

Uniwersytet Papieski Jana Pawła II w Krakowie

Wydział Filozoficzny

Katedra Historii i Filozofii Nauki



Anna Sarosiek

Biosemiotyczna koncepcja inteligencji i jej zastosowanie w tworzeniu inteligentnych maszyn

Rozprawa doktorska napisana pod kierunkiem
dr hab. Pawła Polaka, prof. UPJPII

Kraków 2020

Opis bibliograficzny pracy

Anna Sarosiek, *Biosemiotyczna koncepcja inteligencji i jej zastosowanie w tworzeniu inteligentnych maszyn*, praca doktorska napisana pod kierunkiem doktora habilitowanego Pawła Polaka, Kraków, WF UPJPII, rok 2020, stron 2010.

Abstrakt

W pracy przedstawiono zagadnienie kształtowania się inteligencji sztucznych systemów w oparciu o teorie biosemiotyczne. Analiza filozoficznych założeń sztucznej inteligencji obnaża fundamentalny problem, który sięga najstarszych koncepcji myślenia o działaniu umysłu i maszyny. Problem ten dotyczy możliwości stworzenia w sztucznym systemie semantyki analogicznej do ludzkiej. Z perspektywy biosemiotyki, znaczenie jest jedną z najważniejszych kategorii kształtujących działanie inteligentnych istot. W tej perspektywie inteligentne działanie jest ucieleśnione, więc kategorie percepcji pozostają połączone z anatomiczną budową podmiotu a jego racjonalność jest związana ze światem.

Systemy sztucznej inteligencji zazwyczaj nie uwzględniają semiotycznej perspektywy, zakładając, że świat organizmu i maszyny mogą być identyczne i nadając maszynie semantykę stworzoną przez człowieka. Biosemiotyka pokazuje, że inteligencja jest kształtowana jako cecha konkretnego gatunku. W tej perspektywie inteligentne zachowanie jest wynikiem funkcjonowania w systemie znaków oraz odniesienia się do gatunkowej historii i doświadczeń podmiotu.

Celem pracy jest wykazanie, że sztucznie skonstruowany system może realizować semiotyczne działania. Bazując na koncepcji Jakoba von Uexküll'a uzasadnione zostało, pod jakimi warunkami możliwe jest stworzenie sztucznej inteligencji wyposażonej w fenomenalny, subiektywny świat. Praca przedstawia podstawy kształtowania się inteligentnych zachowań sztucznych układów z perspektywy biosemiotyki oraz rozważa możliwości powstania inteligencji maszynowej jako nowej formy inteligencji z perspektywy biosemiotycznej.

Zastosowano metodę analizy porównawczej.

Wykorzystano 283 pozycje bibliograficzne.

Słowa kluczowe:

- a) imienne: Jakob von Uexküll, Steven Harnad, Tom Ziemke, Claus Emmeche, Thomas Sebeok, Jesper Hoffmeyer, John Searle, Hubert Dreyfus, Rodney Brooks
- b) rzeczowe: biosemiotyka, sztuczna inteligencja, radykalny konstruktywizm, autonomia, intencjonalność, adaptacja, ewolucja, podmiot poznawczy, świat fenomenalny, Umwelt, kręgi funkcjonalne, agent, sieci neuronowe, roboty, algorytmy, ALife

Summary

Biosemiotik concept of intelligence and its application in constructing intelligent machines

The philosophical study of artificial intelligence reveals a fundamental problem following the oldest concepts of thinking about minds and machines. The problem affects every possibility of creating artificial semantic similar to human semantic. In the biosemiotics perspective, the meaning is the most important categories forming the activity of intelligent beings. Intelligent actions are embodied, so perception remains connected to the anatomical structure of the subject and its rationality is associated with the world as well.

Artificial intelligence designers usually do not take into account the biosemiotic perspective, but they consider the world of the organism and the world of the machine as equal. Biosemiotics shows that intelligence is developed as the distinctiveness of a different species. Intelligent behaviour is the result of the functioning signs and references through the subject's history and experiences.

The purpose of the work is to show that an artificially constructed system can implement semiotic actions based on the concept of Jakob von Uexküll. It is possible to design artificial intelligence equipped with a phenomenal, subjective world. The document presents the basics of intelligent behaviour of artificial systems from the perspective of biosemiotics and considers the possibility of the emergence of machine intelligence as a new form.

Keywords

- c) Jakob von Uexküll, Steven Harnad, Tom Ziemke, Claus Emmeche, Thomas Sebeok, Jesper Hoffmeyer, John Searle, Hubert Dreyfus, Rodney Brooks
- a) biosemiotics, artificial intelligence, radical constructivism, autonomy, intentionality, adaptation, evolution, cognitive object, phenomenal world, Umwelt, functional circles, agent, artificial networks, robots, algorithms, ALife

Spis treści

Spis treści.....	4
Wstęp.....	6
1 Najważniejsze problemy programu badań sztucznej inteligencji.....	11
1.1 Inteligencja i sztuczna inteligencja – problem pojęciowy	11
2.1.....	15
2.2 GOFAI – Dobra Tradycyjna Sztuczna Inteligencja formalnych systemów symbolicznych	15
1.1.1 Symboliczna SI	16
1.1.2 Koneksjonizm jako próba ominięcia problemów symbolicznej SI.....	22
1.1.3 Zagadnienie ugruntowania symboli w systemach SI	26
2.3 Ucieleśniona sztuczna inteligencja – programy w świecie rzeczywistym.....	31
1.4 Fizyczne systemy symboliczne a problemy nadawania znaczenia	40
1.5 Nowe problemy i możliwe kierunki rozwoju programu SI	44
2 Zagadnienie inteligencji w perspektywie biosemiotycznej	50
2.1 Koncepcja semiozy w biologii.....	51
2.2 Zagadnienie znaczenia w pracach Jakoba von Uexküllä	58
2.3 Ucieleśnienie i usytuowanie	69
2.4 Autopoeza i autonomia.....	74
2.5 Semiotyczne rusztowania w semiosferze.....	76
2.6 Podmiot jako „semiotyczne ja”	81
2.7 Intencjonalność jako funkcja biologiczna	85
2.7.1 Działania intencjonalne jako czynnik ewolucyjny	90
2.7.2 Koncepcja nowej przyczynowości biologicznej	94
2.8 Biosemiotyczna koncepcja inteligencji	97

3	Modelowanie sztucznej inteligencji w biosemiotycznych ramach koncepcyjnych	102
3.1	Możliwości powstania inspirowanych biosemiotycznie inteligentnych systemów	103
3.1.1	Status semiotyczny sztucznych systemów.....	103
3.1.2	Koncepcja ucieleśnienia i usytuowania sztucznych systemów.....	108
3.1.3	Problem autonomii sztucznych systemów	111
3.1.4	Znaczenie zgodności strukturalnej percepcji i działania.....	115
3.1.5	Biosemiotyczna propozycja rozwiązania problemu ugruntowania symboli.....	117
3.1.6	Rola rusztowań poznawczych i rozproszonego poznania	120
3.2	Badania nad inteligentnymi agentami	122
3.2.1	Architektura subsumpcyjna Rodneya Brooksa.....	123
3.2.2	Badania nad komunikacją synergiczną	127
3.2.3	Działanie sieci neuronowych w perspektywie biosemiotycznej.....	129
3.2.4	Metody ewolucyjne i robotyka adaptacyjna.....	133
3.2.5	Koncepcja działania filtra semiotycznego	139
3.2.6	Propozycja rozwiązania problemu intencjonalności agentów.....	142
4	Radykalny konstruktywizm jako paradygmat rozwoju badań nad sztuczną inteligencją	148
4.1	Sztuczne Umwelty – perspektywa powstania sztucznych fenomenalnych światów	153
4.2	Konsekwencje przyjęcia radykalnego konstruktywizmu – odmienność inteligencji maszynowej.....	158
4.3	Sztuczne życie jako najbardziej perspektywiczny program dla rozwoju sztucznej inteligencji?	165
4.4	Integracja sztucznej i ludzkiej inteligencji.....	170
	Zakończenie	182
	Bibliografia	193

