Google DeepMind sugeruje wręcz porzucenie myślenia o AGI jako pojedynczym progu na rzecz opisu systemów poprzez ogólność ich kompetencji, poziom wykonania oraz autonomię. Autorzy postulują: „focus on capabilities, not processes” – patrzymy przede wszystkim na zdolności, a nie na sposób ich powstania – proponując poziomy zaawansowania zamiast jednego metafizycznego momentu „narodzin AGI”. Dlatego wszelkie rozważania o podziale funkcji między ludzkością a przyszłą AGI muszą łączyć dwa porządki: aktualne wyniki badań nad AI oraz hipotezy dotyczące systemu znacznie bardziej ogólnego.

To zastrzeżenie paradoksalnie wzmacnia intuicję obecną w ujęciu Johna Saneia. Materiał źródłowy opisuje przejście od wartości opartej na sile fizycznej, przez wartość opartą na intelekcie, ku hipotetycznej epoce, w której – skoro „intelekt staje się towarem powszechnym” – kluczowe znaczenie nabierają mądrość, obecność, intuicja i zdolność działania w niepewności. Jako opis empirycznie ustalonej ewolucji byłoby to zbyt daleko idące, jednak jako hipoteza o zmianie struktury rzadkości jest niezwykle interesujące.

Praca jako proces ciągłego rozpinania i zszywania zadań

Jeśli zdolność generowania tekstu, kodu, analizy, projektu czy prognozy stanie się powszechnie tania, wzrośnie względna wartość tych funkcji, których nie da się sprowadzić do samego wygenerowania odpowiedzi. Najważniejszym zasobem świata po AGI może zatem nie być

inteligencja, lecz prawo do wyznaczania celu tej inteligencji. Ekonomia pracy sugeruje, że punktem wyjścia nie powinien być podział na zawody „ludzkie” i „maszynowe”. Współczesny task-based model, rozwijany m.in. przez Daroną Acemoglu, Fredricą Konga i Pascuala Restrepo, traktuje produkcję jako wiązkę zadań, które można przydzielić albo ludziom o określonych kompetencjach, albo kapitałowi. Technologia może zwiększyć ludzką produktywność w dotychczasowym zadaniu, zautomatyzować je, podnieść wydajność kapitału lub stworzyć zupełnie nowe obszary aktywności. Najnowsze modele idą jednak jeszcze dalej.

Althoff i Reichardt wykazują w pracy z czerwca 2026 roku, że AI potrafi jednocześnie automatyzować, augmentować i upraszczać zadania, zmieniając wymagany profil kompetencji, a tym samym przewagę komparatywną pracowników. Joshua Gans wskazuje natomiast, że AI może przesuwać granice zawodów nawet bez pełnej automatyzacji: samo obniżenie kosztu przekazywania kontekstu między uczestnikami procesu może doprowadzić do ponownego grupowania zadań w całkowicie nowe profesje. Prowadzi to do pierwszego istotnego wniosku cywilizacyjnego: przyszłość pracy prawdopodobnie nie będzie historią zastępowania zawodów, lecz historią ciągłego rozpinania i ponownego zszywania zadań.

Lekarz nie musi „zostać zastąpiony przez AGI”, aby jego zawód stał się zupełnie inny. Wystarczy, że system przejmie wyszukiwanie literatury, interpretację obrazowania, tworzenie dokumentacji, różnicowanie hipotez i monitorowanie parametrów. Profesor może nadal istnieć, choć maszyna przejmie generowanie materiałów, testowanie, tłumaczenia, indywidualne ćwiczenia i znaczną część przekazywania faktów. Prawnik pozostanie prawnikiem, nawet jeśli ogromna część wyszukiwania orzecznictwa, porównywania umów i przygotowywania pierwszych wersji pism zostanie wykonana syntetycznie. Zawód zachowa nazwę, podczas gdy jego wewnętrzna anatomia może zostać wymieniona niemal całkowicie. Dotychczasowa empiria GenAI dostarcza już pierwszych sygnałów o tym mechanizmie.

