

Ludzka przewaga w epoce AI: odpowiedzialne podejmowanie decyzji według metody Cheryl Strauss Einhorn



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje pułapkę „algorytmicznej wyroczni”, ostrzegając przed myleniem płynności generowania odpowiedzi przez AI z rzeczywistym rozumieniem i odpowiedzialnością. Przywołując koncepcję „Human Edge” Cheryl Strauss Einhorn, autor podkreśla, że podczas gdy AI może wspierać operacje poznawcze, ostateczna decyzja musi pozostać aktem ludzkim, osadzonym w wartościach i kontekście społecznym. Tekst przedstawia ośmiostopniowy proces rozwiązywania problemów, w którym człowiek zachowuje ster nad technologią. W obliczu standardów LIST i OECD, podejście Einhorn jest prezentowane nie jako poradnik produktywności, lecz jako dyscyplina poznawcza niezbędna do ochrony ludzkiej autonomii przed wygodną abdykacją na rzecz systemów optymalizacyjnych.

The article analyzes the 'algorithmic oracle' trap, warning against confusing AI's fluency in generating answers with actual understanding and responsibility. Invoking Cheryl Strauss Einhorn's 'Human Edge' concept, the author emphasizes that while AI can support cognitive operations, the final decision must remain a human act, embedded in values and social context. The text presents an eight-step problem-solving process in which humans maintain control over technology. In light of LIST and OECD standards, Einhorn's approach is presented not as a productivity guide, but as a cognitive discipline essential for protecting human autonomy against convenient abdication to optimization systems.

W dobie powszechnej fascynacji sztuczną inteligencją wpadamy w pułapkę „algorytmicznej wyroczni”, myląc płynność językową maszyny z rzeczywistym rozumieniem i odpowiedzialnością. Niniejszy artykuł stawia tezę, że ludzka przewaga w

erze AI nie polega na rywalizacji z technologią w zakresie szybkości przetwarzania danych, lecz na rygorystycznym zachowaniu sprawczości poprzez świadomy proces decyzyjny. Autor analizuje ośmioetapową metodę Cheryl Strauss Einhorn – od precyzyjnej definicji problemu i rozpoznania motywacji, przez zarządzanie kontekstem i walkę z uprzedzeniami, aż po włączenie interesariuszy i osiągnięcie osobistego przekonania (conviction). Tekst dowodzi, że AI powinna pełnić rolę potężnego sparingpartnera i lustra poszerzającego pole widzenia, a nie protezy myślenia. Kluczowym przesłaniem jest rozróżnienie między operacją poznawczą (rekomendacją), którą może wygenerować algorytm, a aktem odpowiedzialności (decyzją), który pozostaje wyłącznie ludzkim obowiązkiem. To wezwanie do intelektualnej dyscypliny, która chroni nas przed wygodną abdykacją z roli decydenta na rzecz sformatowanego raportu.

SŁOWA KLUCZOWE

ludzka przewaga

odpowiedzialne podejmowanie decyzji

metoda Cheryl Strauss Einhorn

algorytmiczna wyrocznia

human edge

proces decyzyjny w epoce AI

błąd automatyzacji

higiena poznawcza

metoda AREA

analiza przedśmiertna premortem

system społeczno-techniczny

pułapki optymalizacji

zarządzanie ryzykiem AI

sprawczość ludzka

Koniec mitu algorytmicznej wyroczni

Koniec mitu algorytmicznej wyroczni. Największym nieporozumieniem epoki sztucznej inteligencji nie jest to, że człowiek zbyt wiele oczekuje od maszyny, lecz że zbyt mało oczekuje od siebie. AI przeniknęła do kultury organizacyjnej, biznesowej, naukowej i administracyjnej jako narzędzie niemal magiczne: szybkie, elokwentne, niewyczerpane, pozornie cierpliwe i neutralne.

Wystarczy zapytać, a odpowiada. Doprecyzować – poprawia. Polecić – produkuje warianty, streszczenia, scenariusze, tabele, strategie czy symulacje. Maszyna nie marszczy brwi, nie mówi „nie wiem”, nie prosi o kawę i nie uraża się, gdy po raz setny prosimy o „bardziej elegancką wersję”.

Idealny urzędnik? Idealny analityk? Idealny doradca? Właśnie tu zaczyna się kłopot. Narzędzie, które odpowiada zawsze, łatwo pomylić z narzędziem, które wie. To, co mówi płynnie, z tym, co rozumie. System porządkujący informacje – z podmiotem ponoszącym odpowiedzialność za ich użycie. A przecież między generowaniem rekomendacji a podjęciem decyzji istnieje różnica zasadnicza: rekomendacja jest operacją poznawczą, decyzja zaś aktem odpowiedzialności. AI może uczestniczyć w pierwszej; druga pozostaje ludzka.

Właśnie dlatego książka Cheryl Strauss Einhorn „The Human Edge: Smarter Decisions in the Age of AI” jest interesująca nie jako kolejny poradnik o „sprytnych promptach”, lecz jako próba obrony ludzkiej sprawczości przed wygodną abdykacją. Opis wydawcy przedstawia ją jako wezwanie do działania dla tych, którzy chcą z pomocą AI przeprowadzić, a nie jedynie podążać za technologią. To rozróżnienie jest kluczowe: można bowiem korzystać z AI jak z protezy myślenia albo jak z narzędzia intensyfikującego własny osąd. Różnica między jednym a drugim to różnica między człowiekiem prowadzonym przez system a tym, który potrafi go zaprząć do własnego celu.

Einhorn trafnie wydobywa z tej koncepcji rdzeń, rozkładając proces rozwiązywania złożonych problemów na osiem „unikalnie ludzkich” kroków: definicję problemu, motywację, kontekst, badania, analizę, zwalczanie uprzedzeń, włączanie interesariuszy oraz podejmowanie ostatecznej decyzji. W każdym z tych etapów AI może pomagać, ale nie powinna przejmować steru. Maszyna przyspiesza pracę na danych, wskazuje wzorce i symuluje ryzyka, lecz nie potrafi wiedzieć, co jest dla nas naprawdę ważne, jakie kompromisy są dopuszczalne, a jakie byłyby zdradą celu, ani kto poniesie konsekwencje decyzji, gdy ekran zgaśnie. To punkt wyjścia do poważnej rozmowy o AI: nie od zachwyty nad szybkością, lecz od pytania o odpowiedzialność.

Współczesna kultura technologiczna ma skłonność do mylenia przyspieszenia z postępem. Zakłada się, że jeśli coś można zrobić szybciej, to można to zrobić lepiej; że skoro system analizuje więcej danych niż człowiek, jego rekomendacja jest bardziej racjonalna. Gdy model brzmi pewnie, użytkownik ulega pokusie przeniesienia tej pewności z formy wypowiedzi na jej treść. To psychologicznie zrozumiałe, lecz intelektualnie ryzykowne. Błyskotliwy sofista także mówi płynnie, ale gładkość zdania nie gwarantuje prawdy, a elegancja rekomendacji nie jest dowodem mądrości.

Dlatego koncepcja „human edge” powinna być odczytywana nie jako sentymentalna pochwała człowieka przeciw maszynie, lecz jako dyscyplina poznawcza. Ludzka przewaga nie polega na tym, że zawsze myślimy lepiej od AI – często tak nie jest. Bywamy leniwi, stronniczy, pyszni i ślepi na dane zaprzeczające naszym przekonaniom. Przewaga polega na czymś innym: tylko człowiek rozpoznaje, że decyzja nie jest czystym rachunkiem optymalizacyjnym, lecz aktem osadzonym w wartościach, relacjach i ryzykach, których nie da się zredukować do rankingu wariantów.

Tu książka Einhorn wpisuje się w szerszy paradygmat odpowiedzialnej sztucznej inteligencji. Amerykański NIST opracował AI Risk Management Framework jako ramę zarządzania ryzykiem dla jednostek i organizacji. Nie chodzi w nim o fetyszyzację technologii, lecz o rozumienie AI jako systemu społeczno-technicznego, którego skutki zależą od danych, kontekstu wdrożenia i nadzoru ludzi. Podobnie zasady OECD akcentują, że systemy AI powinny respektować rządy prawa, godność, autonomię i sprawiedliwość społeczną.

Jest to istotne, ponieważ Einhorn nie proponuje prywatnej sztuczki produktywnościowej, lecz dotyka kwestii ustrojowej: kto ma władzę nad decyzją w epoce narzędzi podsuwających gotowe odpowiedzi? W małej skali jest to pytanie menedżera czy prawnika; w dużej – instytucji publicznych, sądów i mediów.

Jeżeli nie zachowamy umiejętności definiowania problemu i kwestionowania odpowiedzi maszyny, AI nie tyle odbierze nam wolność, ile pomoże nam ją oddać – schludnie i w ładnie sformatowanym raporcie. Pierwszy krok, czyli definicja problemu, jest zatem kluczowy. Zła definicja działa jak źle ustawiony kompas: można iść sprawnie i imponować tempem marszu, ale nadal zmierzać w niewłaściwym kierunku.

Pułapki optymalizacji i pierwsze kroki Einhorn

Notatka źródłowa ilustruje to na przykładzie Jamesa, właściciela kawiarni, który początkowo poprosił AI o znalezienie "świetnej powierzchni komercyjnej do wynajęcia w hrabstwie". Otrzymał propozycje logiczne z punktu widzenia ogólnego zapytania, lecz niepraktyczne w realiach życia: zbyt odległe, nieuwzględniające zespołu, lokalizacji konkurencji, dojazdu czy warunków biznesowych. Dopiero gdy precyzyjnie zdefiniował problem — wskazując obecną lokalizację, metraż, dostęp do transportu publicznego, czas dojazdu i konieczność unikania bezpośredniej konkurencji — AI stała się narzędziem użytecznym. Ten przykład jest prosty, ale niesie głębokie znaczenie.

AI nie naprawia niejasności ludzkiego celu. Ona ją skaluje. Jeśli człowiek zada pytanie mętne, otrzyma odpowiedź mętnie trafną: formalnie pomocną, praktycznie jałową. Jeśli pominie ograniczenia, model ich nie odgadnie lub zrobi to błędnie. Jeśli użytkownik nie wie, czy chce maksymalizować zysk, minimalizować stres, chronić zespół, budować reputację, ograniczać ryzyko czy zachować czas dla rodziny, system nie wybierze "właściwej" odpowiedzi. Wybierze tę pasującą do otrzymanych danych i statystycznych wzorców, na których został zbudowany. To za mało, gdy stawką są realne decyzje.

Pierwszy krok Einhorn jest zatem niemal antytechnologiczną lekcją korzystania z technologii: zanim zapytasz maszynę, zapytaj siebie. Nie dlatego, że introspekcja jest zawsze nieomylna — nie jest; człowiek potrafi świetnie okłamywać samego siebie i w tej konkurencji ma nad AI wielowiekowe doświadczenie. Jednak bez próby samookreślenia AI zostaje zmuszona do zgadywania. A zgadywanie przez system o wysokiej płynności językowej jest szczególnie niebezpieczne, bo brzmi jak wiedza. W dawnej biurokracji mówiono: papier wszystko przyjmie. W epoce generatywnej trzeba dodać: model wszystko dopowie.

Drugim krokiem jest motywacja. Definicja problemu odpowiada na pytanie, co chcę rozwiązać; motywacja pyta, dlaczego w ogóle chcę to zrobić i jaka potrzeba stoi za tym wysiłkiem. W notatce źródłowej krok ten przedstawiono jako "rdzennie ludzki", ponieważ AI nie rozumie wartości, emocji, priorytetów życiowych i długoterminowych celów tak, jak doświadcza ich człowiek. Może je opisać, sklasyfikować czy zasymulować, ale nie może ich posiadać. Jeżeli użytkownik nie poda swojego "dlaczego", otrzyma odpowiedzi ogólnikowe, powierzchowne i niekoniecznie zgodne z tym, co naprawdę uważa za sukces.

Przykład Jamesa wraca tu w bardziej dojrzałej formie. Nie chodziło mu po prostu o otwarcie drugiej kawiarni, lecz o przyspieszenie niezależności finansowej i przygotowanie do emerytury — bez nadmiernego ryzyka, które mogłoby zniszczyć jego spokój i plany życiowe. Gdyby kierować się wyłącznie suchą optymalizacją, atrakcyjny lokal przy ruchliwej ulicy mógłby wydać się najlepszą decyzją. Jednak w świetle jego motywacji lepszy okazał się wariant spokojniejszy, tańszy i mniej ryzykowny.