Wstęp

Systemy sztucznej inteligencji są wyjątkowymi obiektami do badania procesów semantycznych. Oferują możliwość systematycznej, dogłębnej analizy złożonych interakcji system-środowisko, w których większość parametrów jest znanych, ponieważ systemy te zostały zaprojektowane i skonstruowane w kontrolowany sposób. Rozwój sztucznej inteligencji warunkowany jest równocześnie w dużym stopniu zrozumieniem zjawisk powstawania znaczenia w wyniku ucieleśnionej interakcji systemu i środowiska. Celem powyższej pracy jest analiza procesów powstawania inteligencji rozumianych jako wyniku przetwarzania znaków i powstawania znaczenia w odniesieniu do sztucznych systemów. Zostanie tu zbadana rola symboli i innych znaków, które odgrywają istotną rolę w procesie kształtowania się inteligencji. Wykorzystane zostanie podejście wypracowane na gruncie biosemiotyki z jej perspektywą etologiczną, które jest najbardziej zaawansowanym obszarem badań nad powstawaniem inteligentnych zachowań. Podejście biosemiotyczne, nie zawężające koncepcji inteligencji jedynie do gatunku *Homo sapiens* rozszerzone zostanie na domenę maszyn, używających skutecznych systemów przetwarzania znaków. Stanowi to istotne *novum* rozprawy, zakreślając możliwości stworzenia nowego paradygmatu badań nad sztuczną inteligencją.

Zagadnienia semantyki rozważane były od dawna w kontekście systemów adaptacyjnych i ewoluujących. Od czasu przełomowej w dziedzinie biologii pracy Jakoba von Uexküll'a wielu biologów, etologów, psychologów i filozofów było zainteresowanych wyjaśnieniem, w jaki sposób powstają znaczenia w interakcji system-środowisko. Z drugiej strony brak semantyki sztucznych systemów jest najczęściej podnoszonym problemem w obrębie filozofii sztucznej inteligencji i jest źródłem najważniejszych argumentów krytycznych formułowanych wobec niej. Tym bardziej dziwi, jak niewielka ilość publikacji i badań w dziedzinie sztucznej inteligencji w pozytywny sposób dotyczy tej fundamentalnej kwestii. Stanowisko badaczy w niewielkim stopniu odnosi się do badania inteligencji jako cechy charakterystycznej dla życia. Tymczasem szeroko rozumiane badania biologiczne ukazują, że dynamika interakcji system-środowisko wymusza na podmiocie poznawczym stosowanie coraz to bardziej inteligentnych działań, związanych z jego rozwojem i adaptacją. Aktywność żywych organizmów realizuje w sobie funkcje, których zrozumienie i odtworzenie w systemie sztucznym wydaje się kluczowe dla dalszego rozwoju systemów sztucznej inteligencji, jak:

ukierunkowanie na cel, intencjonalność, autonomia działań, samoodniesienie, etc. Biosemiotyka w wyraźny sposób przedstawia podstawowe mechanizmy tych funkcji, zatem już dawno powinna zostać wykorzystana w dziedzinie sztucznej inteligencji. Racjonalne jest więc, aby techniczne podejście do problemu ugruntowania symboli i tworzenia semantycznych powiązań powinno zostać dokonane na bazie solidnych podstaw teoretycznych biosemiotyki, np. dla pojęć takich jak znaczenie lub odniesienie.

Współcześnie w obrębie badań kognitywistycznych istnieje zgoda co do tego, że interakcja sensomotoryczna ze środowiskiem ma fundamentalne znaczenie dla poznania. Z historycznego punktu widzenia docenienie roli sensomotorycznego ciała w poznaniu jest przynajmniej częściowo spowodowane koncepcyjną zmianą w badaniach nad sztuczną inteligencją, szczególnie w modelach obliczeniowych, robotycznych i w sztucznym życiu. Jednak większość badań nad ucieleśnioną sztuczną inteligencją, w szczególności prace nad ugruntowaniem symboli i związanymi z tym problemami, redukuje ciało do zwykłego interfejsu sensomotorycznego dla procesów wewnętrznych, które są nadal realizowane w standardowych modelach obliczeniowych opartych na tradycyjnym podejściu. Wydaje się, że badania sztucznej inteligencji przeszły tylko z obliczeniowego do robotycznego funkcjonalizmu. Jest to poważne ograniczenie dla badań, ponieważ postrzeganie ciała fizycznego jako interfejsu sensomotorycznego dla obliczeniowego umysłu wciąż warunkuje rozwój dziedziny. Dlatego też staramy się tu przedstawić argumenty, że istnieje wiele powodów, aby rozwijać sztuczną inteligencję na koncepcyjnym fundamencie biosemiotyki.

Biosemiotyka będzie tu rozumiana tutaj szerzej niż się to zwykle czyni, jako teoria przetwarzania znaków w różnego rodzaju systemach, zarówno naturalnych jak i sztucznych. Należy zastrzec, że niniejsza analiza jest przede wszystkim oparta na tartusko-kopenhaskiej szkole biosemiotyki, która stara się adaptować amerykańską tradycję semiotyczną Charlesa Sandersa Peirce. W niniejszej rozprawie podejście biosemiotyczne wspomnianej szkoły zostało rozszerzone na obiekty sztuczne, odpowiednio dostosowując oryginalne koncepcje. Jako metodę zastosowano tu analizę tekstów, aby wskazać na główne paradygmaty podejścia do interesujących nas zagadnień obecne zarówno w tekstach biosemiotyków jak i badaczy sztucznej inteligencji. W części poświęconej aplikacji koncepcji biosemiotycznych do sztucznej inteligencji skorzystano również z elementów testu komutacyjnego, który pozwolił zidentyfikować istotne dla tej pracy charakterystyczne rodzaje procesów przetwarzania znaków, które są istotne dla przebiegu procesu rozwoju inteligencji. Intencją autora było pokazanie w jaki sposób użycie znaków jednego systemu przez inny system może wpłynąć na przekazanie pierwotnej relacji semiotycznej. Ponadto, umieszczając oryginalne relacje znaków

w różnych kontekstach, dało się uzasadnić, że znaki biologiczne mogą zostać skutecznie użyte przez sztuczne systemy.

Techniki sztucznej inteligencji stanowią rdzeń nowych technologii we współczesnym społeczeństwie, dlatego atrakcyjne może wydawać się wykorzystanie do ich rozwoju rozwiązań inspirowanych biologicznie w obszarze działań semiotycznych. W kontekście sztucznych systemów pojęcia związane z semantyką zbyt często przyjmują formę metaforyki, która niezmiennie odnosi się do obliczeń. Najważniejsze pytania związane z możliwością użycia koncepcji biosemiotycznych w projektowaniu nowego rodzaju maszyn, są następujące:

- Czy założenia filozoficzne biosemiotyki wydają się być wystarczająco ogólne, dla badania użycia znaków i znaczenia w sztucznych systemach?
- Czy procesy przetwarzania znaków w sztucznych systemach można postrzegać jako specyficzne rodzaje semiozy?
- Czy inteligencja jest wynikiem procesów specyficznych dla różnych systemów i czy w związku z tym, w przypadku inteligencji sztucznych systemów można w adekwatny sposób używać pojęć stworzonych w naukach kognitywnych albo w informatyce?

W związku z powyższym, celem tej pracy jest przedstawienie możliwości adaptacji i implementacji podejścia biosemiotycznego i włączenie go do aktualnej debaty o rozwoju sztucznej inteligencji. W celu uzyskania odpowiedzi na powyższe pytania, praca została podzielona na cztery rozdziały, które realizują następujące zadania.

Pierwszy rozdział ukazuje rozwój programu sztucznej inteligencji, z wyszczególnieniem jej problemów wynikających z braku implementacji właściwej interakcji sensomotorycznej ze środowiskiem oraz problemów związanych z krytowanym brakiem własnej semantyki sztucznych systemów. Rozdział jest podzielony na część poświęconą klasycznemu podejściu sztucznej inteligencji (ang. GOFAI – *Good Old-Fashioned Artificial Intelligence*), część prezentującą paradygmat ucieleśnienia w badaniach sztucznej inteligencji (ang. EAI – *Embodied Artificial Intelligence*) oraz część wskazującą na ich wspólne problemy.