Badanie Brynjolfssona, Li i Raymond obejmujące 5172 pracowników obsługi klienta wykazało średnio 15-procentowy wzrost produktywności po wprowadzeniu asystenta AI. Korzyści były najwyższe w grupie osób mniej doświadczonych i słabszych – tutaj wzrost sięgał około 30 procent. Doświadczenie, zakodowane wcześniej w praktykach najlepszych pracowników, zostało w pewnym sensie rozproszane przez system na większą część organizacji. ILO w przeglądzie z czerwca 2026 roku odnotowuje coraz więcej takich wzrostów produktywności, podkreśla jednak ich nierównomierność, brak dowodów na masową redukcję zatrudnienia oraz ryzyka dotyczące młodszych pracowników, autonomii i jakości organizacji pracy. Z tego rodzi się jeden z najbardziej niedocenianych problemów ery AGI: paradoks drabiny kompetencyjnej.

Jeśli maszyna przejmie proste zadania wykonywane dotychczas przez juniorów, przedsiębiorstwo uzyska natychmiastową korzyść produktywnościową. Jednak to właśnie wykonywanie tych prostych zadań było sposobem, w jaki junior po dziesięciu latach stawał się seniorem. Automatyzując dół drabiny, można przypadkowo zniszczyć mechanizm reprodukcji ekspertów potrzebnych na jej szczycie. Hipoteza do weryfikacji H1: im silniej AGI zautomatyzuje zadania początkujących, tym ważniejsze stanie się świadome projektowanie sztucznych ścieżek terminowania i praktyki. Nie wiemy jeszcze, czy efekt ten będzie dominujący.

AI może równie dobrze działać jak prywatny nauczyciel i dramatycznie przyspieszać rozwój kompetencji. Wspomniane badanie obsługi klienta pokazało, że pracownicy korzystający z AI szybciej przesuwali się po krzywej uczenia, a część korzyści utrzymywała się nawet podczas awarii systemu. Problem naukowy nie brzmi zatem „czy AI zabije uczenie?”, lecz: w jakich architekturach pracy AI zastępuje doświadczenie, a w jakich je przyspiesza?

Od sprawczości operacyjnej do konstytucyjnej i zbiorowej

Jeszcze ważniejszego ostrzeżenia dostarcza metaanaliza Michelle Vaccaro, Abdullaha Almaatouqa i Thomasa Malone'a opublikowana w *Nature Human Behaviour*. Autorzy przeanalizowali 106 eksperymentów i 370 efektów dotyczących współpracy ludzi z AI. Wynik jest wyjątkowo niewygodny dla popularnej ideologii "centaura": kombinacja człowiek-AI była przeciętnie lepsza niż człowiek działający samotnie, ale gorsza niż ten z dwóch aktorów, który był silniejszy i pracował samodzielnie. W zadaniach decyzyjnych łączenie człowieka z AI wiązało się średnio ze stratą, podczas gdy znacznie bardziej obiecujące wyniki pojawiały się w obszarach twórczych. Autorzy kładą szczególny nacisk na konieczność właściwego zaprojektowania podziału pracy pomiędzy partnerami. To odkrycie powinno zostać zapisane dużymi literami w każdej teorii świata AGI.

Człowiek plus AI nie staje się automatycznie inteligentniejszy od samej maszyny. Obecność człowieka w pętli może wręcz pogorszyć decyzję, jeśli nie wie on, kiedy systemowi zaufać, kiedy mu się sprzeciwić i których części zadania w ogóle nie powinien dotykać. Podobnie maszyna może osłabiać osąd eksperta, gdy jej rekomendacja tworzy kotwicę poznawczą lub fałszywe poczucie pewności. Badanie 1411 specjalistów z obszaru HR i bankowości pokazało na przykład, że osoby nadzorujące AI były równie skłonne podążać za radą systemu sprawiedliwego, co systemu generującego dyskryminujące rekomendacje.

"Human in the loop" nie jest zatem magiczną formułą bezpieczeństwa. Prawdziwy problem dotyczyć będzie architektury delegacji. Nauka proponuje tu dojrzsze rozrózenie między operative agency a evaluative agency. W pracy z 2026 roku poświęconej meaningful human oversight autorzy sugerują, aby system AI posiadał szeroką sprawczość operacyjną – wykonywał pracę, budował rozwiązania i poszukiwał ścieżek działania – podczas gdy człowiek zachowuje sprawczość ewaluacyjną: zdolność zrozumienia rezultatu na wymaganym poziomie, zakwestionowania go, zatrzymania procesu i wzięcia odpowiedzialności za konsekwencje. Podobny kierunek reprezentuje literatura dotycząca Meaningful Human Control; nie utożsamia ona ludzkiej kontroli z koniecznością ręcznego zatwierdzania każdej mikroczynności, lecz z zachowaniem odpowiedzialności, przejrzystości i realnej możliwości interwencji w system.