Tu ujawnia się granica algorytmicznej racjonalności: to, co "optymalne" w abstrakcyjnym modelu, może być niewłaściwe w konkretnym życiu. To rozróżnienie ma kluczowe znaczenie dla ekonomii, prawa, zarządzania i polityki publicznej. Instytucje często udają, że podejmują decyzje na podstawie czystych wskaźników: firmy mówią o efektywności, administracje o procedurach, uczelnie o parametryzacji, platformy o optymalizacji, a politycy o "twardych danych", gdy akurat pasują one do ich tezy. Tymczasem każda decyzja zawiera ukryty porządek wartości. Optymalizować można tylko względem czegoś; maksymalizować można tylko konkretną wartość; minimalizować można pewien rodzaj kosztu, zwykle pomijając inny. AI nie rozwiązuje tego problemu — ona jedynie zmusza do większej uczciwości: powiedz, co naprawdę liczysz, a pokażę ci wynik. Nie powiesz — policzę za ciebie cudzą filozofię ukrytą w danych.

Trzecim krokiem jest kontekst, czyli środowisko, w którym zapada decyzja. Notatka źródłowa trafnie opisuje go jako "ogólny stan zdrowia decyzji". To celna metafora: decyzja pozbawiona kontekstu przypomina diagnozę lekarską wydaną na podstawie jednego objawu. Kaszel może oznaczać infekcję, alergię, przemęczenie, zanieczyszczone powietrze, stres albo chorobę przewlekłą.

Podobnie decyzja biznesowa, zawodowa czy instytucjonalna nie istnieje w próżni. Zależy od miejsca, czasu, ludzi, relacji władzy, zasobów, ograniczeń prawnych, etapu życia, kultury organizacyjnej i tego, kto realnie ponosi konsekwencje. Einhorn proponuje tu ramę "sytuacyjności", obejmującą lokalizację, ludzi, własność decyzji i etap życia. To praktyczne i filozoficznie nośne rozróżnienie. Lokalizacja przypomina, że rozwiązania zależą od miejsca: inne działają w centrum metropolii, inne w małym mieście, organizacji rozproszonej czy zespole pracującym twarzą w twarz. Ludzie przypominają, że decyzja ma adresatów i ofiary uboczne. Własność decyzji wskazuje, kto ma prawo rozstrzygnąć sprawę i kto będzie z niej rozliczany. Etap życia uświadamia, że ten sam poziom ryzyka inaczej wygląda dla młodego przedsiębiorcy, inaczej dla osoby blisko emerytury, rodzica małego dziecka czy lidera instytucji, który buduje dziedzictwo zamiast kolejnego slajdu do prezentacji.

Kontekst i pułapka nieskończonych badań

AI operuje „poza twoim ekosystemem”, dopóki użytkownik go nie opisze. To zdanie powinno wisieć nad każdym stanowiskiem pracy, przy którym używa się narzędzi generatywnych. Model nie zna nieformalnych konfliktów w zespole, chyba że mu je wyłożymy. Nie zna lokalnej kultury organizacyjnej. Nie wie, kto posiada realne weto, choć formalnie go nie ma. Nie dostrzega, że pracownicy magazynu nie ufają nowemu systemowi, że dział sprzedaży obchodzi procedury, że partner biznesowy gra na czas, a klient mówi „tak”, podczas gdy jego ciało krzyczy „uciekaj”. AI może pomóc sformułować pytania o te zjawiska, lecz nigdy nie zastąpi ludzkiego rozpoznania terenu. Mapa nie jest krajobrazem, a model nie jest spotkaniem. Dopiero na tym fundamencie sensownie pojawia się czwarty krok: badania.

AI błyszczy w gromadzeniu informacji, porównywaniu opcji, streszczaniu dokumentów czy syntetyzowaniu źródeł. Jednak badania pozbawione celu i kontekstu są jedynie elegancką formą błędzenia. Notatka źródłowa wskazuje, że Einhorn dzieli proces badawczy na przeglądanie, zawężanie i głębokie zanurzenie: najpierw poznajemy krajobraz, potem stosujemy kryteria, a na koniec badamy ścieżki najbardziej dopasowane. Brzmi to banalnie tylko dla tych, którzy nigdy nie utonęli w nadmiarze informacji. Dziś bowiem toni się nie z braku danych, lecz z braku kryteriów ich odrzucania. W epoce AI badanie staje się więc nie tyle sztuką znajdowania informacji, co sztuką kończenia poszukiwań. To jeden z najbardziej niedocenianych aspektów ludzkiej przewagi.

Maszyna się nie męczy. Może generować kolejne raporty, warianty i ostrzeżenia w nieskończoność. Człowiek natomiast musi rozpoznać moment, w którym dodatkowa informacja nie zwiększa już jasności, lecz karmi lęk przed decyzją. Należy przejść do analizy wtedy, gdy dysponujemy kilkoma

wiarygodnymi źródłami, dane pokrywają się z priorytetami, potrafimy uzasadnić wybór, a kolejne wyszukiwania przynoszą jedynie powtórzenia. Jest to kluczowe, gdyż współczesny paraliż decyzyjny często nosi kostium odpowiedzialności: „Jeszcze sprawdzam”, „Jeszcze analizuję”, „Jeszcze poproszę model o dodatkowe scenariusze”. Czasem jest to rozsądek. Częściej – elegancko ubrany strach. W pewnym momencie brak decyzji również staje się decyzją, tyle że podjętą tchórzliwie, bez podpisu i bez protokołu.

AI może ten mechanizm wzmacniać, oferując nieskończony dostęp do kolejnych „jeszcze”, z których każde brzmi profesjonalnie. W końcu nie tak nie maskuje bezzadności jak dobrze sformatowany dokument roboczy. Prowadzi to do wniosku zasadniczego: AI nie zastępuje procesu decyzyjnego, lecz wystawia go na próbę. Człowiek, który nie umie definiować problemów, będzie szybciej produkował złe odpowiedzi. Ten, który nie zna własnej motywacji, sprawniej zoptymalizuje cudze cele. Kto pomija kontekst, otrzyma rekomendacje piękne i nieprzydatne. A kto nie potrafi przerwać badań, zamieni narzędzie poznania w maszynę do odkładania odpowiedzialności.

Nie jest to jednak argument przeciwko AI, lecz za jej dojrzałym użyciem. Jest ona zbyt potężna, by korzystać z niej bezmyślnie. Potrzebuje człowieka nie jako biernego operatora, lecz jako decydenta: kogoś, kto wie, o co pyta, dlaczego to robi i kiedy odpowiedź należy zakwestionować. Badania dostarczają materiału, ale dopiero analiza nadaje mu znaczenie. To rozróżnienie uderza w bardzo współczesne złudzenie, że więcej informacji automatycznie przybliży nas do mądrości. Nie przybliża. Informacja jest surowcem, dane są masą, a raport jedynie uporządkowaną masą. Wiedza zaczyna się dopiero wtedy, gdy ktoś potrafi określić, co z tej masy wynika, czego nie wolno z niej wywnioskować i gdzie kończy się fakt, a zaczyna dekoracyjny szum.

AI może błyskawicznie przesiewać dokumenty i liczyć prawdopodobieństwa, ale nie rozstrzygnie, co w konkretnej sytuacji ma sens. Analiza jest piątym krokiem procesu, lecz w istocie stanowi pierwszy moment, w którym widać rzetelność wcześniejszych etapów. Jeśli problem został źle zdefiniowany, analiza będzie elegancką egzegezą pomyłki. Jeśli motywacja była niejasna, zacznie optymalizować cele, których nikt świadomie nie wybrał. Przy pominiętym kontekście stanie się laboratoryjnie czysta i życiowo bezużyteczna. A jeśli badania były chaotyczne, analiza wyprodukuje sens pozorny, podobny do horoskopu dla zarządu: brzmi trafnie, bo jest dostatecznie ogólny.

W źródle analiza zostaje określona jako krok „unikalnie ludzki”. Jej celem jest nadanie sensu i właściwa interpretacja danych. AI wykonuje ciężką pracę obliczeniową – przetwarza zbiory, identyfikuje wzorce i prognozuje trendy – ale to człowiek musi nadać temu ostateczny kierunek.

Rola człowieka w procesie analizy danych

Człowiek nadaje analizie kierunek i znaczenie, gdyż to on ocenia, czy wnioski odpowiadają jego celom, wartościom i priorytetom. To punkt zasadniczy: analiza nie jest samym przetwarzaniem informacji, lecz interpretacją ukierunkowaną przez pytanie. Gdy pytanie jest ubogie, nawet najlepsza analiza pozostanie uboga. Gdy jest błędne, może stać się instrumentem doskonalenia błędu. Jeśli podszyte jest lękiem – wyniki będą służyć lękowi; jeśli próżnością – dane zostaną ułożone w procesję ku tezie, którą decydent pokochał, zanim jeszcze zaczął myśleć. Człowiek ma zatem obowiązek nie tylko pytać AI, ale poddawać rewizji własne pytania: skąd się wzięły, komu służą, co zakładają i czego nie dopuszczają.

Einhorn zwraca uwagę na konieczność dopasowania typu analizy do celu. Inaczej bada się liczby, inaczej opinie czy ryzyka prawne, inaczej dynamikę zespołu, sentyment klientów, skutki finansowe bądź zgodność regulacyjną. Autorka wskazuje na szerokie spektrum narzędzi, które można zlecić AI: od analiz opisowych i diagnostycznych, przez predykcyjne, aż po analizy sentymentu, przestrzenne czy finansowe. Nie istnieje „analiza w ogóle” – zawsze jest to analiza czegoś konkretnego, pod określonym kątem, dla danego celu i przy użyciu wybranej metody. Choć brzmi to technicznie, w praktyce jest głęboko humanistyczne: oznacza bowiem, że nie wolno udawać neutralności tam, gdzie dokonuje się wyboru soczewki.

Analiza finansowa może wskazać rozwiązanie najtańsze, podczas gdy społeczna ujawni, że to samo rozwiązanie niszczy zaufanie zespołu. Analiza prawna uzna działanie za dopuszczalne, lecz reputacyjna pokaże, że dopuszczalność nie oznacza roztropności. Analiza przestrzenna wskaże dogodną lokalizację, a demograficzna ukaże, że za trzy lata otoczenie będzie zupełnie inne. Z kolei analiza sentymentu ujawni nastroje niewidoczne w twardych wskaźnikach. Tylko barbarzyńca poznawczy mówi: „dane pokazały”. Dane same w sobie niczego nie pokazują. Trzeba je zapytać, obejrzeć pod światło, zestawić i sprawdzić, czy przypadkiem nie odpowiadają na pytanie, którego nikt rozsądny nie powinien był zadać.

Właśnie dlatego Einhorn ostrzega przed opieraniem się na jednym typie analizy. Człowiek zwykle wybiera ten rodzaj danych, z którym czuje się najbezpieczniej: zwolennik liczb ignoruje narrację, uciekający w intuicję lekceważy statystykę. Prawnik dostrzeże ryzyko formalne, marketingowiec przekaz, finansista koszt, a działacz społeczny wpływ na ludzi. Problem polega na tym, że rzeczywistość ma złośliwy zwyczaj istnienia jednocześnie w wielu wymiarach; nie daje się grzecznie zamknąć ani w arkuszu kalkulacyjnym, ani w anegdocie. Dobry proces decyzyjny wymaga zatem syntezy różnych analiz. W przykładzie wyboru uczelni analiza porównawcza może wskazać szkołę z najniższym czesnym i najlepszym odsetkiem absolwentów, ale dopiero analiza sentymentu ujawni

skargi studentów na izolację. To nie jest szczegół – dla jednego kandydata taka cena będzie akceptowalna, dla innego stanie się czynnikiem niszczącym całe doświadczenie edukacyjne. Dane twarde mówią, co jest mierzalne; dane miękkie często mówią, co będzie odczuwalne. A człowiek nie żyje w tabeli, tylko w konsekwencjach tabeli.

Metoda AREA jako higiena poznawcza

Ten moment prowadzi bezpośrednio do metody AREA – ustrukturyzowanej ramy podejmowania decyzji, opracowanej przez Cheryl Strauss Einhorn w celu ochrony przed błędami poznawczymi, skrótami myślowymi i martwymi punktami. AREA ma za zadanie spowolnić decydenta, oddzielić źródła informacji i zmusić go do spojrzenia na problem z wielu perspektyw, zanim podejmie on jakiegokolwiek działanie. Akronim AREA odnosi się do czterech perspektyw: Absolute, Relative, Exploration and Exploitation oraz Analysis. W polskim porządku można je określić jako: informacje pierwotne, informacje wtórne, eksploracja i eksploatacja oraz analiza.