Rozdział drugi jest próbą stworzenia biosemiotycznej koncepcji inteligencji, uwzględniająca idee prekursorów biosemiotyki. Uwaga została skupiona na teorii *Umweltu* i działania kręgów funkcjonalnych Jakoba von Uexküll'a oraz oparte na niej koncepcje tworzone przez jego kontynuatorów. Dobór koncepcji biosemiotycznych podporządkowany został głównemu celowi pracy, czyli możliwościom zastosowania ich w tworzeniu sztucznej inteligencji. W tej części zawarto opis pojęć, które niezależnie bada biosemiotyka, a które mogą i powinny być zastosowane do badania inteligencji w ogóle.

W trzecim rozdziale zademonstrowane zostały możliwości zastosowania biosemiotycznych rozwiązań w sztucznych systemach, a zwłaszcza powiązanych z inteligencją pojęć takich jak: sztuczna semioza, ucieleśnienie, autonomia i autopoeza, podmiotowość, znaczenie, intencjonalność. Pierwsza część z tych rozważań ma wyłącznie wymiar teoretyczny, natomiast druga część przedstawia dotychczasowe próby aplikacji koncepcji biosemiotycznych oraz gotowe modele, które mogą zostać zaimplementowane.

W ostatnim rozdziale zastosowano idee radykalnego konstruktywizmu do uzasadnienia możliwości zastosowania koncepcji biosemiotycznych w sztucznych systemach. Następnie wskazano możliwe ścieżki rozwoju nieantropocentrycznej sztucznej inteligencji oraz przeprowadzono analizę możliwych sposobów rewizji paradygmatu badań nad sztuczną inteligencją. W ten sposób rozprawa poza realizacją celów czysto filozoficznych wskazuje również na możliwości zastosowania wypracowanego tu podejścia w praktyce badań nad sztuczną inteligencją, wpisując się w nurt filozofii stosowanej (ang. *applied philosophy*) stanowiący współcześnie jedną z ważniejszych determinant rozwoju współczesnych technik cyfrowych.

Wyniki uzyskane w rozprawie podsumowane zostały w zakończeniu. Praca zawiera obfita bibliografię składającą się z 283 pozycji. Powodem tak wysokiej liczby odniesień do literatury jest fakt, że problem inteligencji jako takiej nie był dotychczas szczególnie opisywany w tekstach biosemiotycznych. W celu znalezienia zgodnej biosemiotycznej koncepcji inteligencji należało się przyjrzeć wielu pozycjom bibliograficznym z tej dyscypliny i ustalić biosemiotyczną koncepcję zachowań inteligentnych, co stanowi kolejny wkład niniejszej rozprawy.

Kilka słów należy również poświęcić kwestii terminologii. Badania biosemiotyczne przedstawiane są głównie w języku angielskim i niemieckim, a część pojęć nie posiada wciąż polskich odpowiedników. W dysertacji starano się przetłumaczyć te konstrukcje w pełni tak, aby jak najwierniej oddać ukryte w nich intuicje filozoficzne, co nie zawsze dawało pojęcia zgodne z wymogami translacji i językoznawstwa, lecz tak użyte pojęcia wydawały się najdogodniejsze pod kątem prowadzonych analiz. W miejscach, w których nie do końca było możliwe znalezienie sensownego odpowiednika w języku polskim pozostawiano oryginalny zapis lub dokonano przypisu wyjaśniającego intencje i założenia definicyjne autora. Przykładem trudności pojęciowych występujących w analizowanych zagadnieniach niech będzie choćby dyskusja przedstawiona w paragrafie 2.7.1. niniejszej pracy.

W końcu wyjaśnić wypada również stanowisko autora pracy. Pod względem metodologicznym praca może być traktowana jako interdyscyplinarna, łączy bowiem podejście

filozoficzne, kognitywistyczne, biosemiotyczne i inżynierskie. Wspomniany interdyscyplinarny obszar poddany został analizie z punktu widzenia filozofii oraz kognitywistyki, zarówno z uwagi na cele niniejszej rozprawy, jak i z uwagi na moje wykształcenie akademickie. Jednocześnie uważam, że ze wszystkich dyscyplin związanych w tym momencie z programem sztucznej inteligencji biosemiotyka ma największą żywotność i potencjał heurystyczny, szczególnie w odniesieniu do fundamentalnych założeń filozoficznych SI. Biorąc to pod uwagę, praca ta oferuje spojrzenie na inteligencję i sztuczną inteligencję z zewnątrz, szczególnie z punktu widzenia biologii teoretycznej i filozofii kontynentalnej reprezentowanej przez tradycje francuskie i niemieckie.

1 Najważniejsze problemy programu badań sztucznej inteligencji

1.1 Inteligencja i sztuczna inteligencja – problem pojęciowy

Badania w dziedzinie sztucznej inteligencji borykają się od samego początku z problemem braku klarownej i uzgodnionej definicji samego pojęcia "inteligencja". Rozważania nad inteligencją opierają się na poszukiwaniu obiektywnych własności organizmu, które mogą zostać zmierzone. Definicje pojęcia "inteligencja" są tworzone poprzez znajdowanie przykładów paradygmatycznych i tworzenie list właściwości. Są to zazwyczaj charakterystyki opisujące własności i funkcje działania żywych istot prowadzące do efektywnego działania w świecie. Inteligentnym zachowaniem określa się adekwatne i skuteczne zachowanie prowadzące do osiągnięcia konkretnego celu. W zależności od konwencji oraz aktualnych ustaleń definicję rozszerza się o kolejne desygnaty, redukuje się zaś te uznane za błędne. Brak jest jednak definicji równoznacznej, ponieważ pojęcie inteligencji jest nieokreślone i niejednolite. Potoczne i naukowe opisy inteligencji się mnożą. Jednak w obecnym ujęciu określa się ją jako umiejętność myślenia, wnioskowania, podejmowania decyzji, rozpoznawania wzorców, przewidywania efektów własnych działań, umiejętność zapamiętywania, uczenia się, rozwiązywania problemów, postrzegania, ale również zdolność użycia języka, umiejętności ruchowe, korzystanie z intuicji, bycie kreatywnym i świadomym, czyli ogólne radzenie sobie w świecie¹.

Sposoby opisywania inteligencji nieustająco zmieniają się z biegiem czasu. To co je łączy, to poszukiwanie charakterystycznych własności, konkretnych wyznaczników, które usprawiedliwiałoby nazwanie zachowania inteligentnym. Rozważania czym miałyby być te umiejętności również przyjmują różne kierunki. Na początku badań nad sztuczną inteligencją, od połowy lat pięćdziesiątych do końca lat osiemdziesiątych, wszystkie metody opierały się na rozważaniach symbolicznych reprezentacji wysokiego poziomu. Za przejawy inteligencji uważano głównie umiejętność logicznego wnioskowania, która przybierała różne konkretne zastosowania: zdolność gry w szachy, rozwiązywania zadań matematycznych lub ścisłej dedukcji. Badania nad sztuczną inteligencją w dużej mierze opierały się na potocznym bądź psychologicznym rozumieniu tego pojęcia. Potocznie za inteligencję uważa się najczęściej zdolność rozwiązywania problemów praktycznych,

1 R. Pfeifer, C. Scheier, *Understanding intelligence*, MIT Press, 2001.

zdolności językowe lub kompetencje społeczne². W psychologii najczęściej uważano ją za sprawność determinującą efektywność działań, wykorzystującą procesy poznawcze, która jest kontrolowana przez nadrzędną zdolność umysłu, odpowiadającemu umysłowemu poziomowi jednostki, oraz zdolności specjalne, odpowiadające za sprawność działania w zakresie konkretnych dziedzin czy rodzajów zadań³. Program badań zaś zawsze bazował na aktualnej wiedzy i przekonaniach dotyczących fenomenu inteligencji.