Tak może wyglądać podstawowy podział funkcji pomiędzy AGI a ludzkością. AGI przejmuje coraz większą część sprawczości wykonawczej: wyszukiwanie, modelowanie, symulowanie, diagnozowanie, generowanie, harmonogramowanie, monitorowanie, optymalizację, koordynację i być może samodzielne realizowanie długich sekwencji działań. Człowiek oraz zbudowane przez ludzi instytucje przesuwają się w stronę sprawczości konstytucyjnej: wyznaczania granic, definiowania celów, rozstrzygania sporów o wartości, określenia praw, ustalania odpowiedzialności i decydowania, czego w ogóle nie należy optymalizować.

Nie wynika to z dowodu, że AGI nigdy nie opanuje rozumowania moralnego. Być może będzie robiła to znakomicie. Wynika to jednak z czegoś innego: legitymizacja nie jest testem inteligencji. System może podać trafniejszą diagnozę polityki publicznej niż parlament i nadal nie posiadać demokratycznego mandatu. Może trafniej przewidzieć wyrok niż sędzia, a mimo to nie być źródłem konstytucyjnej władzy sądowniczej. Może przedstawić bardziej konsekwentną argumentację moralną niż człowiek, ale sama przewaga argumentacyjna nie daje prawa do decydowania o cudzym życiu. UNESCO ujmuje tę granicę kategorycznie: AI "can never replace ultimate human responsibility and accountability". Nie jest to twierdzenie z zakresu neuronauki czy informatyki o metafizycznej naturze maszyny, lecz decyzja normatywna dotycząca tego, gdzie chcemy lokować odpowiedzialność.

Możemy delegować wykonanie, ale nie powinniśmy automatycznie delegować legitymizacji. W tym punkcie perspektywa jednostkowa okazuje się niewystarczająca. Ludzkość nie jest pojedynczym mózgiem konkurującym z większym mózgiem AGI. Jest zbiorowością miliardów ludzi, instytucji, kultur, państw, przedsiębiorstw, laboratoriów, systemów prawnych i wspólnot pamięci. Jej najważniejsza inteligencja już dziś w dużej mierze jest inteligencją zbiorową. Nauka funkcjonuje dlatego, że wiedza jest rozproszona między badaczy, czasopisma i pokolenia. Rynek koordynuje

informacje posiadane przez miliony aktorów. Demokracja – w idealnym modelu – pozwala konfrontować konkurencyjne rozpoznania Dobra Wspólnego.

Język sam w sobie jest ogromnym magazynem poznawczym, którego żaden pojedynczy człowiek nie stworzył. Badacze kolektywnej inteligencji traktują LLM-y nie tylko jako nowe indywidualne narzędzia, lecz jako technologię mogącą przebudować sposób agregowania, transmitowania i porządkowania społecznej wiedzy. Nature Human Behaviour zwraca uwagę, że systemy te mogą jednocześnie wspierać i zaburzać zbiorową zdolność rozwiązywania złożonych problemów. Inna literatura wyróżnia trzy funkcje, w których AI może wzmacniać inteligencję zbiorową: pamięć zbiorową, uwagę zbiorową i rozumowanie zbiorowe.

Jeżeli powstanie AGI, może ona okazać się czymś znacznie większym niż "pracownik bez ciała". Może stać się warstwą poznawczą cywilizacji: interfejsem pomiędzy ogromnymi zasobami wiedzy, systemem tłumaczenia między dyscyplinami, pamięcią dostępną na żądanie, generatorem symulacji, koordynatorem wielkich przedsięwzięć i mechanizmem odkrywania nieoczywistych powiązań. Hipoteza do weryfikacji H2: przełomem AGI nie będzie przede wszystkim automatyzacja pracy jednostek, lecz reorganizacja inteligencji zbiorowej ludzkości. Byłaby to zmiana porównywalna nie tyle z pojawieniem się lepszego pracownika, co z wynalezieniem pisma, druku, telekomunikacji lub Internetu – technologii, które zmieniały nie pojedynczy akt poznania, lecz sposób, w jaki wiedza przepływa przez cywilizację. Jednocześnie istnieje ciemna strona tej możliwości.