Ten schemat jest prosty tylko na pierwszy rzut oka; w istocie zawiera on bardzo dojrzałą teorię higieny poznawczej. Nakazuje oddzielić to, co pochodzi bezpośrednio ze źródła, od tego, co jest już jedynie jego interpretacją. Wymusza wyjście poza znane kanały informacji, rozmowy z ludźmi i poszukiwanie nieoczywistych danych, by następnie zwrócić ostrze krytyki przeciw sobie: własnym uprzedzeniom, założeniom, wygodom i przyzwyczajeniom. Dopiero potem pozwala przejść do syntezy. Ta kolejność jest kluczowa, gdyż współczesny człowiek często analizuje, zanim naprawdę zbadał. Najpierw formuje opinię, a dopiero potem szuka dowodów. Najpierw czuje, kto ma rację, a następnie dobiera argumenty. Najpierw chce konkretnego wyniku, potem zamawia u świata uzasadnienie.

AI może przyspieszyć ten proces do prędkości niebezpiecznej. Wystarczy sformułować tendencyjne pytanie, aby otrzymać odpowiedź, która będzie tendencyjnie przydatna. Model nie zawsze zauważy i powie: „Drogi użytkowniku, twoje pytanie jest tezą przebraną za zapytanie”. A szkoda, bo czasem byłaby to najbardziej potrzebna funkcja: elektroniczny policzek epistemiczny, mały alarm przeciw samozadowoleniu.

Metoda AREA działa jak antidotum na tę skłonność. Etap Absolute każe wrócić do podstaw: dokumentów, danych pierwotnych, bezpośrednich obserwacji, rozmów ze źródłem i materiałów nieprzetworzonych. W świecie opinii z drugiej ręki jest to gest niemal konserwatywny: zanim uwierzysz interpretatorom, zobacz, co właściwie interpretują. Etap Relative pozwala poszerzyć pole o źródła wtórne – ekspertów, raporty, komentarze, opracowania i cudze punkty widzenia. Etap Exploration and Exploitation otwiera proces na twórcze poszukiwanie, ale także na konfrontację z

własnymi błędami. Ostatni etap, Analysis, nie jest więc pierwszą reakcją, lecz dojrzałą syntezą po przejściu przez różne porządki poznania.

Tu ujawnia się głębsza wartość koncepcji Einhorn: ona nie uczy tylko, jak używać AI, ale przede wszystkim, jak nie zostać przez nią rozleniwionym poznawczo. To różnica zasadnicza. Sztuczna inteligencja może wzmocnić najlepsze praktyki intelektualne człowieka, ale może też uczynić z niego istotę jeszcze bardziej powierzchowną: kogoś, kto już nie czyta, lecz prosi o streszczenie; nie porównuje źródeł, lecz pyta o „najważniejsze wnioski”; nie formułuje hipotezy, lecz wybiera najbardziej przekonującą wersję wygenerowanego tekstu. AI może być warsztatem myślenia albo fast foodem poznania. Ta sama technologia, dwie różne antropologie.

Zwalczanie uprzedzeń i pułapka automatyzacji

Szczególnie istotny jest szósty krok procesu, poświęcony zwalczaniu uprzedzeń. Notatka źródłowa wskazuje, że problem błędów poznawczych w epoce AI ma charakter dwuwymiarowy: błędy popełniają zarówno ludzie, jak i maszyny. Choć ich źródła są różne, skutki mogą się nakładać i wzajemnie potęgować. Człowiek ulega efektowi potwierdzenia, zakotwiczenia, dostępności czy automatyzacji; AI natomiast bywa obciążona błędami danych, algorytmicznymi lub wdrożeniowymi. Największe zagrożenie pojawia się w momencie, gdy uprzedzony człowiek spotyka uprzedzony system i obaj zaczynają sobie wzajemnie przytakiwać. To jedna z najcięższych diagnoz współczesności.

Przed erą AI człowiek mógł błędzić samodzielnie. Robił to z rozmachem, godnym lepszej sprawy. Szukał potwierdzeń, ignorował niewygodne fakty, ufał pierwszemu wrażeniu i z uporem bronił decyzji, które należało porzucić trzy zebrania wcześniej. AI wprowadza jednak nowy element: potrafi nadać ludzkiemu błędowi pozory obiektywności. Jeśli system potwierdzi moje przecucie, łatwo uznać, że „to nie ja, to analiza”. Gdy model wskaże kandydata, którego i tak chciałem wybrać, można ogłosić zwycięstwo neutralnej technologii. Jeśli narzędzie oceni ryzyko zgodnie z moimi obawami, mogę nazwać własny lęk predykcją.

W tym sensie błąd automatyzacji jest szczególnie groźny. Polega on na bezkrytycznym zaufaniu systemom, zwłaszcza gdy działają sprawnie i beznamiętnie. Człowiek chętnie oddaje ciężar odpowiedzialności temu, co brzmi pewnie – dawniej była to pieczęć urzędu, autorytet eksperta czy majestat procedury; dziś bywa nim syntetyczny głos modelu. Przykładem jest Alex, prawnik, który użył AI do znalezienia argumentów i otrzymał nieistniejący cytat oraz fikcyjną sprawę. Błąd nie polegał jedynie na tym, że AI „halucynowała”, lecz na tym, że człowiek potraktował tę halucynację

jako materiał procesowy. Maszyna zmyśliła, człowiek uwierzył. Katastrofa poznawcza zwykle potrzebuje zespołu.

To rozróżnienie jest kluczowe w debacie publicznej. Zbyt często mówi się o ryzyku AI tak, jakby problemem była sama technologia, a człowiek pozostawał niewinną ofiarą algorytmu. To wygodna bajka. W rzeczywistości ryzyka powstają na styku: człowieka, który chce szybko; organizacji, która chce tanio; systemu działającego statystycznie i danych historycznie obciążonych. AI nie tyle tworzy te patologie, co pozwala im działać szybciej, szerzej i ciszej.

Drugim przykładem jest Pam, administratorka w kuratorium, testująca system typujący uczniów zagrożonych problemami w nauce. System zbyt często wskazywał osoby z konkretnych szkół, ponieważ trenowano go na danych nieuwzględniających różnorodności populacji, m.in. uczniów dwujęzycznych czy dotkniętych problemami społeczno-ekonomicznymi. Gdyby Pam zaufała systemowi bezkrytycznie, utrwaliłaby istniejące nierówności. AI może zatem nie tylko popełnić błąd, ale nadać starej niesprawiedliwości nowoczesną twarz. Dane historyczne nie są neutralnym zapisem świata, lecz zapisem pełnym wykluczeń i instytucjonalnych ślepot. Jeśli model uczy się na przeszłości, reprodukuje ją pod postacią prognozy. Wtedy system nie mówi: „tak było”, lecz: „tak prawdopodobnie będzie”. Tu zaczyna się polityczna i etyczna waga AI – predykcja może stać się samospełniającym wyrokiem.

Einhorn proponuje cztery strategie ochronne: konfrontowanie oczekiwań, testowanie narzędzia, zadawanie pytań „dlaczego?” oraz weryfikację krzyżową. W praktyce wymaga to wyrobienia w sobie nawyku poznawczej nieufności. Nie cynizmu, bo cynizm jest tylko lenistwem w eleganckim płaszczu. Chodzi o nieufność metodyczną: zanim przyjmę wynik, zapisuję własne oczekiwania. Gdy wynik je potwierdza, pytam, czy go nie sprowokowałem. Gdy im przeczy, sprawdzam, czy nie wskazuje czegoś istotnego. Proszę AI o uzasadnienie, zmieniam parametry, porównuję odpowiedzi i konsultuję się z ludźmi.

Warto zatrzymać się przy pytaniu „dlaczego?”. To jedno z najprostszych i najbardziej niewygodnych pytań w pracy z AI. Modele generatywne odpowiadają płynnie, ale wymuszenie uzasadnienia ujawnia, czy mamy do czynienia z realnym łańcuchem rozumowania, czy jedynie z ładnie ułożonym rezultatem. Pytanie to rozbija iluzję magicznego pudełka i obnaża luki: brak źródeł, skoki logiczne czy nieadekwatne analogie. Dobrze używana AI nie powinna być maszyną do dawania odpowiedzi, lecz partnerem w produkowaniu lepszych pytań.

Weryfikacja krzyżowa jest natomiast elementem, bez którego korzystanie z AI pozostaje intelektualnym hazardem. Nie chodzi o sprawdzanie każdej drobnostki, lecz o rozpoznawanie twierdzeń wysokiego ryzyka: danych liczbowych, przepisów prawa, cytatów czy sygnatur orzeczeń.

Tam AI nie może być ostatnim słowem; może być pierwszą mapą – czasem dobrą, a czasem taką z pięknie narysowanym mostem, którego nigdy nie zbudowano.

W tym miejscu warto połączyć analizę, metodę AREA i walkę z błędami poznawczymi w jedną zasadę: decyzja wspomagana przez AI musi mieć procedurę oporu wobec własnej wygody. To zdanie może brzmieć surowo, ale jest praktyczne. Każda organizacja używająca AI powinna pytać: gdzie w naszym procesie jest miejsce na zakwestionowanie rekomendacji? Kto ma obowiązek sprawdzić źródła i kto odpowiada za dane wejściowe?

Odpowiedzialność autora a pułapka algorytmicznego alibi

Kto ma prawo stwierdzić, że wynik jest formalnie poprawny, lecz społecznie głupi? Kto decyduje o momencie zakończenia analizy? Kto bierze odpowiedzialność w chwili, gdy sformułowanie „system zasugerował” przestaje być argumentem, a staje się alibi? Bez tych pytań AI przeobrazi się w idealne narzędzie konformizmu. Pozwoli na bezpieczne stwierdzenia: „tak wskazał model”, „tak wyszło z analizy”, „algorytm uznał” lub „dane pokazały”. Często jednak za takimi zdaniami kryje się człowiek, który nie chce powiedzieć wprost: „to ja wybieram”, „to ja akceptuję ryzyko”, „to ja pomijam czyjeś interesy” lub „to ja uznaję ten koszt za dopuszczalny”. Jedną z najważniejszych funkcji koncepcji Einhorn jest więc przywrócenie decyzji twarzy. Decyzja nie jest emanacją systemu; decyzja ma autora.

Arkusze Geparda, zaprojektowane jako krótkie ćwiczenia i organizatory graficzne wspierające proces decyzyjny, wpisują się w tę samą filozofię. Ich nazwa odwołuje się do geparda nie ze względu na jego szybkość, lecz zdolność do gwałtownego hamowania i zmiany kierunku. To trafna metafora pracy z AI.

W świecie czczącym prędkość przewaga może polegać na umiejętności zatrzymania się. Nie każde przyspieszenie jest rozwojem, nie każda odpowiedź postępem, a nie każdy kolejny wariant przybliża do rozwiązania. Czasem najmądrzejszym ruchem jest pauza: sprawdzenie problemu, motywacji, kontekstu, źródeł i samego siebie. Arkusze Geparda pełnią rolę „strategicznych przystanków”. Zmuszają użytkownika, by nie pędził za pierwszą odpowiedzią AI, lecz nadał kierunek pracy maszyny. Ich wartość nie tkwi w biurokratycznej formalizacji procesu, ale w ochronie przed poznawczym poślizgiem. Każdy, kto pracował z generatywną AI, zna ten mechanizm: jedno pytanie prowadzi do drugiego, odpowiedź do kolejnego wariantu, a ten do następnego dokumentu. Po

godzinie człowiek dysponuje ogromem materiału, tracąc jednocześnie jasność co do celu początkowego. Arkusz działa tu jak hamulec.

Pyta: czego szukasz? dlaczego? w jakim kontekście? czego jeszcze nie wiesz? co może być błędem? kto zostanie dotknięty? kiedy skończysz? Zdolność do takiego hamowania jest dziś jedną z kluczowych kompetencji decyzyjnych. Nowoczesność kocha dynamikę, lecz instytucje umierają nie tylko z bezruchu, ale i od nieprzemyślanego ruchu. Firmy wdrażają systemy niezrozumiałe dla pracowników. Administracje cyfryzują procedury, których wcześniej nie uprościły. Szkoły kupują narzędzia, zanim ustalą cele dydaktyczne.

Organizacje społeczne komunikują strategię bez konsultacji z tymi, którzy mają je realizować. Politycy ogłaszają rozwiązania, zanim nazwą problem. Potem wszyscy dziwią się, że „dobry pomysł” napotkał opór. Jednak to nie opór był problemem, lecz pycha procesu. Analiza i AREA uczą czegoś przeciwnego: pokory strukturalnej. Nie chodzi tu o skromność jako ozdobę charakteru, lecz o pokorę wobec złożoności świata. Każda decyzja ma więcej wymiarów, niż widzi pierwszy entuzjasta. Każdy system ma interesariuszy dysponujących wiedzą niedostępną dla zarządu. Każde narzędzie posiada swój kontekst wdrożenia, każdy model – założenia i ograniczenia. Każdy człowiek ma swoje biasy, a każda organizacja ślepe pola. AI może pomóc je ujawnić, o ile człowiek nie używa jej jedynie do potwierdzania własnej wygody.