Sztuczna inteligencja (SI) jest nauką i inżynierią tworzenia inteligentnych maszyn oraz inteligentnych programów komputerowych⁴. Dziedzinie tej stawia się również za cel zrozumienie ludzkiej inteligencji przy pomocy komputerów. Istnieją dwa podejścia do sztucznej inteligencji. Pierwsze z nich to silna sztuczna inteligencja; termin ten oznacza, że sztuczny system ma posiadać umiejętności analogiczne do ludzkich: świadomość, zdolność myślenia, umiejętność używania języka, użycie wiedzy i doświadczeń przy rozwiązywaniu problemów. W silnym wersji SI system może być postrzegany jako umysł i to szczególne ujęcie jest rozważane w tej pracy. Istnieje również pojęcie słabej sztucznej inteligencji, które oznacza system poprawnie przetwarzający dane – taki system umiejętnie wykonuje zadania i przedstawia zadowalające dane wyjściowe. W drugim ujęciu refleksja filozoficzna nie ma większego znaczenia: jeśli programy zachowują się tak, „jakby były inteligentne”, cel został osiągnięty. Dyscyplina sztucznej inteligencji nie ogranicza się do metod, których skutki przypominają te obserwowane w psychologii i/lub biologii. W sztucznych systemach szuka się obliczeniowej zdolności do osiągania określonych celów w świecie rzeczywistym. Mimo tego komputacyjnego podejścia nie istnieje definicja sztucznej inteligencji, która nie jest zależna od definicji inteligencji żywych organizmów.

Arthur Jensen, jeden z czołowych psychologów, zajmujących się badaniem ludzkiej inteligencji, postawił hipotezę, że wszyscy ludzie mają takie same intelektualne mechanizmy, a różnice w inteligencji są związane z „ilościowymi warunkami biochemicznymi i fizjologicznymi”⁵. Miał na myśli takie umiejętności jak uwaga, percepcja, uogólnianie, uczenie się, pamięć, język, myślenie i rozwiązywanie problemów. Programy komputerowe mają dużą szybkość działania i olbrzymią pamięć, rozwiązują

2 *Inteligencja*, [w:] <https://encyklopedia.pwn.pl/haslo/inteligencja;3915042.html>.

3 C. Spearman, „*General Intelligence, Objectively Determined and Measured*,” *The American Journal of Psychology*, t. 15, nr 2, 1904, ss. 201–292.

4 M. Flasiński, *Wstęp do sztucznej inteligencji*, Warszawa 2011.

5 A.R. Jensen, *The g factor and the design of education*, „*Intelligence, instruction, and assessment: Theory into practice*”, 1998, ss. 111–131.

trudne problemy, ale ich umiejętności odpowiadają tym intelektualnym mechanizmom, które zostały wystarczająco dobrze opisane i zrozumiane, by projektanci sztucznej inteligencji mogli je implementować w sztucznych systemach. Część intelektualnych umiejętności nie jest jednak dostępna dla sztucznych systemów. Związane jest to niewątpliwie z faktem, że naukom kognitywnym wciąż nie udało się dokładnie określić, jakie są ludzkie możliwości. Podejmując temat inteligentnego zachowania żywych organizmów, rzeczywiście zastanawiamy się nad mechanizmami leżącymi u podstaw takiego działania. W pierwszej kolejności wyróżniane są właściwości, które zdają się być kluczowe dla fenomenu inteligencji. Nie rozważa się jej jako całości zachowań, lecz próbuje się wyodrębnić te cechy, które wydają się istotne i decydujące o inteligentnym działaniu.

Podobnie jest w badaniach SI. Ważną rolę odgrywa również to, że inteligentne funkcjonowanie organizmów postrzegamy przez pryzmat działań, które wydają nam się możliwe do odwzorowania w symbolicznym systemie. Powstają zatem programy komputerowe lub roboty, które wykazują cechy inteligentnego zachowania. Jednak jak zaobserwował Rodney Brooks „badacze sztucznej inteligencji zauważają, że SI jest często pozbawiana należnego jej sukcesu. Ci którzy odrzucają SI twierdzą, że problem rozwiązany przez algorytm, nie jest prawdziwym problemem SI. Dlatego SI nigdy nie odnosi sukcesu”⁶. Za przykład może posłużyć komputer IBM – Deep Blue, który zwyciężył w szachy z mistrzem świata Garrijem Kasparowem⁷. Nie został nazwany inteligentnym, mimo że umiejętność wygrywania w tej grze była (i wciąż jest) postrzegana jako jeden z czołowych przejawów inteligencji. Paradoksalnie, wielki przegrany tego pojedynku, z racji tych samych umiejętności, jest uważany za nieprzeciętnie inteligentnego człowieka⁸. O Deep Blue mówiono, że jego wygrana możliwa była tylko dzięki szybkości przeprowadzania obliczeń (mógł zbadać do 200 milionów możliwych pozycji szachowych na sekundę) oraz modyfikacjom, które pozwalały zaadaptować program do stylu gry Kasparowa. Ta ilustracja pokazuje, że mimo prób odtwarzania poszczególnych cech

6 „Artificial intelligence researchers are fond of pointing out that AI is often denied its rightful successes. AI detractors claim that since the problem is solvable by an algorithm, it is not really an AI problem. Thus, AI never has any successes.”, tłum. własne. R.A. Brooks, *Intelligence without representation*, „Artificial Intelligence”, t. 47, nr 1-3, 1991, ss. 139-159.

7 IBM100 - Deep Blue,

<http://www.03.ibm.com/ibm/history/ibm100/us/en/icons/deepblue/> (2017).

8 R. Weijermars, *Building Corporate IQ – Moving the Energy Business from Smart to Genius: Executive Guide to Preventing Costly Crises*, Springer Science & Business Media 2011, s. 4.

inteligencji, zawsze poszukujemy ich w szerszym kontekście, którego brakuje sztucznym systemom.

Inteligencja jako zdolność organizmu do złożonych reakcji na bodźce środowiska jest mocno powiązana ze zdolnością do myślenia. Jest to poniekąd kartezjańskim dziedzictwem zachodniej kultury. Dla Kartezjusza istniały dwie oddzielne substancje: cielesna i mentalna⁹. Podział ten zrodził historyczny problem, w jaki sposób między dwoma systemami – ciałem i umysłem – mogą przebiegać wzajemne relacje, skoro każda z nich może istnieć bez drugiej. Jednym z głównych wyzwań stawianych przez ten dylemat jest pytanie, w jaki sposób myśl, czyli coś, co dzieje się w niematerialnym umyśle, może wpływać na materialne ciało. Zagadnienie dualizmu substancjalnego do dziś jest rozważane we współczesnej nauce jako problem ciało-umysł. Kartezjusz twierdził też, że ciała zwierząt i ludzi są prostymi mechanizmami, ponieważ są poddane czynnikom zewnętrznym i działają dzięki siłom mechanicznym. Założeniem kartezjańskiego dualizmu jest, że materialne ciało posiada określone właściwości fizyczne a ich opis jest zapewniony przez matematyczne przyrodoznawstwo¹⁰. Jedynie umiejętność myślenia, która przynależy tylko człowiekowi, pozwala na racjonalne i świadome kierowanie własnym ciałem. Samoświadomość zaś jest skutkiem sposobu rozumienia siebie poprzez przyczynowość lub logikę oraz warunkiem poznania. Ten pogląd sprawił, że umysł ludzki postrzegany jest jako swoista wartość a wszystkie inne istoty są mierzone jego miarą.

Większość ludzi prawdopodobnie zgodziłaby się, że zjawiska psychiczne, takie jak myślenie, pochodzi z procesów w mózgu, jak również z tym, że to umysł kontroluje tego typu działania. Jednak opis tej aktywności umysłowej przysparza pewnych fundamentalnych problemów. Dotychczas ten rodzaj działania umysłu nie został kompleksowo opisany. Mimo że wszyscy mamy całkiem jasne pojęcie o tym, co rozumiemy przez termin „myślenie”, pozostaje on słabo zdefiniowany. Myślenie rozumiane jako możliwość indukowania, wnioskowania, szukania analogii, przeprowadzania obliczeń i rozpoznawania wzorów w dużej mierze pokrywa się w opisie z inteligentnym zachowaniem. Dlaczego więc Deep Blue, który spełnił opisane wyżej wymagania i wykonał oczekiwane działania, nie został nazwany inteligentnym? Podobnie

9 R. Descartes, *Rozprawa o metodzie*, T. Żeleński (tłum.), Warszawa 1952.

10 R. Descartes, *Medytacje o pierwszej filozofii*, M. Ajdukiewicz, K. Ajdukiewicz (tłum.), Kęty 2001.

inne programy naśladowujące mechanizmy poznawcze człowieka również nie uzyskują tego przydomka. Ciekawy jest przykład systemów eksperckich, które odkrywają wzorce i znajdują najlepsze z możliwych rozwiązań. Żaden program diagnostyczny nie zostanie jednak uznany za inteligentniejszy od przeciętnego lekarza, nawet jeśli jego baza wiedzy będzie większa, praca szybsza a diagnozy trafniejsze. Umiejętność rozwiązywania problemów stawiana jest zawsze wysoko wśród cech inteligentnych organizmów. Związana jest ze zdobytą wcześniej wiedzą i możliwością jej kreatywnego wykorzystania. Programy rozpoznają emocje, chatboty budują wielokrotnie złożone, poprawne gramatycznie zdania. Istnieją również programy tworzące obrazy i muzykę. Ponadto możliwości sztucznych systemów, szczególnie te związane z możliwościami obliczeniowymi, często znacznie przekraczają umiejętności ludzkie.