Paradoks AI: wzrost wydajności kosztem różnorodności nauki

AI może zwiększyć indywidualne możliwości każdego naukowca, jednocześnie ograniczając różnorodność nauki jako całości. Analiza 41,3 miliona publikacji opublikowana w Nature ujawniła właśnie taki paradoks: badacze korzystający z narzędzi AI publikowali więcej i częściej byli cytowani, jednak na poziomie zbiorowym zakres podejmowanych tematów ulegał zmniejszeniu, a wzajemne zaangażowanie naukowców słabło. Autorzy odnotowali 4,63-procentowe zawężenie przestrzeni tematów oraz spadek wzajemnego zaangażowania o 22 procent.

Granica autonomii jako ostatni bastion ludzkiej sprawczości

Wniosek ten ma ogromne znaczenie dla hipotezy AGI. Inteligencja może bowiem maksymalizować sukces jednostek, jednocześnie ograniczając eksploracyjność całego systemu. Jeśli wszyscy zaczniemy korzystać z podobnych modeli, danych i zoptymalizowanych rekomendacji, społeczeństwo może stać się niezwykle produktywnie w badaniu tego, co już uznaje za obiecujące, lecz utraci zdolność do wchodzenia na ścieżki dziwne, niepopularne czy pozornie nieracjonalne. Prowadzi to do hipotezy, że największym poznawczym ryzykiem AGI może nie być brak wiedzy, lecz epistemiczna monokultura. W takim świecie jedną z najważniejszych funkcji człowieka byłoby nie dostarczanie „lepszego rozwiązania” niż maszyna, ale wytwarzanie heterogeniczności: stawianie innych pytań, wyznaczanie odmiennych celów, wnoszenie lokalnych doświadczeń, sprzeciwu, kulturowych anomalii oraz problemów, których centralny system nie uznał za wartę optymalizacji. Margines, herezja, prowincja, subkultura i eksperyment zyskują wówczas nową wartość jako źródła poznawczej różnorodności, co nadaje świeże znaczenie Saneiowej figurze outsidera.

W materiale źródłowym Tomorrow Team ma chronić eksplorację, przyjmować pozycję nowicjusza i być mierzony tempem uczenia oraz liczbą strategicznych możliwości, a nie wyłącznie bieżącym ROI. W świecie AGI odpowiednikiem tej zasady na poziomie cywilizacyjnym może być celowa ochrona ludzkiego prawa do nieracjonalnej na pierwszy rzut oka eksploracji. Nie wynika to z założenia, że człowiek jest z definicji bardziej kreatywny od przyszłej maszyny – tego nie wiemy. Chodzi o to, że system poznawczy traci odporność, gdy wszystkie jego elementy zaczynają korzystać z tego samego modelu rzeczywistości. Pojawia się wtedy głębszy podział funkcji: AGI może stać się mistrzem przeszukiwania przestrzeni możliwości, podczas gdy zbiorowość ludzka zachowa funkcję ustanawiania przestrzeni znaczenia.

System może znaleźć tysiąc metod zwiększenia wydajności szkoły, lecz to społeczeństwo musi zdecydować, czym jest wykształcony człowiek. AGI może projektować dziesiątki scenariuszy systemu emerytalnego, ale to obywatele muszą rozstrzygnąć, jaki poziom solidarności międzypokoleniowej uważają za sprawiedliwy. Maszyna zoptymalizuje ochronę zdrowia pod kątem liczby przeżytych lat, kosztów lub jakości życia; to jednak wspólnota polityczna musi określić, czy i w jaki sposób wolno zestawiać te wartości. Nie pytamy już tylko: kto wykona zadanie najlepiej? Pytamy: kto ma prawo zdecydować, że zadanie ma zostać wykonane?

Rozróżnienie to zyskuje na wadze w momencie, gdy AGI otrzyma autonomię. DeepMind słusznie oddziela poziom zdolności od poziomu autonomii: bardzo inteligentny system może służyć jako narzędzie uruchamiane na żądanie, podczas gdy system słabszy może otrzymać relatywnie szeroki

zakres samodzielnego działania. Podział funkcji przyszłości będzie zatem oparty na dwóch osiach: tym, co system potrafi, oraz tym, na co pozwalamy mu bezpośrednio działać. Prowadzi to do kolejnej hipotezy, że najważniejszą granicą ery AGI nie będzie granica kompetencji, lecz granica autonomii. Jeśli AGI okaże się lepsza od człowieka w analizie, planowaniu i przewidywaniu, presja ekonomiczna będzie przesuwiała ją z funkcji doradczej ku wykonawczej.