W tym sensie piąty i szósty krok procesu – analiza oraz zwalczanie uprzedzeń – stanowią serce intelektualnej odpowiedzialności. Definicja problemu, motywacja, kontekst i badania przygotowują materiał; analiza wydobywa z niego sens, a walka z uprzedzeniami sprawdza, czy ten sens nie jest jedynie dobrze ubraną projekcją. Dopiero po takim przejściu można poważnie mówić o włączaniu interesariuszy. Rozmowa z ludźmi nie powinna być bowiem rytuałem po podjęciu decyzji, lecz elementem jej dojrzewania.

Jeśli analiza pozostaje samotna, może być elegancka, lecz jeśli nie spotka się z interesariuszami, pozostanie ślepa. Tu otwiera się kolejny etap: decyzja wychodzi poza głowę decydenta i ekran narzędzia AI. Wchodzi w świat relacji, oporu, zaufania, nieformalnej wiedzy, lęków, interesów i odpowiedzialności zbiorowej. AI może symulować perspektywy interesariuszy, proponować pytania i przewidywać obiekcje, lecz nie zastąpi tego procesu.

Włączanie interesariuszy jako warunek racjonalności

Maszyna nie może jednak usiąść z człowiekiem przy stole, dostrzec jego reakcji, usłyszeć milczenia, rozpoznać napięcia ani odbudować zaufania tam, gdzie proces przeprowadzono ponad głowami

ludzi. Interesariusze, zaufanie i decyzja jako fakt społeczny – to nierozzerwalna całość. Decyzja nigdy nie jest tak samotna, jak lubi to sobie wyobrażać decydent. Nawet gdy formalnie podpisuje ją jedna osoba, żyją z nią inni: pracownicy, klienci, partnerzy, rodzina, wspólnicy, wyborcy, użytkownicy systemu, pacjenci, uczniowie, mieszkańcy, podwładni, przełożeni, konkurenci, regulatorzy oraz wszyscy ci, którzy nie zostali zaproszeni do stołu, lecz zostaną wezwani do ponoszenia skutków. Właśnie dlatego siódmy krok w procesie Einhorn – włączanie interesariuszy – nie jest grzecznościowym dodatkiem, lecz warunkiem racjonalności decyzji. Decyzja pozbawiona głosów tych, których dotknie, może być formalnie poprawna, analitycznie spójna i organizacyjnie szybka, a mimo to społecznie głucha.

W źródle włączanie interesariuszy zostaje bardzo trafnie odróżnione od zwykłego komunikowania decyzji. To rozróżnienie jest kluczowe. Komunikowanie oznacza: decyzja już zapadła, teraz trzeba ją sprzedać, wyjaśnić, opakować, wypolerować, osłodzić i przeprowadzić przez opór. Włączanie oznacza coś radykalnie innego: zanim wybór zostanie dokonany, należy rozpoznać perspektywy ludzi, których decyzja dotyczy – tych, którzy posiadają wiedzę praktyczną, będą ją wdrażać, mogą ją zablokować lub których zaufanie można stracić bezpowrotnie.

To nie jest różnica proceduralna, lecz antropologiczna. Komunikowanie po fakcie traktuje ludzi jako odbiorców decyzji. Włączanie przed decyzją uznaje ich za współuczestników rzeczywistości, której decydent sam nie widzi w całości. Pierwsze podejście jest typowe dla organizacji hierarchicznych i próżnych: elita wie, lud ma przyjąć. Drugie wymaga pokory poznawczej: ten, kto formalnie posiada władzę, niekoniecznie dysponuje pełnią wiedzy. W praktyce często bywa odwrotnie. Im wyżej siedzi się w strukturze, tym dalej jest się od szczegółu, tarcia, absurdu procedury, błędu interfejsu, nieformalnego obejścia, lęku pracownika i realnego sposobu działania systemu.

AI może być użyteczna na etapie pracy z interesariuszami, ale jej możliwości mają swoje granice. Może podpowiedzieć, kogo warto uwzględnić, pomóc stworzyć mapę wpływów i zależności czy przewidzieć możliwe obiekty różnych grup. Może zasugerować pytania do partnera, zespołu, klientów, pracowników, beneficjentów lub regulatorów. Potrafi przygotować kilka wersji komunikatu – ostrzejszą, spokojniejszą, bardziej ekspercką bądź empatyczną. Może odegrać rolę sceptycznego odbiorcy albo sfrustrowanego użytkownika, pomagając decydentowi wyjść poza własną głowę. Nie może jednak wykonać za niego rozmowy. AI potrafi symulować perspektywę człowieka, ale nie jest człowiekiem z interesem, lękiem, pamięcią, doświadczeniem upokorzenia, pozycją w zespole, niewypowiedzianą obawą i twarzą, która zdradza więcej niż słowa.

Model może wygenerować listę pytań do pracowników magazynu, lecz nie zauważy, że przy trzecim pytaniu zapadła cisza, bo wszyscy boją się powiedzieć prawdę w obecności kierownika. Może

przygotować komunikat do działu sprzedaży, ale nie wyczuje ironicznego uśmiechu człowieka, który od lat słyszy o „nowych usprawnieniach”, a potem sam łąta ich skutki. Może zagrać rolę partnera negocjacyjnego, lecz nie dostrzeże zmęczenia, urażonej godności, nieufności ani tego krótkiego momentu, w którym ton głosu decyduje o tym, czy rozmowa zmierza ku porozumieniu, czy ku protokołowi rozbieżności. Źródło ujmuje to precyzyjnie: ludzką przewagą na tym etapie jest empatia, dostosowanie tonu w czasie rzeczywistym, budowanie relacji i rozpoznanie ukrytej dynamiki władzy.

AI nie usiądzie z kimś przy stole. Nie zobaczy uniesionych brwi ani założonych rąk. Nie wyczuje, czy interesariuszy należy włączyć już teraz, czy najpierw trzeba uporządkować własne cele, aby nie rozmyć procesu. Zbyt wczesne włączenie może rozproszyć decyzję; zbyt późne – może ją pogrzebać.

Właśnie tutaj widać, że teoria decyzji nie jest tylko nauką o wyborze między wariantami, lecz także nauką o zaufaniu. Bez zaufania nawet najlepszy projekt staje się obcym ciałem w organizmie społecznym. Można go wdrożyć siłą procedury, ale nie można zmusić ludzi, by naprawdę wzięli za niego odpowiedzialność. Będą wykonywać minimum, sabotować milczeniem lub obchodzić system bocznymi ścieżkami. Będą mówić „tak” na spotkaniach i „to nie przejdzie” na korytarzu. Każdy, kto zna realne życie instytucji, wie, że oficjalna akceptacja bywa najtańszą walutą świata. Prawdziwa zgoda kosztuje więcej: wymaga wysłuchania, wpływu i poczucia, że człowiek nie jest meblem w sali wdrożeniowej.

Przykład Tracey pokazuje tę prawidłowość z niemal laboratoryjną prostotą. Dyrektorka operacyjna wdraża nowy system zarządzania zapasami, konsultując decyzję z bezpośrednim zespołem badawczym. Z perspektywy centrum decyzyjnego wszystko wygląda rozsądnie: problem rozpoznany, narzędzie wybrane, wdrożenie zaplanowane. Dopiero po uruchomieniu okazuje się, że pracownicy magazynu mają problemy z interfejsem, a dla działu sprzedaży system jest niekompatybilny. Pominięto tych, którzy mieli z systemem żyć na co dzień. Skutek? Frustracja, opóźnienia, błędy i organizacyjna utrata zaufania. Ta historia nie opisuje awarii technologii, lecz awarię procesu decyzyjnego. System mógł być dobry, analiza poprawna, a intencja racjonalna. Porażka polegała na tym, że wiedza praktyczna została uznana za drugorzędną.

To częsty błąd elit organizacyjnych: mylenie wiedzy abstrakcyjnej z wiedzą pełną. Tymczasem ten, kto projektuje system, widzi architekturę, a ten, kto go używa codziennie, widzi tarcie. Decyzja mądra musi znać jedno i drugie. Bez architektury mamy chaos; bez wiedzy o tarcu mamy papierową utopię. W tym sensie włączanie interesariuszy nie jest demokratycznym rytuałem dla dobrego samopoczucia, lecz metodą pozyskiwania wiedzy niedostępnej z centrum. Pracownik pierwszej linii zna rzeczy, których nie zna zarząd. Klient zna irytacje, których nie zna dział produktu. Pacjent zna skutki procedury, których nie widzi ministerialny wskaźnik. Uczeń zna

absurd regulaminu, którego nie przewidzi kuratorium. Obywatel zna realne działanie urzędu, którego nie da się odtworzyć z prezentacji o cyfryzacji. Jeżeli tych głosów zabraknie przed decyzją, pojawią się one po niej – jako opór, reklamacje, obojętność, skargi, odejścia, cynizm albo memy. A memy bywają najbardziej demokratyczną formą audytu.

Drugi przykład – historia Solomona – pokazuje włączanie interesariuszy w skali prywatnej. Solomon rozważa przyjęcie nowej oferty pracy. Mógłby potraktować sprawę jako indywidualny rachunek: pensja, prestiż, rozwój, elastyczność i przyszłość zawodowa. Zamiast tego chce porozmawiać z partnerką, mentorem i przyjacielem z branży. Co ważne, nie zaczyna od komunikatu: „mam lepszą pracę i będzie więcej pieniędzy”. Pyta AI, jakie pytania zadać partnerce, aby zrozumieć wpływ decyzji na wspólne priorytety. Dzięki temu rozmowa nie zaczyna się od autoprezentacji zwycięzcy, lecz od wysłuchania obawy Sandy o to, że będą spędzać ze sobą mniej czasu. Dopiero wtedy Solomon może odpowiedzieć na realny lęk, wskazując, że nowa praca daje elastyczny grafik i możliwość pracy z domu.

Ten przykład jest subtelny: pokazuje, że AI może pomóc człowiekowi przygotować się do empatii, ale nie może jej wykonać. Może podsunąć pytania, lecz nie usłyszy odpowiedzi w imieniu Solomona. Może wskazać, że partnerka może martwić się o czas, bliskość czy zmianę rytmu życia, ale nie sprawi, że Solomon naprawdę uzna te obawy za ważne. To jest jego zadanie.

AI może otworzyć drzwi do lepszej rozmowy, ale nie przejdzie przez nie za człowieka. W organizacjach i instytucjach publicznych ta lekcja ma jeszcze większe znaczenie. Decyzje dotyczące technologii, restrukturyzacji, polityk publicznych, regulacji, edukacji, zdrowia, transportu czy rynku pracy bardzo często przedstawiane są jako rozwiązania techniczne. A gdy coś zostanie nazwane technicznym, natychmiast znika z pola moralnego sporu: „To tylko system”, „To tylko procedura”, „To tylko optymalizacja”, „To tylko automatyzacja”.

Odpowiedzialność w relacjach z interesariuszami

Nieprawda. Za każdym „tylko” kryje się czyjaś praca, czas i dostęp, ale także wykluczenie, koszt, obowiązek, utrata kontroli lub przesunięcie władzy. Techniczność bywa parawanem aksjologii. Dlatego siódmy krok Einhorn należy rozumieć szerzej niż jako praktykę zarządzania zmianą; to w istocie teoria odpowiedzialnej decyzji w świecie złożonych zależności.

Interesariusze nie są przeszkodą w procesie, lecz częścią rzeczywistości, której ten proces dotyczy. Ich uwzględnienie nie musi oznaczać spełnienia wszystkich oczekiwań – byłaby to naiwność, a nie odpowiedzialność. Czasem decyzja musi zostać podjęta mimo sprzeciwu, gdy interesy są sprzeczne

lub gdy trzeba rozstrzygnąć między kosztem krótkoterminowym a dobrem długoterminowym. Czasem lider musi po prostu powiedzieć „nie”. Istnieje jednak zasadnicza różnica między decyzją podjętą po wysłuchaniu stron a decyzją narzuconą z pozycji wygodnej niewiedzy.

Włączanie interesariuszy pełni trzy funkcje. Pierwszą jest funkcja poznawcza: ujawnia informacje, których decydent nie posiada. Druga to funkcja relacyjna: buduje zaufanie i poczucie współodpowiedzialności. Trzecia zaś – legitymizacyjna: sprawia, że nawet niepopularna decyzja może zostać uznana za uczciwą, jeśli ludzie widzą, że ich głos nie był jedynie dekoracją. Jest to kluczowe w epoce narastającej nieufności wobec instytucji. W świecie, w którym podejrzewamy, że decyzje zapadają „gdzieś”, „przez kogoś” i „na podstawie algorytmów”, bez naszego udziału, procedura włączania staje się jednym z ostatnich narzędzi odbudowy wiarygodności.