W tym rozdziale niezbędne jest przedstawienie wybranych historycznych uwarunkowań i podstawowych założeń SI, ponieważ wskaże to istotne problemy badań nad sztuczną inteligencją. Część poświęcona tradycyjnej sztucznej inteligencji zostanie przedstawiona w możliwie zwięzły sposób, ponieważ ten temat był wielokrotnie omawiany i nie ma potrzeby, żeby szeroko opisywać wielokrotnie już dyskutowane tematy. Kolejny podrozdział dotyczy ucieleśnionej sztucznej inteligencji i założeń, które przyświecały jej twórcom oraz sposobom konstruowania ucieleśnionych systemów. Następnie przedstawiony zostanie sposób działania fizycznego systemu symbolicznego ze szczególnym podkreśleniem jego pozytywnych aspektów jak i podstawowych wad. Na końcu rozdziału zostaną wskazane kierunki rozwoju programu sztucznej inteligencji.

1.2 GOFAI – Dobra Tradycyjna Sztuczna Inteligencja formalnych systemów symbolicznych

W 1950 r. brytyjski matematyk Alan Turing opublikował doniosły artykuł w czasopiśmie filozoficznym *Mind* zatytułowany „Maszyny liczące i inteligencja”¹¹. Idea przetwarzania informacji w sposób obliczeniowy zostały uznane za podstawową funkcję inteligencji. Opisany przez niego (i nadal dominujący) model obliczeń powstał dzięki obserwacji zachowań rachmistrza – ludzkiego komputera, wykonującego obliczenia za

11 A.M. Turing, *Computing machinery and intelligence*, „Mind”, 1950, ss. 433–460.

pomocą pióra i papieru oraz przestrzegającego ustalonych reguł. Mimo iż Turing modelował to, co robi osoba, a nie to, co ona myśli, jego metafora stała się samą definicją obliczeń. Używając przykładu maszyny analitycznej Babbage'a, dowodził, że takie obliczenia są niezależne od medium. Model obliczeniowy Turinga potraktowany został później jako metafora mózgu, a z nim również utrwalił przekonanie, że ludzki mózg jest maszyną Turinga, wykonującą obliczenia i kontrolującą ludzkie działania.

Artykuł Turinga w pewnym sensie wyznaczył początek badań sztucznej inteligencji, choć pojęcie to pojawiło się znacznie później. Autor postawił w nim kluczowe pytanie otwierające wspomniany tekst: „Proponuję rozważyć pytanie: Czy maszyny mogą myśleć?”¹². Zaproponował w nim test na inteligencję lub myślenie, który nazwał grą w naśladownictwo. Jej celem było sprawdzenie czy osoba badana nie widząca swojego interlokutora, może odróżnić człowieka od komputera. Gra polegała na wpisywaniu pytań w terminalu komunikacyjnym i odgadywaniu, kto odpowiada na pytania: człowiek czy program komputerowy. Jeśli komputer miałby zdolność naśladowania człowieka, można by założyć, że potrafi myśleć. Test ten, wszedł do literatury pod nazwą Testu Turinga. Wzbudził wiele kontrowersji i dyskusji na temat możliwości zbudowania maszyny zdolnej do zaliczenia testu. Turing przewidywał, że w ciągu 50 lat pojawią się komputery, które zdadzą ten egzamin. To twierdzenie stało się pierwszym z wielu fałszywych przewidywań dotyczących sztucznej inteligencji.

1.2.1 Symboliczna SI

Projekt badań i nazwa sztuczna inteligencja powstały dopiero w 1956 roku na interdyscyplinarnej konferencji w Dartmouth, New Hampshire – dwa lata po śmierci Alana Turinga. Do udziału w tym wydarzeniu John McCarthy zaprosił wielu naukowców badających szeroki zakres zaawansowanych zagadnień, takich jak: teoria złożoności, symulacja języka, sieci neuronowe, abstrakcyjne dane otrzymywane z wejść sensorycznych maszyn, związek losowości z myśleniem twórczym i uczeniem się. W Dartmouth obecni byli czołowi badacze tej nie ukonstytuowanej wcześniej dziedziny:

¹² *Ibid.*

Marvin Minsky, Nathaniel Rochester, Claude Shannon, Allen Newell i wielu innych¹³. Już w 1956 roku zauważono, że szybkość i funkcjonalność komputerów elektronicznych podwaja się co osiemnaście miesięcy, a stopa wzrostu nie wykazuje oznak spowolnienia¹⁴. Konferencja była jedną z pierwszych poważnych prób rozważenia konsekwencji tej wykładniczej krzywej. Wielu uczestników było przekonanych, że dalsze postępy w powiększaniu się elektronicznej prędkości, pojemności i programowaniu doprowadzą do tego, że komputery zyskają zasoby inteligencji jak istoty ludzkie. W Dartmouth skryształizowały się koncepcje, które dały początek nowej dziedzinie jako obszarowi badań interdyscyplinarnych i do dziś stanowią tło intelektualne dla wszystkich badań informatycznych. Postawiono sobie za cel stworzenie dziedziny nauki, która będzie rozwijała badania nad budową inteligentnych maszyn. Konferencja była pierwszym spotkaniem poświęconym takiemu tematowi. Stworzyła fundamenty dla idei, która od dawna wpływała na badania i rozwój inżynierii, matematyki, informatyki, psychologii i wielu innych dziedzin. Mimo tego, całościowy przedmiot dysput był na tyle nowy, że należało stworzyć nowy termin: sztuczna inteligencja.

Inteligencja była wówczas postrzegana jako zdolność pojmowania czegoś rozumem, lecz trwały badania nad nowoczesnym rozumieniem tego pojęcia. Na nową dyscyplinę SI znacząco wpływały tzw. czynnikowe teorie inteligencji. Charles Spearman a później Philip Vernon twierdzili, że istnieje pewna zdolność umysłu, która jest nadrzędnym czynnikiem wpływającym na poziom umysłowy człowieka. Uznali też, że jest on odpowiedzialny za szczególne umiejętności w pewnych dziedzinach lub działaniach. Spearman wykorzystał metody matematyczne w swoich badaniach. Twierdził, że zachowanie podczas wykonywania zadań jest wynikiem zestawu ograniczeń lub wymagań, które wynikają z instrukcji, określonych na wielu poziomach abstrakcji¹⁵. To, co połączyło fundatorów nowej dyscypliny naukowej, to przekonanie, że każdy aspekt uczenia się lub jakakolwiek inna cecha inteligencji mogą być w zasadzie tak dokładnie opisane, że można skonstruować maszynę, która będzie je mogła symulować¹⁶. Dyskusja wokół projektu

13 J. McCarthy, M.L. Minsky, N. Rochester, *A Proposal for The Dartmouth Summer Research Project on Artificial Intelligence*, 1955.

14 *Ibid.*

15 C. Spearman, „General Intelligence,” *Objectively Determined and Measured...*, *op. cit.*

16 J. McCarthy, M.L. Minsky, N. Rochester, *A Proposal for The Dartmouth Summer Research Project on Artificial Intelligence...*, *op. cit.*

skupiała się więc na zagadnieniu, w jaki sposób komputer może symulować ludzkie myślenie i procesy poznawcze.