Pojawi się pokusa: skoro system poprawnie proponuje decyzję, pozwólmy mu ją realizować; skoro robi to bezbłędnie tysiąc razy, zmniejszmy częstotliwość nadzoru; a gdy nadzór stanie się kosztowną formalnością – zrezygnujmy z niego. Ostatecznym problemem nie będzie zatem pytanie „czy AI jest inteligentna?”, lecz czy człowiek zachowuje praktyczną możliwość przywrócenia kontroli w momencie, gdy trajektoria systemu zacznie odbiegać od wartości, które miała realizować. Właśnie temu służy rozwijająca się literatura Meaningful Human Control. Nauka podpowiada jednocześnie, że ten podział nie powinien być statyczny. Badania nad dynamicznym przydziałem zadań w zespołach człowiek-AI postulują podejście zależne od kontekstu, ryzyka, aktualnych kompetencji obu stron oraz konieczności zachowania ludzkiej sprawczości.

Nowy podział pracy w trójkącie człowiek-AGI-instytucja

Jest to logiczna konsekwencja "poszarpanej granicy technologicznej". Eksperyment z 758 konsultantami pokazał, że AI potrafi radykalnie zwiększyć jakość i szybkość wykonania zadań mieszczących się w jej aktualnych kompetencjach, jednak ta granica przebiega nierówno nawet między pozornie podobnymi czynnościami. W świecie AGI zasada ta mogłaby zostać przeniesiona z poziomu firmy na poziom całej cywilizacji: zadania powinny migrować między ludźmi a systemami w zależności od aktualnej przewagi, lecz funkcje konstytucyjne winny migrować znacznie wolniej. Możemy błyskawicznie zmienić podmiot analizujący dane, ale znacznie ostrożniej powinniśmy zmieniać podmiot posiadający prawo do ostatecznej decyzji o wolności, życiu, własności, wojnie czy rozkładzie praw. Wreszcie AGI może przededefiniować sam sens ekonomicznej przewagi komparatywnej. Klasyczna intuicja Ricardo przypomina, że podział pracy może istnieć nawet wtedy, gdy jeden aktor jest absolutnie lepszy we wszystkim, o ile względne koszty alternatywne pozostają różne.

W relacji człowiek-AGI mechanizm ten może jednak napotkać dodatkowe ograniczenia. Zadania pozostaną przy człowieku nie tylko dlatego, że jest w nich względnie lepszy, ale ponieważ wymaga tego prawo, preferencja społeczna, odpowiedzialność, fizyczna relacja, zaufanie albo samo znaczenie czynności. Rodzic może otrzymać od AGI znakomitą analizę sposobu rozmawiania z dzieckiem, lecz nikt rozsądny nie wyciąga z tego automatycznego wniosku, że funkcja rodzica

powinna zostać zoptymalizowana poprzez zastąpienie go systemem. Ekonomia podziału zadań zderzy się więc z ekonomią znaczenia. Hipoteza do weryfikacji brzmi: im tańsza stanie się inteligencja wykonawcza, tym większa część wartości ekonomicznej przesunie się ku odpowiedzialności, zaufaniu, relacji, legitymizacji i własności celu. Może to wytworzyć paradoks.

AGI będzie potrafiła wykonać większość pracy koniecznej do powstania rezultatu, a człowiek pozostanie centralny dlatego, że odpowiada za to, który rezultat ma otrzymać prawo wejścia do świata. Rola człowieka przesunie się z producenta ku kuratorowi możliwości, z wykonawcy ku konstytucjonalistce systemu, z posiadacza wiedzy ku właścicielowi zobowiązania. Nie jest jednak przesądzone, czy taki porządek powstanie.

Możliwa jest również odwrotna trajektoria: koncentracja infrastruktury AGI w rękach niewielu organizacji może pogłębić nierówność władzy poznawczej i gospodarczej. Możliwe jest wypłukanie ludzkich kompetencji na skutek nadmiernego uzależnienia od systemów, ujednolicenie kultury czy pojawienie się nowych zawodów, których obecnie nie umiemy nazwać. Nauka nie zna jeszcze bilansu. Dlatego najbardziej odpowiedzialną teorią przyszłości nie jest ta, według której "człowiek zachowa X, bo maszyna nigdy nie będzie potrafiła X". Tego typu twierdzenia są epistemicznie słabe. Lepsza jest inna konstrukcja.