AI może tu służyć jako wsparcie przygotowawcze. Może pomóc stworzyć mapę interesariuszy: określić, kto wdroży decyzję, kto posiada wiedzę bezpośrednią, kto ma prawo formalnego lub nieformalnego weta, kto poniesie koszt, a kto zyska. Pozwala zidentyfikować tych, którzy mogą zostać pominięci, choć powinni być zaproszeni do rozmowy. Notatka wskazuje, że Einhorn proponuje Arkusze Geparda służące właśnie mapowaniu kluczowych osób, dostosowaniu komunikatów do ich obaw oraz projektowaniu systemu informacji zwrotnej. Taka mapa nie może być jednak zwykłym ćwiczeniem z zarządzania projektem; powinna być mapą odpowiedzialności.

W każdej decyzji należy zadać pytania, które pełnią rolę testu przyzwoitości procesu: Kto będzie żył z konsekwencjami? Kto nie ma głosu, choć poniesie koszt? Kto wie coś z praktyki, czego nie wie autor strategii? Kto może zablokować wdrożenie nie ze złośliwości, lecz dlatego, że plan ignoruje jego warunki pracy? Kto zostanie postawiony przed faktem dokonanym? I kto usłyszy „to dla waszego dobra” z ust ludzi, którzy sami nie doświadczą skutków tej decyzji? W każdej instytucji pytania te działają jak latarka wniesiona do piwnicy. Nagle staje się widoczne to, co wcześniej wygodnie nazywano „szczegółami operacyjnymi”.

Interesariusze nie są grupą jednolitą. Niektórzy dysponują władzą formalną, inni praktyczną. Jedni mogą wydać zgodę, inni opóźnić wdrożenie; jedni dysponują kapitałem, inni reputacją lub wiedzą. Są interesariusze widoczni i niewidoczni, głośni i milczący, uprzywilejowani i zależni. Uporządkowany proces decyzyjny powinien szczególnie chronić tych ostatnich, gdyż systemy – zarówno ludzkie, jak i algorytmiczne – mają tendencję do wzmacniania głosów już silnych. AI, trenowana na dostępnych danych, również częściej „widzi” tych, o których pozostawiono najwięcej śladów w dokumentach.

Symulacja interesariuszy a realna podmiotowość

Milczenie grup słabszych w modelu łatwo staje się niewidzialnością. Prowadzi to do istotnego problemu etycznego: symulowanie interesariuszy przez AI może być pomocne, ale grozi zastąpieniem realnego słuchania substytutem. Decydent może wówczas stwierdzić: „poprosiłem model, żeby odegrał rolę pracowników, klientów, obywateli, pacjentów, rodziców czy uczniów”.

To dobry początek, lecz nie całość. Symulacja nie jest konsultacją. Model roli społecznej nie zastąpi społeczności, a persona marketingowa nie stanie się człowiekiem. Wygenerowana obiekcja nie jest żywym sprzeciwem. AI może pomóc przygotować się do rozmowy, ale jeśli staje się pretekstem, by jej uniknąć, zmienia się z narzędzia empatii w instrument eleganckiego paternalizmu. Można to nazwać „paternalizmem predykcyjnym”: zamiast pytać ludzi, przewidujemy, co zapewne powiedzieliby, gdybyśmy ich zapytali. Wersja miękka brzmi profesjonalnie: dysponujemy danymi, modelami, segmentami, profilami zachowań i mapami ryzyka. Wersja brutalna jest prosta: nie musimy was słuchać, bo system już was zasymulował. To nie jest odpowiedzialna technologia; to stara pycha w nowym interfejsie.

Dlatego włączanie interesariuszy wymaga rozróżnienia między przewidywaniem reakcji a uznaniem podmiotowości. Reakcję można przewidzieć bez szacunku. Uznanie podmiotowości wymaga natomiast dopuszczenia, że drugi człowiek może powiedzieć coś nieprzewidzianego, czego nie chcemy usłyszeć i czego nie da się łatwo wpisać w plan. Właśnie dlatego prawdziwe konsultacje są niewygodne. Gdyby były jedynie potwierdzeniem tezy, można by je zastąpić ankietą z trzema pytaniami i tortem na spotkaniu zespołu. Prawdziwa rozmowa otwiera możliwość korekty; bez niej staje się teatrykiem partycypacji.

Przykładem dojrzałego podejścia jest działanie Tracey, która przy wdrażaniu kolejnego systemu wykorzystwała AI do zaprojektowania grywalizacji – „Wyzwania Produktowności”. Pracownicy podzieleni na grupy sami wymyślali rozwiązania potencjalnych problemów z nowym narzędziem, zdobywali drobne nagrody i czuli realny wpływ na decyzję. Technologia nie zastąpiła tu uczestnictwa, lecz pomogła je zaprojektować. Nie udawała pracowników, ale stworzyła warunki, w których mogli oni wnieść własną wiedzę. To subtelna, lecz decydująca różnica: AI nie była „głosem ludzi”, lecz infrastrukturą procesu, w którym ludzie mogli przemówić. W tym modelu maszyna nie kolonizuje relacji społecznych, lecz wspiera ich organizację. To właściwy kierunek: AI jako fundament mądrzejszej rozmowy, a nie jej zamiennik.

Z perspektywy nauk społecznych można uznać, że siódmy krok Einhorna broni decyzji przed technokratycznym skrzywieniem. Technokracja nie polega na korzystaniu z wiedzy eksperckiej – ta jest niezbędna. Zaczyna się ona wtedy, gdy ekspert lub system uznaje świat społeczny za przeszkodę

we wdrożeniu rozwiązania logicznego. Ludzie stają się wówczas „oporem”, „czynnikiem ryzyka” czy „zmienną miękką”. To język wygodny, lecz niebezpieczny; odbiera ludziom status współtwórców rzeczywistości i zamienia ich w tarcie operacyjne. Einhorn proponuje rozsądniejszą alternatywę: AI może wspierać rozumienie interesariuszy, ale to człowiek musi prowadzić proces uznania, rozmowy i odpowiedzialności.

Nie jest to romantyczna pochwała spotkań. Każdy, kto uczestniczył w źle prowadzonych konsultacjach, wie, że ludzie potrafią mówić chaotycznie i obok tematu. Jednak lekarstwem na złe słuchanie nie jest brak słuchania, lecz struktura: właściwe pytania, dobór osób, jasny cel oraz uczciwe określenie, co podlega zmianie, a co nie. Jeśli uczestnicy po konsultacjach nie wiedzą, co stało się z ich głosem, uznają, że byli jedynie dekoracją – i zwykle mają rację.

Siódmy krok prowadzi zatem do pytania o instytucjonalną dojrzałość. Organizacje niedojrzałe komunikują decyzje po fakcie i dziwią się oporowi. Przeciętne konsultują, lecz nie zmieniają niczego istotnego. Dojrzałe potrafią włączyć interesariuszy na etapie, w którym ich wiedza ma jeszcze wpływ na rozwiązanie. Wymaga to odwagi, gdyż otwiera proces na krytykę, ale krytyka przed decyzją jest tańsza niż katastrofa po wdrożeniu. Lepiej usłyszeć trudne pytanie w sali konferencyjnej, niż odczuć jego skutki w budżecie, sądzie czy masowym odejściu ludzi.

W tym miejscu warto wrócić do pojęcia „human edge”. Ludzka przewaga nie oznacza, że człowiek jest zawsze bardziej empatyczny od maszyny – bywa przeciwnie: AI potrafi wygenerować uprzejmniejszy komunikat niż niejeden dyrektor. Przewaga polega na tym, że tylko człowiek może zostać moralnie zobowiązany przez odpowiedź drugiego człowieka. AI rozpozna wzorzec emocji; człowiek poczuje ciężar czyjejś obawy. AI przygotuje argument; człowiek zrezygnuje z części planu, rozumiejąc jego ukryty koszt. AI może symulować sprzeciw, ale tylko człowiek może uznać ten sprzeciw za uprawniony. To jest fundamentalna różnica między inteligencją instrumentalną a odpowiedzialnością relacyjną.

Przekonanie i odpowiedzialność Głównego Decydenta

Pierwsza pyta: jak osiągnąć cel? Druga pyta: czy cel, droga i koszt są do obrony wobec tych, których dotyczą? Pierwsza jest domeną optymalizacji. Druga to domena etyki, polityki, zarządzania i dojrzałego życia społecznego. AI świetnie wspiera tę pierwszą. Może wspierać drugą, ale nie może jej zastąpić. Włączanie interesariuszy przygotowuje grunt pod ostatni krok: przekonanie, czyli gotowość do działania. Dopiero po zdefiniowaniu problemu, rozpoznaniu motywacji, ustaleniu kontekstu, przeprowadzeniu badań, wykonaniu analizy, zakwestionowaniu błędów poznawczych i wysłuchaniu interesariuszy człowiek może zapytać: czy wiem wystarczająco dużo, by działać?

Nie chodzi o pewność absolutną. Nie o usunięcie każdego ryzyka czy sprawienie, by wszyscy byli zadowoleni. Dojrzała decyzja nie rodzi się z braku niepewności, lecz z przekonania, że niepewność została potraktowana poważnie. Conviction: decyzja ostateczna, premortem i odwaga działania bez gwarancji. Ostatni krok procesu decyzyjnego jest najbardziej ludzki, ponieważ nie da się go w całości zamienić w obliczenie, procedurę czy symulację. Można przygotować dane, rozpoznać motywacje, opisać kontekst, przeprowadzić badania i analizy, zakwestionować błędy poznawcze oraz wysłuchać interesariuszy. Można poprosić AI o warianty, porównania, scenariusze, ryzyka, kontrargumenty, premortem i rekomendacje. Można nawet przez chwilę ulec złudzeniu, że decyzja sama wyłoni się z tej pracy jak wynik równania. Ale tak się nie dzieje. W pewnym momencie człowiek musi zrobić coś, czego maszyna za niego nie zrobi: uznać, że wie dostatecznie dużo, choć nie wie wszystkiego, i wziąć odpowiedzialność za działanie. W źródle ósmy krok nazwano „Podejmowaniem Ostatecznej Decyzji”, „Coming to Conviction” albo „Take the Plunge”. Chodzi o osiągnięcie przekonania i gotowości do przekucia przemyśleń w realne działanie.

Einhorn nazywa człowieka „Głównym Decydentem”, ponieważ sztuczna inteligencja może rekomendować, ale to człowiek będzie żył z konsekwencjami wyboru. AI potrafi dostosowywać sugestie do priorytetów, które wcześniej otrzymała, lecz nie potrafi za człowieka osądzić, czy dany wybór jest naprawdę właściwy w świetle wartości, obaw, relacji, ambicji i kosztów, które nie mieszczą się w czystym rachunku optymalizacyjnym. To rozróżnienie między rekomendacją a przekonaniem jest kluczowe. Rekomendacja może być zewnętrzna; przekonanie musi zostać uwewnętrznione. Rekomendacja mówi: „wariant A wydaje się najlepszy według przyjętych kryteriów”. Przekonanie mówi: „rozumiem kryteria, rozumiem ryzyka, rozumiem konsekwencje i jestem gotów działać”.

Pierwsze może wygenerować system. Drugie musi wypowiedzieć człowiek, choćby tylko wobec siebie. Bez tego decyzja pozostaje zawieszona w pół drogi: niby wybrana, ale jeszcze nie przyjęta; niby racjonalna, ale jeszcze nie oswojona; niby gotowa, ale wciąż odkładana. Einhorn podkreśla, że przekonanie nie oznacza stuprocentowej pewności. To istotne, gdyż kultura zarządzania i polityki często udaje, że dobra decyzja to taka wolna od niepewności. Tymczasem przyszłość nie jest dokumentem do podpisu. Nie da się jej zawczasu przeczytać, poprawić czerwonym długopisem i odesłać do sekretariatu świata. Każda poważna decyzja zawiera ryzyko, a im bardziej jest doniosła, tym mniej przypomina wybór między oczywistym dobrem a oczywistym złem. Zwykle jest to wybór między dobrami częściowymi, kosztami rozłożonymi nierówno, szansami obciążonymi niepewnością i wartościami, które trzeba uporządkować. Przekonanie pojawia się więc nie wtedy, gdy ryzyko znika, ale gdy człowiek rozumie je wystarczająco dobrze, by nie być jego zakładnikiem. To moment, w którym wcześniejsza praca układa się w logiczną całość: problem zdefiniowany, motywacja rozpoznana, kontekst opisany, badania dostarczyły materiału, analiza wydobyła

znaczenie, uprzedzenia zakwestionowane, a interesariusze wnieśli swoje perspektywy. Nadal nie ma gwarancji.