Metody opisanego podejścia określonego skrótem GOFAI¹⁷ pokrywały się z powyższym postrzeganiem inteligencji i zakładały silne twierdzenie, że inteligentne działania wynikają ze zdolności do wewnętrznej automatycznej manipulacji symbolami działającymi na zbiorze stabilnych przechowywanych reprezentacji¹⁸. Oznacza to, że wewnętrzne manipulacje symbolami muszą być interpretowane jako dotyczące świata zewnętrznego oraz muszą być wykonywane przez jakiś podsystem. Przyjęło się twierdzenie Allena Newella, że najważniejszym odkryciem sztucznej inteligencji i informatyki było stworzenie pojęcia fizycznego systemu symbolicznego¹⁹. Była to koncepcja systemów zdolnych do posiadania i manipulowania symbolami, które są możliwe do zrealizowania w fizycznym świecie. Przyjęto, że są to te same symbole, których używa człowiek. Hipoteza Newella twierdziła, że „ludzie są przedstawieniem fizycznych systemów symboli i dzięki temu umysł wchodzi do fizycznego wszechświata”²⁰.

W GOFAI systemy symboliczne odgrywały nadrzędną rolę w inteligentnym działaniu. Z punktu widzenia SI inteligentne działanie wymaga oznaczenia prawdopodobnych sytuacji, co zapewniają systemy symboli. Ludzkie umysły działają poprzez tworzenie i manipulowanie obiektywnymi reprezentacjami środowiska i wykonują zadania tak jak systemy symboliczne. Poznanie jest również postrzegane jako proces manipulowania symbolami i przetwarzania informacji. Zgodnie z tą perspektywą, poznanie jest konsekwencją przyjmowania informacji dostarczanych ze środowiska i formowania z nich reprezentacji, które mogą być przetwarzane w celu dostarczenia logicznych odpowiedzi poprzez działanie. Ta metafora zakłada pogląd, że inteligencja jest obliczeniowa i powstaje w wyniku działania wewnętrznego zaprogramowanego komputera lub ludzkiej maszyny do przetwarzania informacji, jak również, że mózg przetwarza symbole, które są ze sobą powiązane, tworząc reprezentacje świata zewnętrznego.

17 *Good Old-Fashioned Artificial Intelligence* – Dobra Tradycyjna Sztuczna Inteligencja J. Haugeland, *Artificial Intelligence: The Very Idea*, MIT Press 1989.

18 *Ibid.*, s. 113.

19 A. Newell, H.A. Simon, *Computer science as empirical inquiry: Symbols and search*, „Communications of the ACM”, t. 19, nr 3, 1976, ss. 113–126.

20 “[...] humans are instances of physical symbol systems, and, by virtue of this, mind enters into the physical universe” tłum. własne, A. Newell, *Physical symbol systems*, „Cognitive science”, t. 4, nr 2, 1980, s. 136.

W kognitywnym, obliczeniowym ujęciu inteligencję uznawano więc za zdolność do skutecznego manipulowania symbolami. Zakładano, że rozumowanie ludzkie jest podobne do wyrafinowanych obliczeń komputerowych i formalnych działań matematycznych. Abstrahowano od umysłu oraz od biologicznego mózgu, próbując zbadać w ten sposób właściwości inteligentnego myślenia niezależnie od możliwych fizycznych realizacji systemu. Inteligentne działanie opierałoby się więc o następujące reguły:

- a) umysł i ciało są rozłączne;
- b) myślenie polega na manipulowaniu abstrakcyjnymi reprezentacjami świata;
- c) reprezentacje mogą być wyrażone w języku formalnym;
- d) i mogą zostać zaimplementowane w sztucznym systemie.

Systemy komputacyjnej sztucznej inteligencji nie próbują naśladować sposobu, w jaki mózg faktycznie działa, ale próbują odwzorować to działanie za pomocą procesu efektywnego wyszukiwania odpowiedzi w systemie. Zakładają, że skuteczne wykonywanie działań to inteligencja, bez względu na to, jak jest wykonywana²¹. Jednak pierwotną intencją sztucznej inteligencji było nie tylko opracowanie inteligentnych algorytmów, ale także zrozumienie naturalnych form inteligencji, które manifestują się podczas interakcji w świecie rzeczywistym.

Postrzeganie inteligencji jako konkretnych logiczno-matematycznych umiejętności umysłu oprowadziło to do ukształtowania dziedziny sztucznej inteligencji w formie, która pomijała środowiskowe uwarunkowania. Wnioskowanie, podejmowanie decyzji, wyprowadzanie wniosków indukcyjnych, rozwiązywanie problemów logicznych, umiejętność przeprowadzania obliczeń i użycie pamięci to działania, które mogą być przeprowadzone przez sztuczne systemy. Inteligencja rozumiana jako obliczeniowa umiejętność doprowadziła do powstania nauki o wytwarzaniu inteligentnych maszyn i programów komputerowych. Program badawczy oparty na założeniach podyktowanych przez takie rozumienie inteligencji, która może zostać bądź to rozłożona na poszczególne czynniki, bądź sprowadzana do pojedynczych mechanizmów nie mógł jednak ominąć pewnych fundamentalnych trudności.

21 J. Mingers, *Self-producing systems: implications and applications of autopoiesis*, New York 1995, s. 191.

Krytyka założeń SI i metafory komputerowej została przedstawiona przez Johna Searle'a, Huberta Dreyfusa i wielu innych²². Searle wystąpił przeciwko koncepcji stwierdzającej, że stany i procesy mentalne mogą być przedstawione wyłącznie z perspektywy syntaktycznej. W swoim eksperymencie myślowym²³ wskazywał, że ludzki umysł jest nie tylko formalną strukturą, lecz przede wszystkim posiada semantyczną treść. Nigdzie w procesie syntaktycznego użycia symboli przez sztuczny system nie pojawia się możliwość zrozumienia ich znaczenia, tj. zrozumienia ich treści semantycznej. Innymi słowy, formalny charakter programów komputerowych polega na przeprowadzaniu sekwencji syntaktycznych prowadzących do wnioskowania o interpretacji symboli lub przypisaniu tym symbolom znaczeń²⁴. Searle dowodził również, że systemom komputerowym brakuje mocy przyczynowej procesów umysłowych. Zatem oprogramowanie komputerowe nie ma własnej intencjonalności, która umożliwiałaby wpływanie na otoczenie²⁵. Wniosek Searle'a jest taki, że programom formalnym brakuje fundamentalnych i wystarczających podstaw do wywoływania zjawisk psychicznych. Obliczeniowe podejście, które zakłada istnienie poziomu symbolicznego między procesami neuronalnymi i stanami intencjonalnymi jest obarczone błędem. Dokładne zrozumienie ludzkiego umysłu musi być oparte na traktowaniu go jako zjawiska biologicznego i zrozumieniu stanów psychicznych jako wyniku procesów mózgowych. Cel ten nie może zostać osiągnięty bez założenia istotności semantycznej treści ani bez zrozumienia intencjonalności jako nieodłącznej cechy ludzkiego umysłu.

Znaczącą perspektywą kontestującą możliwość symulowania ludzkiego umysłu przez SI była fenomenologia. Według Huberta Dreyfusa metafora przedstawiona przez klasyczną teorię sztucznej inteligencji nie jest jedynym sposobem na przetworzenie informacji. Sposób przetwarzania danych przez mózg jest analogowy i globalny i nie polega na przypisywaniu symbolu do każdego bitu informacji. Dreyfus dowodził, że

22 Patrz R.A. Brooks, *Achieving Artificial Intelligence through Building Robots*, 1986; P.J. Denning, *The science of computing: Is thinking computable?*, „American Scientist”, t. 78, nr 2, 1990, ss. 100–102; H. Dreyfus, S. Dreyfus, *Making a Mind Versus Modelling the Brain: Artificial Intelligence Back at a Branchpoint*, „Daedalus”, t. 117, nr 1, 1988, ss. 185–197; S.R. Graubard, *The artificial intelligence debate*, MIT press Cambridge (Ma) 1988; J.R. Searle, *Minds, brains, and programs*, „Behavioral and brain sciences”, t. 3, nr 03, 1980, ss. 417–424.

23 Patrz Eksperyment Chińskiego Pokoju w: J. Searle, *Umysły, mózgi i programy*, [w:] *Filozofia umysłu*, B. Chwedeńczuk (red.), t. II, Warszawa 1995, ss. 301–324.