Nie pytajmy, co AGI nigdy nie będzie potrafiła zrobić. Pytajmy, czego nawet bardzo zdolnej AGI nie chcemy oddać bez politycznej, moralnej i instytucjonalnej decyzji.

Adaptacyjność Saneia spotyka się tutaj z problemem Dobra Wspólnego. Jeśli AGI rzeczywiście stanie się egzemplifikacją kolejnej rewolucji cywilizacyjnej, wyzwaniem ludzkości nie będzie wygranie z nią konkursu inteligencji. Tak jak człowiek nie pokonał silnika parowego w podnoszeniu ciężarów i nie próbował odzyskać sensu życia przez trenowanie silniejszych mięśni, tak przyszła cywilizacja nie powinna bronić godności człowieka poprzez zmuszanie go do produkowania analiz szybciej niż AGI.

Jej zadaniem będzie zbudowanie nowego społecznego podziału funkcji. AGI może stać się pamięcią, symulatorem, konstruktorem wariantów, koordynatorem, tłumaczem, narzędziem odkrywania, operatorem i wykonawcą. Ludzie mogą coraz bardziej stawać się twórcami celów, nosicielami praw, uczestnikami relacji, źródłami legitymizacji, autorami zobowiązań i wspólnotami decydującymi, które spośród milionów możliwych przyszłości są warte realizacji. Instytucje będą natomiast potrzebne po to, aby utrzymywać pomiędzy tymi dwiema sferami kontrolę, odpowiedzialność, pluralizm i możliwość korekty.

Hipoteza do weryfikacji zakłada, że docelowym ustrojem epoki AGI nie będzie "człowiek kontra maszyna", lecz trójkąt człowiek-AGI-instytucja. Człowiek bez AGI może być poznawczo zbyt słaby

wobec złożoności świata. AGI bez instytucji może być potężna, ale pozbawiona społecznie uzgodnionej legitymizacji. Instytucja bez zdolności adaptacyjnej może z kolei próbować rządzić nowym światem przy użyciu map sporządzonych dla poprzedniego.

Właśnie ten trójkąt może stać się elementarną komórką przyszłej cywilizacji. A więc nie koniec człowieka, lecz koniec człowieka rozumianego wyłącznie jako najwydajniejszy znany nośnik inteligencji. Jeżeli AGI nadejdzie, ludzkość może zostać zmuszona do niezwykłego aktu unlearning: porzucenia przekonania, że jej godność wynika z dominacji poznawczej. W zamian będzie musiała przypomnieć sobie znacznie starszą ideę, którą cywilizacja techniczna częściowo przesłoniła – że rozum jest środkiem, a nie ostatecznym celem; że wiedza nie określa sama tego, co dobre; że możliwość nie jest jeszcze powinnością; i że największa odpowiedzialność przypada temu, kto może uczynić najwięcej. Być może zatem najbardziej przewrotna konsekwencja AGI będzie miała charakter humanistyczny. Im inteligentniejsze staną się maszyny, tym mniej człowiek będzie mógł usprawiedliwiać się samą inteligencją. Pozostanie pytanie o mądrość. Pozostanie pytanie o odpowiedzialność. Pozostanie pytanie o Dobro Wspólne.

Okaże się, że prawdziwy podział pracy nowej cywilizacji nie przebiega pomiędzy człowiekiem, który myśli, a maszyną, która wykonuje. Będzie przebiegał między coraz potężniejszym systemem odkrywania tego, co możliwe, a ludzkością, która musi nauczyć się znacznie trudniejszej sztuki: wspólnie rozstrzygać, co z tego, co możliwe, powinno zostać urzeczywistnione.

Być może największym paradoksem ery sztucznej inteligencji będzie to, że im potężniejsze staną się maszyny, tym mniej będziemy mogli usprawiedliwiać swoją wartość samą inteligencją. Pozostaniemy z pytaniem o mądrość i odpowiedzialność, których nie da się zoptymalizować algorytmicznie. Ostatecznie prawdziwy podział pracy nowej cywilizacji nie przebiegnie między tym, kto myśli, a tym, kto wykonuje, lecz między systemem odkrywającym to, co możliwe, a ludźmi decydującymi, co z tych możliwości jest w ogóle warte urzeczywistnienia.