Jest jednak jasność. A jasność w świecie złożonych decyzji jest często więcej warta niż złudzenie pewności. Właśnie tu pojawia się niebezpieczeństwo paraliżu decyzyjnego. Notatka źródłowa trafnie wskazuje, że jeśli mimo wyłonienia się faworyta człowiek bez końca przegląda te same dane albo prosi AI o kolejne symulacje, grozi mu decyzja przez bierność. Zwłoka również jest decyzją, choć zwykle nie ma odwagi przedstawić się pod własnym imieniem. Można stracić ofertę pracy, okno rynkowe, zaufanie zespołu, moment polityczny czy szansę naprawy relacji nie dlatego, że wybrało się źle, ale dlatego, że nie wybrało się wcale. Paraliż decyzyjny jest szczególnie zdradliwy w epoce AI, ponieważ potrafi nosić maskę profesjonalizmu. Dawniej człowiek zwlekał, mówiąc: „muszę to jeszcze przemyśleć”. Dziś może powiedzieć: „poproszę jeszcze model o dodatkową analizę scenariuszową, rozszerzoną weryfikację ryzyk i porównanie wrażliwości parametrów”. Brzmi poważnie. Czasem jest poważne. Ale czasem jest tylko nowoczesną formą chowania się pod biurkiem.

Nie każde dodatkowe badanie zwiększa odpowiedzialność; czasem ją odracza. Dlatego dojrzały proces decyzyjny musi posiadać nie tylko mechanizmy gromadzenia informacji, ale i moment ich zakończenia. Człowiek powinien wiedzieć, kiedy kolejne dane przynoszą nowe rozumienie, a kiedy jedynie powtarzają znane napięcia. Powinien umieć odróżnić roztropność od lęku. Roztropność pyta: czego jeszcze naprawdę nie wiemy, a co może zmienić decyzję? Lęk pyta: co jeszcze mogę sprawdzić, żeby nie musieć wybierać? Roztropność szuka brakującego elementu; lęk produkuje nieskończony aneks. W tym miejscu szczególnego znaczenia nabiera technika premortem, którą Einhorn proponuje jako narzędzie wzmacniania przekonania. Premortem, czyli analiza „przedśmiertna”, polega na wyobrażeniu sobie przyszłości, w której nasza decyzja okazała się katastrofą, i cofnięciu się myślowo do pytania: co mogło do tego doprowadzić? To przeciwieństwo zwykłego postmortem. Premortem pozwala przeprowadzić sekcję zwłok decyzji, zanim ta umrze naprawdę. Brzmi makabrycznie, ale bywa aktem troski – lepiej rozciąć złudzenie na stole analitycznym niż zbierać szczątki projektu po wdrożeniu.

Wskazane są trzy podejścia do premortem z pomocą AI. Pierwsze to analiza SWOT: mocne i słabe strony, szanse i zagrożenia. Drugie to symulacje scenariuszy, w których AI odgrywa rolę konkurenta, klienta, użytkownika czy regulatora. Trzecie to przewidywanie łańcuchów porażek, czyli efektu domina: mały błąd w oprogramowaniu prowadzi do złych recenzji, co odstrasza klientów i pogarsza przepływy pieniężne, frustrując w końcu inwestorów. Siła premortem polega na odwróceniu psychologii sukcesu. Zamiast pytać: „dlaczego to się uda?”, pyta: „jak to może się nie udać?”. To pytanie jest bezcenne, ponieważ człowiek zakochany we własnym planie widzi przede

wszystkim ścieżkę zwycięstwa: konferencję prasową, zadowolony zespół i wykresy wzrostu. Premortem nakazuje zdjąć z planu świecący lakier i obejrzeć zawiasy oraz miejsca, w których konstrukcja może pęknąć. Jest metodą poznawczego pesymizmu w służbie praktycznej nadziei.

AI nadaje się do premortem szczególnie dobrze, gdyż szybko generuje alternatywne scenariusze porażki. Może zagrać sprytnego konkurenta wykorzystującego nasze słabości, niezadowolonego klienta irytującego się ceną lub interfejsem, regulatora pytającego o zgodność z przepisami czy pracownika, który widzi, że plan jest niewykonalny przy istniejących zasobach. Może wcielić się w rolę dziennikarza szukającego luki reputacyjnej lub użytkownika pominiętego przez standardowy proces projektowania. Trzeba jednak pamiętać, że premortem z AI nie jest wrózeniem z algorytmicznych fusów ani prorocstwem, lecz ćwiczeniem wyobraźni krytycznej. Jego celem nie jest przewidzenie dokładnej przyszłości, lecz ujawnienie kruchych punktów decyzji. Jeżeli AI wygeneruje dziesięć scenariuszy porażki, nie oznacza to, że wszystkie są jednakowo prawdopodobne.

Decydent musi je ocenić i zestawić z wiedzą ludzi oraz realiami. AI pomaga mnożyć hipotezy; człowiek musi oddzielić te poważne od fantazji, sygnał od szumu, ryzyko strategiczne od literackiej katastrofy. Dobrze przeprowadzony premortem nie osłabia decyzji – przeciwnie, hartuje ją. To jedno z najważniejszych przesłań notatki źródłowej: identyfikacja ryzyka pozwala przygotować bezpieczniki i plany awaryjne, a więc wzmacnia gotowość do działania. Człowiek, który widział możliwe porażki i ma na nie odpowiedzi, nie działa z naiwnym entuzjazmem, lecz z trzeźwą odwagą. Różnica jest zasadnicza. Naiwny entuzjazm mówi: „na pewno się uda”. Trzeźwa odwaga mówi: „wiem, co może pójść źle, przygotowałem zabezpieczenia i mimo to uważam, że warto działać”.

W zarządzaniu, prawie czy polityce publicznej ta różnica ma ogromne znaczenie. Zbyt wiele decyzji upada nie dlatego, że ich autorzy nie mieli wizji, ale dlatego, że mieli wyłącznie wizję – bez mapy ryzyk, planu awaryjnego i pokory wobec rzeczywistości. Wizja bez premortem bywa piękną katastrofą w stanie embrionalnym. Z drugiej strony premortem bez decyzji staje się rytuałem straszenia. Celem nie jest więc ani bezkrytyczny optymizm, ani paraliżujący pesymizm, lecz przekonanie po próbie ognia. Decyzja ostateczna wymaga pogodzenia dwóch pozornie sprzecznych postaw: sceptycyzmu i sprawczości. Sceptycyzm bez sprawczości prowadzi do cynicznego bezruchu; sprawczość bez sceptycyzmu do brawury.

Einhorn proponuje drogę trzecią: używaj AI, aby myśleć szerzej, szybciej i krytyczniej, ale nie traktuj jej jako schronu przed odpowiedzialnością. Poproś maszynę o najtrudniejsze kontrargumenty, ale nie chowaj się za nimi w nieskończoność. Poproś o scenariusze porażki, lecz nie uznawaj samego istnienia ryzyka za powód rezygnacji. Poproś o rekomendację, ale nie myl jej z

decyzją. Ma to szczególne znaczenie w instytucjach wykorzystujących AI do przygotowywania decyzji administracyjnych czy prawnych, gdzie pokusa odpowiedzialności rozproszonej jest ogromna. Decyzja przygotowana przez system, zaakceptowana przez analityka i podpisana przez przełożonego może stać się decyzją bez autora. W razie sukcesu autorów znajdzie się wielu; w razie porażki pozostanie mgła procesu. Koncepcja „Głównego Decydenta” jest odpowiedzią na tę mgłę: ktoś musi wiedzieć, że to on bierze odpowiedzialność za ostateczny wybór.

Nie oznacza to autorytaryzmu. Przeciwnie – im poważniejsza decyzja, tym bardziej powinna korzystać z analiz i narzędzi AI. Ale po zakończeniu procesu nie wolno udawać, że wynik sam wypłynął z procedury. Procedura może uzasadniać decyzję, ale jej nie uosabia. Modele wspierają wybór, lecz nie ponoszą odpowiedzialności moralnej. Dane informują, ale nie zastępują osądu. Organizacja dojrzała to taka, która potrafi połączyć partycypację i analizę z jasno przypisaną odpowiedzialnością.

W wymiarze osobistym „conviction” również nie oznacza ślepej pewności siebie, lecz zgodę na działanie po uczciwej pracy myślowej. Człowiek rozważający zmianę pracy czy założenie firmy nigdy nie uzyska pełnej gwarancji. Może jednak zyskać świadomość, że nie działa z impulsu, nie ucieka przed lękiem i nie ignoruje kontekstu ani możliwych porażek. To wystarczy, by decyzja była poważna. Nie wystarczy, by była nieomylna – ale nieomyślność nie jest kategorią ludzkiego życia; jest domeną bajek i złych prezentacji sprzedażowych.

Trzeba powiedzieć rzecz niewygodną: część ludzi nie chce AI jako narzędzia lepszego myślenia, lecz jako usprawiedliwienia gorszej odpowiedzialności. Chcą, aby system wskazał kogo zwolnić lub gdzie inwestować, by wybór wyglądał na obiektywny. To nie jest technologia jako partner poznawczy, lecz technologia jako pralnia sumienia. Można tam wrzucić brudną decyzję, a wyjąć pachnący raport o optymalizacji. Tylko rzeczywistość po pewnym czasie wystawia rachunek.

Koncepcja Einhorn sprzeciwia się takiemu użyciu AI. Jej ośmioetapowy proces, Arkusze Geparda i metoda AREA tworzą etykę decydowania z technologią. Definicja problemu chroni przed złe ustawionym pytaniem; motywacja przed optymalizacją cudzych celów; kontekst przed praktyczną głupotą; badania przed ignorancją; analiza przed chaosem danych; praca nad uprzedzeniami przed potwierdzaniem błędów głosem maszyny; interesariusze przed społeczną ślepotą, a conviction przed paraliżem i ucieczką od odpowiedzialności. Każdy z tych kroków odpowiada na inną patologię współczesnego myślenia: chaos celów, pustkę sensu, technokratyczne oderwanie od świata czy pychę poznawczą. Dlatego książka Einhorn to podręcznik odzyskiwania decyzyjnej dojrzałości w epoce, która daje nam narzędzia szybsze niż nasze sumienie.

Najpiękniejszą metaforą pozostają Arkusze Geparda. Gepard nie wygrywa tylko dlatego, że jest szybki, lecz dlatego, że potrafi hamować i skręcać. W epoce AI to zdanie powinno być maksymą cywilizacyjną. Szybkość bez zdolności hamowania jest katastrofą o wysokiej wydajności. Organizacja, która generuje dokumenty, ale nie stawia pytań, będzie produkować więcej błędów w krótszym czasie. Lider uzyskujący rekomendacje, lecz nieznający własnej motywacji, będzie skuteczniej realizował cudze schematy. Instytucja wdrażająca systemy bez słuchania interesariuszy będzie cyfryzować alienację. Człowiek, który potrafi pytać AI, ale nie potrafi zdecydować, pozostanie użytkownikiem narzędzia, a nie autorem życia.

Odwaga decyzji i protokół sprawczości

Decyzja ostateczna wymaga aktu odwagi. Nie tej teatralnej, karmiącej się wielkimi deklaracjami i zdjęciami na tle slajdów, lecz odwagi cichszej: zdolności powiedzenia „wybieram” po rzetelnym sprawdzeniu alternatyw. Powiedzenia „biorę odpowiedzialność” po wysłuchaniu wszystkich zastrzeżeń. To odwaga uznania: „wiem, że mogę się mylić”, a mimo to działanie w przekonaniu, że bierność byłaby błędem znacznie poważniejszym. W tym sensie conviction nie jest pychą, lecz pewnością wystarczającą. Nie absolutną, lecz operacyjną; nie dogmatyczną, lecz dojrzałą. Na końcu procesu człowiek nie powinien pytać AI: „co mam zrobić?”, lecz raczej: „czy pomogłaś mi zobaczyć to, czego sam nie widziałem?”.

Taka jest właściwa relacja. AI jako lustro poszerzające pole widzenia, a nie suweren decyzji. Jako sparingpartner, a nie kapłan wyroczni. Narzędzie premortem, a nie producent alibi. Przyspieszenie badań, a nie zastępstwo sensu. Pomoc w myśleniu, a nie zwolnienie z myślenia. Dlatego ludzka przewaga w epoce sztucznej inteligencji nie będzie polegała na odrzuceniu technologii – to byłaby nostalgia, nie strategia. Nie będzie też polegała na ślepym zachwycie nad automatyzacją. To byłby infantylizm z dostępem do API.