24 *Ibid.*

25 J.R. Searle, *The rediscovery of the mind*, Cambridge 1992, ss. 79–80.

należy rozważyć sposób działania mechanizmów biologicznych w mózgu, które nie wyglądają jak obliczeniowe przetwarzanie danych²⁶. Starał się udowodnić, że ludzkie zachowanie w teorii SI jest rozważane poprzez pryzmat poziomu przetwarzania informacji, które znacząco różni się o neuronowego. Człowiek, wbrew założeniom behawioryzmu, nie jest jedynie maszyną udzielającą odpowiedzi na zewnętrzne bodźce. Ten błąd pojawia się, gdy nie istnieje rozróżnienie między poziomem neuronowym a fenomenologicznym. Zachowanie rozróżnienia między dwoma poziomami usuwa możliwość skrajnej redukcji wskazującej, że do człowieka można podejść z perspektywy fizycznej, jako urządzenia przetwarzającego dane; jego zachowanie nie może być wyjaśnione w tych kategoriach, ale tylko przez rozważenie fenomenologicznego wymiaru człowieczeństwa²⁷. Dreyfus przywołuje również argumenty przeciwko epistemologicznym założeniom SI. Formalne podejście przyjmuje, że jeśli człowiek jako maszyna do przetwarzania danych jest częścią świata fizycznego, podlega prawom fizyki ujętych w formułach matematycznych. Według Dreyfusa ten argument opiera się na pomieszaniu praw fizycznych i zasad przetwarzania danych. Żywe istoty nie przetwarzają wyłącznie formalnych symboli tak jak komputer, ale również inne informacje dotyczące właściwości fizycznych lub chemicznych analogowego świata²⁸.

Holistyczny charakter wiedzy podkreślany przez fenomenologię, oznacza, że żywe organizmy nie poznają rzeczy w odosobnieniu, ale zaczynają relację z nowymi przedmiotami lub istotami, posiadając już o nich pewną wiedzę. W ich umyśle istnieje rodzaj niejawnej znajomości kontekstu, który jest niezbędnym warunkiem zrozumienia tego, co się dzieje w świecie. Komputer nie może symulować wiedzy kontekstualnej bez rozpoznania istotnych cech danego kontekstu. Takie zadanie przekracza możliwości systemów formalnych, ponieważ zakłada możliwość sformalizowania całej ludzkiej wiedzy²⁹. Tak więc kontekstualizacja pozostaje wyłączną cechą żywych organizmów dzięki bezpośredniej interakcji ze światem. Odpowiednie zrozumienie działania umysłu należy zacząć od zrozumienia struktur fenomenologicznych, za pomocą których może on odnosić się do świata. Do poznania nie należy podchodzić w kategoriach struktur logicznych, ale zaczynając od ciała, które otwiera dostęp do świata. Pozwoliło to

26 H.L. Dreyfus, *What computers still can't do: a critique of artificial reason*, MIT press 1992, s. 72.

27 *Ibid.*, ss. 89–90.

28 *Ibid.*, ss. 106–107.

29 *Ibid.*, s. 100.

zrozumieć, że inteligencja jest produktem dynamicznej interakcji między ciałem a środowiskiem, w którym poznawczy podmiot jest usytuowany, mający bezpośrednie zrozumienie świata, który nie może być sformalizowany. Te cechy można uchwycić tylko za pomocą metod fenomenologii, do których redukcjonistyczna teoria obliczeniowa nie może mieć dostępu.

1.2.2 Koneksjonizm jako próba ominięcia problemów symbolicznej SI

Mimo powyżej opisanej celnej krytyki, naukowcy większą uwagę poświęcili próbom obrony GOFAI niż perspektywie fenomenologicznej. Sprzyjało temu pojawienie się prac nad nową postacią inteligentnych systemów inspirowanych biologicznie. W nurcie tym – zwanym koneksjonizmem – starano się modelować mechanizmy poznawcze w sztucznych sieciach neuronowych. Mimo, że w prace opierały się na badaniach nad działaniem mózgu, czyli reprezentowały podsystemowe spojrzenie na powstawanie reprezentacji, nie kwestionowało w ogóle działania sztucznych systemów symbolicznych. Sztuczne sieci neuronowe (SSN) powstały już w pierwszej połowie XX wieku, kiedy to po raz pierwszy zaproponowano modele sieci neuronów³⁰.

Renesans zainteresowania sztucznymi sieciami neuronowymi w latach 80. XX wieku w kontekście badań nad sztuczną inteligencją było wynikiem gwałtownego rozwoju neuronauk. Kiedy badacze sztucznej inteligencji zaczęli zdawać sobie sprawę z problemów z tradycyjnym podejściem do przetwarzania symboli, zaczęto wykorzystywać sztuczne sieci neuronowe, które miały za zadanie odzwierciedlać zdolność mózgu do uczenia się, adaptacji, użycia języka oraz tworzenia uogólnień. Działanie SSN opierało się na przetwarzaniu wzorów w miejsce wcześniejszego manipulowania symbolami, co dawało nadzieję, że będą lepiej opisywać naturalne zjawiska psychiczne lub tworzyć systemy eksperckie niż systemy algorytmiczne. Koneksjoniści podjęli próbę zastąpienia starszej metafory komputerowej metaforą mózgową, zakładając, że działania inteligentne są własnością zawartą w strukturze urządzenia. Inspirację czerpano z mózgu, ale tylko na bardzo abstrakcyjnym poziomie.

30 W.S. McCulloch, W. Pitts, *A logical calculus of the ideas immanent in nervous activity*, „The Bulletin of Mathematical Biophysics”, t. 5, nr 4, 1943, ss. 115–133.

Mózg składa się z wielu prostych jednostek przetwarzania (neuronów) połączonych równolegle przez dużą ilość połączeń (aksony i synapsy). Poszczególne jednostki (neurony) są zwykle wrażliwe jedynie na informacje lokalne. Z dużej ilości równoległych połączeń, prostych procesorów i lokalnych interakcji powstaje olbrzymia moc obliczeniowa i zdolność mózgu do rozwiązywania wielu problemów. Sztuczna sieć neuronowa to sieć prostych jednostek obliczeniowych, zwykle zorganizowanych w warstwy. Warstwa wejściowa otrzymuje różne informacje ze świata zewnętrznego. Są to dane, które sieć ma przetwarzać. Z jednostki wejściowej dane przechodzą przez jedną lub więcej ukrytych jednostek. Zadaniem ukrytej jednostki jest przekształcenie danych wejściowych w coś, czego może użyć jednostka wyjściowa. Różne części mózgu są odpowiedzialne za przetwarzanie różnych aspektów informacji a hierarchiczna struktura umożliwia danym przechodzenie przez poprzez różne poziomy neuronów. Podobnie w sieci neuronowej neurony w warstwach są połączone. Połączenia te są ważone: każda jednostka (sztuczny neuron) otrzymuje liczbę wejść numerycznych od innych jednostek, z którymi jest połączona, oblicza z ważonej sumy wartości wejściowych własną wyjściową wartość liczbową zgodnie z pewną funkcją aktywującą i przekazuje tę wartość jako dane wejściowe do innych neuronów. Po drugiej stronie sieci znajdują się jednostki wyjściowe, które reagują na dane, które zostały im przekazane i przetworzone.