Inspiracje

[1]



"The Adaptability Advantage" John Sanei

2026 | John Wiley & Sons | ISBN: 9781394421886

John Sanei to uznany międzynarodowy futurolog, strateg biznesowy i autor bestsellerów, znany ze swojej pracy nad konwergencją neuronauki, technologii i strategii biznesowej. Specjalizuje się w pomaganiu liderom i organizacjom w radzeniu sobie z niepewnością, destabilizacją i przyszłością pracy w erze napędzanej przez sztuczną inteligencję. Sanei jest globalnym ekspertem na Singularity University i członkiem wydziału Duke Corporate Education. Ponadto jest partnerem stowarzyszonym w Copenhagen Institute for Futures Studies. Uznawany za czołowego globalnego futurologa, doradza firmom z listy Fortune 100, rządowi i organizacjom globalnym w zakresie budowania elastycznych, zwinnych kultur. Sanei jest autorem licznych książek poświęconych przywództwu, sposobowi myślenia i strategicznej dalekowzroczności, których celem jest pomoc jednostkom i zespołom w przejściu od zachowań reaktywnych do proaktywnych, celowych działań w szybko zmieniającym się świecie.

John Sanei is an acclaimed international futurist, business strategist, and bestselling author known for his work on the convergence of neuroscience, technology, and business strategy. He specializes in helping leaders and organizations navigate uncertainty, disruption, and the future of work in an AI-powered era. Sanei serves as a global expert at Singularity University and is a faculty member at Duke Corporate Education. Additionally, he is an associate partner at the Copenhagen Institute for Futures Studies. Recognized as a top global futurist, he advises Fortune 100 companies, governments, and global organizations on building adaptable, agile cultures. Sanei is a prolific author who has published numerous books focused on leadership, mindset, and strategic foresight, aiming to help individuals and teams transition from reactive behaviors to proactive, purposeful action in a rapidly changing world.

Singularity University, Duke Corporate Education, Copenhagen Institute for Futures Studies

Literatura uzupełniająca

Książki

Autorzy przywoływani w artykule:



[1] "Expansive"

John Sanei, Erik Kruger

2025 | Jonathan Ball Publishers | ISBN: 9781067228958



[2] "Who Do We Become?"

John Sanei

2022 | Jonathan Ball Publishers | ISBN: 9781776192175



[3] "Magnitiize"

John Sanei

2019 | ISBN: 9781928230533



[4] "What's Your Moonshot?"

John Sanei

2018 | ISBN: 9781928230472



[5] "Foresight"

JOHN. SANEI

2019 | ISBN: 9781928230748



[6] "FutureNEXT"

Iraj Abedian, John Sanei

2020 | ISBN: 9798692623935

Literatura pokrewna (Google Scholar):



[7] "The Ritual Process: Structure and Anti-Structure"

Victor Turner

Transaction Publishers | ISBN: 9780202368634



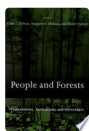
[8] "The Ritual Process"

Victor Turner

1995 | Transaction Publishers | ISBN: 9780202011905

[9] "Institutions, Incentives, and Irrigation in Nepal"

Elinor Ostrom



[10] "People and Forests: Communities, Institutions, and Governance"

Elinor Ostrom

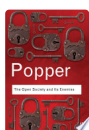
MIT Press | ISBN: 9780262571371



[11] "The Commons in the New Millennium: Challenges and Adaptations"

Elinor Ostrom

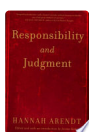
ISBN: 9780262271868



[12] "The Open Society and Its Enemies"

Karl Popper

Routledge | ISBN: 9781136700323



[13] "Responsibility and Judgment"

Hannah Arendt

Schocken | ISBN: 9780307544056



[14] "The Human Condition"

Hannah Arendt

2013 | University of Chicago Press | ISBN: 9780226924571



[15] "Identity and Violence"

Amartya Sen

Penguin Books India | ISBN: 9780141027807

[16] "Skills, Tasks and Technologies: Implications for Employment and Earnings"

Daron Acemoglu

[17] "Distorted innovation: Does the market get the direction of technology right?"