Ludzka przewaga będzie polegać na umiejętności łączenia mocy obliczeniowej z odpowiedzialnym osądem, szybkości z pauzą, rekomendacji z wartościami, symulacji z rozmową i sceptycyzmu z działaniem. Człowiek, który to potrafi, nie konkuruje z AI na liczbę przetworzonych dokumentów, lecz o coś znacznie ważniejszego: o zachowanie autorstwa własnych decyzji. W tym sensie „The Human Edge” nie jest pochwałą człowieka takiego, jaki jest, lecz wezwaniem do bycia lepszym niż własne poznawcze lenistwo. AI może obnażyć nasze słabości: niejasność celów, brak cierpliwości, skłonność do potwierdzania własnych tez, ucieczkę od rozmowy czy lęk przed decyzją. Może jednak pomóc je korygować, jeśli potraktujemy ją jako narzędzie dyscypliny, a nie maszynę do wygodnych odpowiedzi. Ostateczny wybór należy jednak do człowieka. I dobrze. Bo tam, gdzie zaczynają się

konsekwencje, kończy się mit wyroczni. Tam zostaje decydent: jego rozum, wartości, odwaga i jego podpis pod światem, który współtworzy.

Osiem unikalnych ludzkich kroków – pełna rekonstrukcja rdzenia metody – nie może pozostać jedynie tłem tego wyводу. To nie jest dekoracyjna lista ani luźny zbiór porad. To kręgosłup koncepcji Cheryl Strauss Einhorn.

Bez systematycznego wydobycia tych etapów tekst traci architekturę, gdyż to właśnie one wyznaczają granicę między używaniem AI jako narzędzia a oddaniem jej własnej sprawczości. W notatce źródłowej kroki te zostały wymienione jasno: definicja problemu, motywacja, kontekst, badania, analiza, zwalczanie uprzedzeń, włączanie interesariuszy oraz podejmowanie ostatecznej decyzji. Każdemu z nich towarzyszy „Pierwszy bajt wglądu” – pytanie prowadzące decydenta przez etap, którego maszyna nie może za niego przeżyć ani moralnie rozstrzygnąć. Te osiem kroków należy rozumieć jako antyautomatyczny protokół myślenia. Nie chodzi o to, by człowiek rywalizował z AI w szybkości przetwarzania informacji. Przegrałby z kretesem, i to bez godnej sceny finałowej.

Chodzi o wykonanie aktów rozpoznania, których AI nie posiada jako własnego doświadczenia: określenie problemu, rozpoznanie motywacji, osadzenie decyzji w kontekście, nadanie kierunku badaniom, interpretację danych, kontrolę błędów poznawczych (własnych i maszynowych), włączenie interesariuszy oraz wzięcie odpowiedzialności za ostateczny wybór. To jest pełny sens „human edge”: ludzka przewaga nie jest sentymentalnym dodatkiem do technologii, lecz warunkiem jej odpowiedzialnego użycia.

Definicja problemu i rola motywacji

Pierwszym krokiem jest definicja problemu. Pytanie prowadzące brzmi: „Gdybyś wiedział, czego naprawdę chcesz, czy potrafiłbyś o to poprosić lub o to zapytać?”. To pytanie uderza w samo sedno pracy z AI. Większość błędnych odpowiedzi nie wynika z błędu modelu, lecz z niejasności człowieka. Kto nie wie, czego szuka, otrzyma imponująco szybkie przybliżenie własnego chaosu. Definicja problemu wymaga zatem nie tylko nazwania tematu, ale precyzyjnego uchwycenia jego granic, przyczyn, pożądanego rezultatu, ograniczeń, zasobów oraz osób zaangażowanych w sytuację. Przykład Jamesa obrazuje to bardzo wyraźnie.

Gdy właściciel kawiarni poprosił AI o znalezienie „świetnej powierzchni komercyjnej do wynajęcia w hrabstwie”, otrzymał propozycje formalnie zgodne z zapytaniem, lecz praktycznie bezużyteczne: zbyt odległe i nieuwzględniające dojazdu, zespołu, konkurencji czy realiów biznesowych. Dopiero

gdy sam wykonał pracę definicyjną — wskazując lokalizację obecnej kawiarni, metraż, odległość od transportu publicznego, maksymalny czas dojazdu oraz konieczność unikania bezpośredniej konkurencji — AI zaczęła działać użytecznie. Definicja problemu jest zatem pierwszą barierą przeciwko błędowi pozornej kompetencji AI. Algorytm potrafi odpowiedzieć na źle postawione pytanie z taką samą elegancją, z jaką odpowiada na pytanie dobre. W tym właśnie tkwi ryzyko: zła definicja problemu nie krzyczy i nie zapala czerwonej lampki; często wygląda jak profesjonalny prompt, podczas gdy prowadzi w stronę błędnych danych, analiz i decyzji. To jak ustawienie nawigacji na niewłaściwy adres: samochód jedzie sprawnie, głos brzmi spokojnie, trasa jest zoptymalizowana, tylko cel jest błędny.

Sztuczna inteligencja nie zastępuje zatem pierwszego aktu człowieczeństwa w decyzji: umiejętności powiedzenia, o co naprawdę chodzi. Drugim krokiem jest motywacja. Pytanie brzmi: „Czy wiesz, dlaczego rozwiązujesz swój problem i jaka jest nadrzędna potrzeba znalezienia rozwiązania?”. O ile definicja problemu odpowiada na pytanie „co?”, o tyle motywacja wyjaśnia „po co?”. W tym punkcie decyzja przestaje być jedynie zadaniem technicznym, a staje się zadaniem aksjologicznym. Człowiek musi rozpoznać, jaka wartość, potrzeba, lęk, ambicja, zobowiązanie lub wizja przyszłości stoi za jego wyborem. AI może pomóc uporządkować odpowiedzi, ale nie jest w stanie posiadać ludzkiego „dlaczego”.

Einhorn proponuje w tym celu pięć pytań: co ma się wydarzyć po rozwiązaniu problemu, dlaczego jest to ważne właśnie teraz, na jakie kompromisy decydent jest gotów, a na jakie nie, jak wygląda sukces za sześć miesięcy lub rok oraz jakie obawy mogą powstrzymać przed działaniem. To pozornie proste pytania, lecz w praktyce bywają bardziej wymagające niż analiza danych. Łatwiej poprosić model o porównanie rynków, niż przyznać, że w rzeczywistości szukamy bezpieczeństwa. Łatwiej zamówić symulację finansową, niż powiedzieć, że nie chcemy więcej stresu. Łatwiej pytać o optymalizację kariery, niż nazwać lęk przed utratą czasu dla rodziny. Motywacja zmienia sposób odczytywania rekomendacji AI.

W przykładzie Jamesa wybór droższego lokalu przy ruchliwej ulicy mógłby wyglądać optymalnie w świetle danych o natężeniu ruchu. Jednak jeśli jego prawdziwym celem było przyspieszenie niezależności finansowej bez nadmiernego ryzyka, spokojniejsza i tańsza lokalizacja okazywała się bardziej zgodna z jego życiowym projektem. Podobnie Angelika, rozważając przeprowadzkę z Atlanty do Bostonu, nie pytała wyłącznie o atrakcyjność oferty pracy. Jej motywacja obejmowała awans i bezpieczeństwo finansowe, ale także zdrowie psychiczne i czas dla rodziny.

Kontekst, badania i analiza danych

Dopiero tak określona motywacja pozwala przekształcić AI z domniemanej wyroczni w partnera do myślenia. Trzecim krokiem jest kontekst, a pytanie, które mu towarzyszy, brzmi: „Czy znasz środowisko, w którym osadzona jest twoja decyzja?”. To etap, na którym człowiek musi spojrzeć na zewnątrz: na czas, miejsce, ludzi, ograniczenia, instytucje, kulturę organizacyjną i sytuację życiową. W notatce kontekst zostaje nazwany „ogólnym stanem zdrowia decyzji”. Bez niego ryzykujemy leczenie objawu przy jednoczesnym ignorowaniu choroby; możemy podjąć decyzję logiczną, lecz całkowicie nieadekwatną do realnego świata.

Einhorn wprowadza tu kategorię sytuacyjności, obejmującą lokalizację, ludzi, własność decyzji i etap życia. Lokalizacja określa, gdzie decyzja będzie działać; ludzie wskazują, kto odczuje jej skutki; własność decyzji definiuje, kto formalnie i faktycznie odpowiada za wybór. Etap życia z kolei przypomina, że żadna decyzja nie istnieje w próżni – jest zawsze wpisana w biografię decydenta i osób zaangażowanych. Ten sam poziom ryzyka będzie inaczej postrzegany przez młodego lidera budującego karierę, inaczej przez osobę planującą emeryturę, a jeszcze inaczej przez rodzica zależnego od stabilnego rytmu życia.

To właśnie w kontekście ujawnia się ślepotą AI. Model może dysponować wiedzą ogólną, ale nie zna naszego ekosystemu, dopóki go nie opiszemy. Nie wie, kto w organizacji posiada nieformalne weto, że zespół jest wyczerpany kolejną zmianą czy że lokalny rynek rządzi się zwyczajami niewidocznymi w statystykach. Nie rozumie, że decyzja formalnie racjonalna może zniszczyć kruche zaufanie. Sztuczna inteligencja staje się potężnym narzędziem dopiero wtedy, gdy człowiek dostarczy jej kontekst. Bez tego produkuje rozwiązania poprawne abstrakcyjnie i kulawe praktycznie.

Czwartym krokiem są badania. Pytanie prowadzące brzmi: „Jakie informacje pomogłyby ci zrozumieć twój problem?”. Choć brzmi ono technicznie, w rzeczywistości nie pyta o to, jakie informacje istnieją, lecz które z nich są niezbędne do zrozumienia sedna problemu. W epoce AI dostęp do danych przestał być głównym ograniczeniem; wyzwaniem stały się kierunek poszukiwań, selekcja, filtracja, weryfikacja oraz umiejętność zakończenia procesu.

Źródło opisuje badania poprzez metaforę platform streamingowych: przeglądanie, zawężanie i głębokie zanurzenie. Najpierw rozpoznajemy krajobraz możliwości, następnie filtrujemy go według kryteriów, by wreszcie zgłębić najbardziej dopasowane ścieżki. Einhorn wyróżnia tu podejście progresywne – przydatne, gdy nie wiemy, od czego zacząć – oraz ustrukturyzowaną ścieżkę badawczą, w której pytania wynikają z wcześniej ustalonego kontekstu. Badania wspierane przez AI wymagają jednak żelaznej dyscypliny. Model może halucynować, mieszać fakty z fikcją, generować pozorne źródła lub produkować nadmiar materiału.

Dlatego człowiek musi żądać źródeł i rygorystycznie je weryfikować, wychodząc poza własną badawczą strefę komfortu. Kto ufa wyłącznie liczbom, powinien szukać narracji i ludzkich doświadczeń; kto opiera się na historiach – powinien sprawdzić twarde dane. Badania są więc nie tylko gromadzeniem materiału, lecz ćwiczeniem przeciwko własnej jednostronności.

Piątym krokiem jest analiza, której pytanie brzmi: „Jakie znaczenie i sens potrafisz wydobyć ze swoich danych?”. To moment, który odróżnia posiadanie informacji od ich rozumienia. AI potrafi przetwarzać dane, identyfikować wzorce, porównywać warianty i prognozować trendy, jednak to człowiek musi zdecydować, co te wzorce oznaczają w świetle celu, wartości i kontekstu. Analiza nie jest zatem mechaniką obliczeń, lecz sztuką interpretacji.

Einhorn zauważa, że nie każda analiza odpowiada na to samo pytanie. Możemy prowadzić analizę opisową, diagnostyczną, predykcyjną, finansową, sentymentu, przestrzenną czy porównawczą. Sam wybór typu analizy jest już decyzją. Kto ogranicza się do analizy finansowej, dostrzeże koszt i rentowność, ale może przeoczyć reputację, morale zespołu lub skutki społeczne. Kto wybiera tylko analizę sentymentu, może zignorować twarde limity budżetowe. Mądrość polega na zestawianiu różnych perspektyw.

Sz szczególnie istotne jest uwzględnianie danych zaprzeczających – to jeden z kluczowych elementów całej koncepcji. Informacje przeczące naszym preferencjom nie są przeszkodą, lecz materiałem ochronnym. Mogą one uświadomić nam, że zakochaliśmy się w błędnym rozwiązaniu, pominęliśmy istotny czynnik lub pomyliliśmy motywację z ambicją. AI może pomóc wykrywać takie napięcia, ale tylko wtedy, gdy człowiek nie traktuje jej jako służby propagandowej własnej tezy. Model zapytany tendencyjnie często odpowie w ten sam sposób. Tu właśnie zaczyna się odpowiedzialność analityczna. Szóstym krokiem jest zwalczanie uprzedzeń, czyli biases.

Pułapki poznawcze i rola interesariuszy

Jego pytanie brzmi: „Czy wiesz, w których miejscach ty i AI opieracie się na założeniach i wydajecie osądy?”. To jeden z najbardziej nowoczesnych i niepokojących elementów metody, gdyż dowodzi, że błąd nie leży wyłącznie po stronie człowieka ani maszyny. Człowiek operuje skrótami poznawczymi; AI zaś obarczone jest błędami danych, algorytmów i wdrożenia.

Najgroźniejszy jest moment, w którym te dwa porządki zaczynają się wzajemnie wzmacniać. Autor wymienia klasyczne ludzkie pułapki: efekt potwierdzenia, który każe szukać dowodów na poparcie naszych założeń; zakotwiczenie, przywiązujące nas do pierwszej uzyskanej informacji; błąd dostępności, sprawiający, że przeceniamy to, co świeże lub dramatyczne, oraz błąd automatyzacji,

prowadzący do bezkrytycznego zaufania technologii. Po stronie AI występują błędy danych i algorytmów, a także błędy wdrożenia, gdy narzędzie trafia do środowiska innego niż to, w którym było trenowane. Najbardziej niebezpieczna sytuacja powstaje wtedy, gdy człowiek i AI tworzą sojusz złudzeń. Menedżer może faworyzować kandydata, a system oparty na stronniczych danych historycznych tę preferencję potwierdzi. Prawnik może zaufać wygenerowanej sygnaturze sprawy, która w rzeczywistości nie istnieje. Kuratorium może użyć systemu wskazującego uczniów „zagrożonych”, nie dostrzegając, że model obciąża nieproporcjonalnie określone szkoły lub grupy. Wówczas AI nie rozwiązuje problemu, lecz nadaje staremu błędowi nowoczesną pieczęć obiektywności.

Einhorn proponuje cztery strategie ochronne: konfrontowanie oczekiwań, testowanie narzędzia AI, zadawanie pytań „dlaczego?” oraz weryfikację krzyżową. To nie techniczne sztuczki, lecz zasady higieny poznawczej. Zanim przyjmimy odpowiedź AI, powinniśmy wiedzieć, czego się spodziewaliśmy i sprawdzić, czy zmiana parametrów daje podobne wyniki. Należy wymuszać uzasadnienia i porównywać kluczowe informacje z niezależnymi źródłami. W tym kontekście sceptycyzm nie jest przeszkodą w pracy z AI – jest warunkiem profesjonalizmu.

Siódmym krokiem jest włączanie interesariuszy, a towarzyszące mu pytanie brzmi: „Czy znasz wpływ swojej decyzji na interesariuszy?”. Ten etap wyprowadza proces decyzyjny z głowy menedżera i ekranu komputera do świata relacji. Decyzje rzadko zapadają w próżni; dotyczą ludzi, którzy będą je wdrażać, znosić, popierać lub kontestować. Pominięcie interesariuszy może zniszczyć nawet najbardziej logiczny plan, ponieważ wiedza formalna nie zastępuje wiedzy praktycznej. Należy tu wyraźnie odróżnić włączanie interesariuszy od komunikowania im decyzji po fakcie – to drugie jest często jedynie elegancką formą narzucania woli. Włączanie oznacza aktywne słuchanie i zbieranie opinii przed podjęciem decyzji. AI może pomóc wskazać, kogo uwzględnić, przewidzieć obiekcje czy przygotować strategie komunikacyjne, ale nie potrafi usiąść z ludźmi przy stole, odczytać mowy ciała, wychwycić napięcie ani budować zaufania w czasie rzeczywistym.

Przykład Tracey obrazuje koszt pominięcia tego kroku. Dyrektorka wdrożyła system zarządzania zapasami po konsultacjach z zespołem badawczym, lecz bez realnego udziału pracowników magazynu i sprzedaży. System okazał się niekompatybilny z potrzebami użytkowników; nie była to awaria technologii, lecz awaria słuchania. Z kolei historia Solomona pokazuje, że choć AI może pomóc przygotować lepszą rozmowę z partnerką czy mentorem, nie jest w stanie przeżyć tej rozmowy za człowieka.

Ostateczna decyzja i ludzka odpowiedzialność

Ósmym krokiem jest podejmowanie ostatecznej decyzji, czyli conviction. Kluczowe pytanie brzmi tutaj: „Czy jesteś gotowy do podjęcia działania?”. To finał procesu, choć nie w sensie osiągnięcia absolutnej pewności. Einhorn wyraźnie podkreśla, że ostateczna decyzja nie oznacza stuprocentowej wiedzy o przyszłości, lecz wystarczającą jasność, by móc działać. To moment, w którym wcześniejsza praca – definicja, motywacja, kontekst, badania, analiza, walka z uprzedzeniami i rozmowy z interesariuszami – składa się w spójną całość.

Conviction nabiera szczególnego znaczenia w epoce AI, ponieważ narzędzia generatywne mogą paradoksalnie podsycać paraliż decyzyjny. Zawsze można poprosić o kolejną analizę, dodatkowy wariant, nową symulację czy model ryzyka. Należy jednak pamiętać, że zwłoka również jest decyzją – często decyzją o oddaniu szansy, utracie momentu i schowaniu odpowiedzialności pod stertą raportów. Ostateczny wybór wymaga aktu przejścia: nie lekkomyślnego skoku, lecz świadomego kroku po wykonaniu rzetelnej pracy poznawczej.

W tym miejscu pojawia się premortem, czyli analiza przedśmiertna. Einhorn proponuje, by przed ostatecznym wyborem wyobrazić sobie przyszłość, w której decyzja okazała się katastrofą, a następnie cofnąć się myślowo i zapytać, co mogło do tego doprowadzić. AI może wspierać ten proces poprzez analizę SWOT, symulacje scenariuszy i przewidywanie łańcuchów porażek; może odgrywać rolę konkurenta, klienta, regulatora albo sfrustrowanego użytkownika. Celem nie jest jednak osłabienie decyzji, lecz jej zahartowanie. Kto widzi ryzyka i przygotował bezpieczniki, działa dojrzej niż ten, kto idzie naprzód z transparentem entuzjazmu i pustą gaśnicą.

W pełnym ujęciu osiem kroków Einhorn tworzy logiczną sekwencję ludzkiej odpowiedzialności. Definicja problemu pyta, czego naprawdę chcemy. Motywacja – dlaczego to ma znaczenie. Kontekst – w jakim świecie decyzja będzie działać. Badania – jakiej wiedzy potrzebujemy. Analiza – co ta wiedza znaczy. Biases – gdzie człowiek i maszyna mogą się mylić. Interesariusze – kto będzie żył ze skutkami decyzji. Conviction zaś pyta, czy jesteśmy gotowi działać mimo niepewności.

Taki porządek pokazuje, że AI nie odbiera człowiekowi sprawczości, jeśli ten sam jej nie odda. Może natomiast bezlitośnie ujawnić obszary, w których byliśmy dotąd nieprecyzyjni, nieuczciwi wobec własnych motywacji, ślepi na kontekst, leniwi badawczo, powierzchowni analitycznie, podatni na uprzedzenia, głusi na interesariuszy albo tchórzliwi wobec działania. Brzmi to surowo, ale jest wyzwajające, bo jeśli AI obnaża słabości procesu, może też pomóc je naprawić.

Dlatego osiem unikalnych ludzkich kroków należy potraktować jako praktyczną konstytucję decyzji w epoce AI. Każdy z nich wyznacza granicę, której maszyna nie powinna przekraczać jako rzekomy

decydent, a jednocześnie wskazuje przestrzeń, w której AI jest pożyteczna jako narzędzie: może zadawać pytania, porządkować odpowiedzi, wykrywać luki, proponować źródła, symulować obiekcje, testować scenariusze i wspierać refleksję. Ale człowiek musi prowadzić. Nie dlatego, że jest zawsze mądrzejszy, lecz dlatego, że tylko on jest odpowiedzialny.

Szybkość bez zdolności hamowania jest jedynie katastrofą o wysokiej wydajności. W świecie, który ceni dynamikę, prawdziwym luksusem i najwyższą kompetencją staje się odwaga postawienia znaku stop przed ostatecznym wyborem. Czy potrafimy zatem używać AI jako narzędzia do lepszego myślenia, czy jedynie jako pralni sumienia dla decyzji, których nie mamy odwagi podpisać własnym imieniem?

Inspiracje

[1]



"Human Edge" Cheryl Strauss Einhorn

2026 | Cornell Publishing | ISBN: 978-1119931313

Cheryl Strauss Einhorn jest wielokrotnie nagradzaną dziennikarką śledczą, edukatorką oraz założycielką i prezeską Decisive, firmy zajmującej się naukami decyzyjnymi. Jest powszechnie uznawana za twórczynię metody AREA, opartej na dowodach metody, która pomaga jednostkom i zespołom radzić sobie ze złożonymi problemami poprzez kontrolę błędów poznawczych i poprawę osądu. Einhorn spędziła dwie dekady jako dziennikarka śledcza, pisząc dla ważnych czasopism, takich jak „The New York Times”, „Barron's” i „Foreign Policy”. Obecnie jest edukatorką na Uniwersytecie Cornella, a wcześniej wykładała w Columbia Business School. Jej praca koncentruje się na nauce o podejmowaniu decyzji, psychologii podejmowania decyzji oraz roli ludzkiego osądu w erze coraz większego wpływu sztucznej inteligencji. Jest autorką kilku książek, w tym „Problem Solved”, „Investing in Financial Research” i „The Human Edge”.

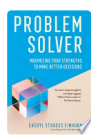
Cheryl Strauss Einhorn is an award-winning investigative journalist, educator, and the founder and CEO of Decisive, a decision-sciences company. She is widely recognized for creating the AREA Method, an evidence-based framework designed to help individuals and teams navigate complex problems by controlling for cognitive bias and improving judgment. Einhorn spent two decades as an investigative journalist, contributing to major publications such as The New York Times, Barron's, and Foreign Policy. She currently serves as an educator at Cornell University and has previously taught at Columbia Business School. Her work focuses on decision science, the psychology of decision-making, and the role of human judgment in an era increasingly influenced by artificial intelligence. She is the author of several books, including Problem Solved, Investing in Financial Research, and The Human Edge.

Cornell University

Literatura uzupełniająca

Książki

Autorzy przywoływani w artykule:



[1] "Problem Solver"

Cheryl Strauss Einhorn

2023 | Cornell University Press | ISBN: 9781501768026



[2] "Investing in Financial Research"

Cheryl Strauss Einhorn

2019 | Cornell Publishing | ISBN: 9781501730955



[3] "Problem Solved"

Cheryl Strauss Einhorn

2017 | Red Wheel/Weiser | ISBN: 9781632659170

Literatura pokrewna (Google Scholar):

[4] "Language in Mind Psycholinguistics 2nd edition Julie Sedivy"

[5] "How We Think"

Marius Ostrowski

[6] "Wrongs Harms and Compensation Adam Slavny"

[7] "Imperium tekstu"

[8] "Agentic Mesh The GenAI-Powered Agent"

Eric Broda

Artykuły naukowe

Artykuły pokrewne (Google Scholar):

[1] "How Do You Decide?"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.7591/cornell/9781501768002.003.0002

[Google Scholar](#)

[2] "Using PSPs to Bolster Strengths"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.7591/cornell/9781501768002.003.0013

[Google Scholar](#)

[3] "The Relationship between PSPs and Life Outlook"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.7591/cornell/9781501768002.003.0016

[Google Scholar](#)

[4] "2. Lexicon, Situationality, and Community"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.1515/9781501768026-005

[Google Scholar](#)

[5] "PSPs and Cognitive Biases"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.7591/cornell/9781501768002.003.0011

[Google Scholar](#)

[6] "How PSPs Color Our Relationship with Data"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.7591/cornell/9781501768002.003.0014

[Google Scholar](#)

[7] "The AREA method: how to make rational decisions"

Cheryl Strauss Einhorn

2021 | The Business & management collection. | DOI: 10.69645/cpsv3438

[Google Scholar](#)

[8] "Cast of Problem Solvers (in Order of Appearance)"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.7591/cornell/9781501768002.002.0014

[Google Scholar](#)

[9] "Hunt Like the Cheetah"

Cheryl Strauss Einhorn

2023 | Cornell University Press eBooks | DOI: 10.7591/cornell/9781501768002.003.0009

[Google Scholar](#)