Daron Acemoglu



[18] "The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies"

Erik Brynjolfsson

W. W. Norton & Company | ISBN: 9780393241259



[19] "On the Principles of Political Economy and Taxation"

David Ricardo

Library of Alexandria | ISBN: 9781465520302



[20] "Guiding Projects through Transformation"

Barrett Williams, ChatGPT

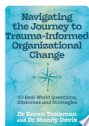
2025 | Barrett Williams



[21] "Organizational Behavior and Human Resource Management for Complex Work Environments"

Belias, Dimitrios, Rossidis, Ioannis, Papademetriou, Christos

2024 | IGI Global | ISBN: 9798369334676



[22] "Navigating the Journey to Trauma-Informed Organizational Change"

Karen Treisman, Mandy Davis

2026 | Jessica Kingsley Publishers | ISBN: 9781805019152



[23] "Human Resources Management"

St. Clements University Academic Staff

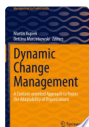
Prof. Dr. Bilal Semih Bozdemir



[24] "Organizational Change"

Maria Vakola, Paraskevas Petrou

2018 | Routledge | ISBN: 9781315386096



[25] "Dynamic Change Management"

Martin Kupiek, Bettina Marcinkowski

2024 | Springer Nature | ISBN: 9783031707063

Artykuły naukowe

Autorzy przywoływani:

[1] "Los ritos de paso"

Arnold van Gennep

2008

[Google Scholar](#)

[2] "Governing the Commons: The Evolution of Institutions for Collective Action"

Barry C. Field, Элинор Остром

1992 | Land Economics | DOI: 10.2307/3146384

[Google Scholar](#)

[3] "Tribute to the Life and Work of Friedrich Hayek"

Karl Popper

1997 | Hayek: Economist and Social Philosopher | DOI: 10.1007/978-1-349-25991-5_14

[Google Scholar](#)

[4] "The Human Condition"

Hannah Arendt, Margaret Canovan, Danielle Allen

1998 | DOI: 10.7208/chicago/9780226586748.001.0001

[Google Scholar](#)

[5] "Philosophical Essays: From Ancient Creed to Technological Man."

Daniel S. Robinson, Hans Jonas

1974 | Philosophy and Phenomenological Research | DOI: 10.2307/2106643

[Google Scholar](#)

[6] "Animal RightsCurrent Debates and New Directions"

Cass R. Sunstein, Martha C. Nussbaum

2005 | Oxford University Press eBooks | DOI: 10.1093/acprof:oso/9780195305104.001.0001

[Google Scholar](#)

[7] "Aristotle, Politics, and Human Capabilities: A Response to Antony, Arneson, Charlesworth, and Mulgan"

Martha C. Nussbaum

2000 | Ethics | DOI: 10.1086/233421

[Google Scholar](#)

Artykuły pokrewne (Google Scholar):

[8] "The Comparative Analysis of the Scoring System Used in BREEAM International New Construction 2016 and the Recent Trends in Housing Sustainability-Related Literature"

Mohsen Sanei

2022 | Prostor | DOI: 10.31522/p.30.1(63).10

[Google Scholar](#)

[9] "Understanding Online Meaning-Making: Normativity, Polycentricity, and Memetic Communication Through a Chronotopic-Scalar Lens"

Taraneh Sanei

2026 | Journal of Sociolinguistics | DOI: 10.1111/josl.70056

[Google Scholar](#)

[10] "General Overview to the Research Programs in Part I"

Arezoo Sanei

2020 | Research and Management Practices for Conservation of the Persian Leopard in Iran | DOI: 10.1007/978-3-030-28003-1_1

[Google Scholar](#)

[11] "Reclaiming Establishment: Identity and the 'Religious Equality Problem'"

Faraz Sanei

2022 | DOI: 10.2139/ssrn.4076046

[Google Scholar](#)

[12] "University–Informal Education Institution Partnerships for Teacher Education"

Jenna Sanei

2019 | Proceedings of the 2019 AERA Annual Meeting | DOI: 10.3102/1435352

[Google Scholar](#)

[13] "Normativity, power, and agency: On the chronotopic organization of orthographic conventions on social media"

Taraneh Sanei

2021 | Language in Society | DOI: 10.1017/S0047404521000221

[Google Scholar](#)

[14] "TEXTURE SEGMENTATION USING SEMI-SUPERVISED SUPPORT VECTOR MACHINES"

SAEID SANEI

2004 | International Journal of Computational Intelligence and Applications | DOI: 10.1142/S1469026804001197

[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: bbo48952-oeob-4b11-9de3-91cob7ce31f3