Sieci neuronowe uczą się mapowania, które polega na tym, że każdemu połączeniu między dwiema jednostkami nadawana jest waga (wartość liczbowa wysłaną z jednego neuronu do drugiego), który moduluje sygnał. Poprzez osłabienie lub wzmocnienie wartości połączenia można dopasować przepływ sygnału między poszczególnymi neuronami, a zmiany poszczególnych wartości pozwalają uzyskać ogólne przykłady mapowania sieci od wejścia do wyjścia. W SSN zastosowano różne techniki uczenia się, które różnią się stopniem informacji zwrotnej oraz poziomem samoorganizacji. Podczas nadzorowanego uczenia z wartościami wejściowymi są dostarczane dane dotyczące prawidłowych wyjść w każdym kroku. Sieć jest informowana jakich wejść użyć i jakie sygnały wyjściowe wytworzyć, ale koordynacja przepływu sygnału między wejściem a wyjściem, zależy od samoorganizacji sieci³¹. Wewnętrzne reprezentacje jak wagi i lub ukryte aktywacje jednostek można uznać za znaki, które są prywatne dla sieci i często nieprzejryste dla obserwatorów zewnętrznych. Zatem, w przeciwieństwie do tradycyjnej sztucznej inteligencji, nie istnieją symboliczne reprezentacje, które odzwierciedlają

31 R. Tadeusiewicz, *Sieci neuronowe*, Akademicka Oficyna Wydawnicza 1993.

wcześniej daną rzeczywistość zewnętrzną. W zamian istnieje tam samoorganizacja adaptacyjnego przepływu sygnałów pomiędzy prostymi jednostkami przetwarzania w interakcji ze środowiskiem, co jest zgodne z poglądem dotyczącym istoty powstawania interaktywnych reprezentacji³². Koneksjonizm zaproponował zatem alternatywne podejście do badania wewnętrznych reprezentacji i sposobu użycia znaków. Reprezentacje o równoległym i rozproszonym charakterze mogą być tworzone w wyniku w interakcji z otoczeniem. Połączenia są adaptacją powstałą w wyniku nabierania doświadczeń.

Koneksjonizm nie był w stanie rozwiązać problemów związanych z interakcją ze światem rzeczywistym. Pomimo pewnego postępu w wykorzystaniu sieci neuronowych, nie był to prawdziwy przełom w przechwytywaniu specjalistycznej wiedzy, budowaniu systemów językowych lub w postrzeganiu środowiska. Skupiając się na próbie odwzorowywania działania biologicznego mózgu, zupełnie pominięto problem działania umysłu oraz ciała. Mimo przełomu, który pojawił się dzięki powstaniu SSN, problem środowiska nadal sprowadzał się do wartości wejściowych i wyjściowych dostarczanych i interpretowanych przez ludzkich projektantów i obserwatorów³³. Sztuczne sieci nie są osadzone w sztucznym systemie analogicznie do sieci neuronowych w żywym organizmie. Badanie zjawisk poznawczych zachodzących w takiej sieci to badanie interakcji między organizmem oddzielonym od świata³⁴. Tworzenie połączeń między wejściami, wyjściami i reprezentacjami wewnętrznymi oraz obiektami, które mają reprezentować, zostaje pozostawione opinii obserwatora. Istnieją jako odwzorowanie obiektów i własnych idei projektanta systemu, nie istnieją jednak autonomiczne połączenia między systemem sztucznej inteligencji a światem, który ma reprezentować. W związku z tym taki system nie posiada własnej semantyki.

Sztuczna sieć neuronowa (SSN) jest próbą symulowania działania sieci neuronów obecnych w mózgu. Dzięki ich zastosowaniu sztuczny system może się uczyć i podejmować decyzje w sposób podobny do ludzkiego. Z perspektywy biosemiotyki

32 M.H. Bickhard, L. Terveen, *Foundational issues in artificial intelligence and cognitive science: Impasse and solution*, Elsevier 1996, t. 109, s. 229,330.

33 Por. A. Clark, *Being there: Putting brain, body, and world together again*, MIT Press 1997, ss. 58–59; G. Dorffner, *Radical connectionism-a neural bottom-up approach to AI*, „Neural Networks and a New Artificial Intelligence”, 1997, ss. 93–132; G. Lakoff, *Smolensky, semantics, and the sensorimotor system*, „Behavioral and Brain Sciences”, t. 11, nr 1, 1988, ss. 39–40.

34 T. Ziemke, *The construction of 'reality' in the robot: Constructivist perspectives on situated artificial intelligence and adaptive robotics*, „Foundations of Science”, t. 6, nr 1–3, 2001, ss. 163–233.

jednak, takie badania dotyczą jednak tylko wewnętrznej części koła funkcyjnego podmiotu, gdzie jednostki wejściowe mogą być porównane z receptorami a jednostki wyjściowe z efektorami. Korzystanie z sieci neuronowej, w której aktywacja jest przekazywana tylko w jednym kierunku, odwzorowanie z wejścia na wyjście zawsze będzie takie samo, ponieważ sieć nauczyła się wykonywać zadanie i nie modyfikuje długości wag połączeń. Te same dane wejściowe będą zawsze odwzorowywane na tych samych wyjściach.

Zastosowanie w sieciach neuronowych informacji zwrotnych używających powtarzających się połączeń, powoduje powstanie nietrywialnej rekurencyjnej sieci neuronowej (RSN)³⁵. Możemy tam rozróżnić informacje zwrotne pierwszego rzędu jak ponowne wykorzystanie wcześniejszych wartości aktywacji neuronalnej jako dodatkowych wejść oraz informacje zwrotne wyższego rzędu, jak dynamiczna adaptacja i modulacja wagi połączeń. W obu przypadkach odwzorowanie z wejścia na wyjście będzie różne w zależności od stanu wewnętrznego sieci, a zatem sztuczny system, korzystając z nabytych doświadczeń, może dawać inne wyniki. Sieć zaczyna interpretować dane wejściowe zgodnie ze swoim wewnętrznym stanem³⁶. W świetle teorii biosemiotycznych, sieć dynamicznie tworzy historię doświadczeń, która pozwala jej subiektywnie postrzegać bodźce³⁷. Krąg funkcjonalny realizowany przez sieć rekurencyjną i przypisywanie znaczenia bodźcom, zmieniają się w czasie.

RSN odgrywają ważną rolę w badaniu i modelowaniu reprezentacji poznawczych. Stanowią długoterminową reprezentację doświadczeń w zakresie wagi połączeń oraz krótkoterminową reprezentację bieżącego kontekstu lub bezpośrednią przeszłość w postaci wewnętrznych informacji zwrotnych. RSN, podobnie jak realne systemy nerwowe, reagują na bodźce środowiskowe zależnie od aktualnego stanu systemu, nie będąc determinowane wyłącznie sygnałem wejściowym³⁸. Podobnie dzieje się w mózgach żywych organizmów, których neurony charakteryzują się silnymi sprzężeniami zwrotnymi. W takich sieciach

35 J.L. Elman, *Distributed representations, simple recurrent networks, and grammatical structure*, „Machine Learning”, t. 7, nr 2–3, 1991, ss. 195–225.

36 J. Elman, *Generalization, simple recurrent networks, and the emergence of structure*, [w:] Mahwah, NJ: Lawrence Erlbaum Associates 1998, s. 6.

37 Patrz T. Ziemke, *The construction of 'reality' in the robot: Constructivist perspectives on situated artificial intelligence and adaptive robotics...*, *op. cit.*

38 M.F. Peschl, A. Riegler, *Does Representation Need Reality?*, [w:] *Understanding Representation in the Cognitive Sciences: Does Representation Need Reality?*, A. Riegler, M. Peschl, A. von Stein (red.), Boston, MA 1999, ss. 9–17.

wewnętrzne struktury odpowiadają za powstawanie zachowania, które jest wyzwalane i modulowane przez środowisko i określane przez stan wewnętrzny. Wpływ na to mają adaptacyjne procesy filogenetyczne i ontogenetyczne, zmieniające architekturę sieci w trakcie uczenia się. Oznacza to, że ukryta warstwa sieci jest już w stanie aktywacji, gdy odbierane są dane wejściowe. Dlatego też dane otrzymywane na wejściu zależą od tej początkowej aktywacji. Powtarzające się działania zapewniają agentowi użycie mechanizmów do ponownego i dynamicznego dostosowania zachowania, a tym samym procedury przypisywania znaczenia przychodzącym bodźcom³⁹. Mimo tego, istnieje wiele ważnych aspektów, w jakie sieci neuronowe różnią się od ucieleśnionych mózgów. Główną różnicą jest brak osadzenia w ciele i środowisku. Bezcielesność sieci neuronowych przemawia na ich niekorzyść, gdy rozważamy sposoby radzenia sobie z problemami w rzeczywistym świecie. Żywe organizmy doskonale radzą sobie z odniesieniem takim jak postrzeganie kierunku, orientacja w przestrzeni, różnica między tym co wewnętrzne a zewnętrzne. Dlatego też muszą być wbudowane w fizyczne systemy, co pozwoliłyby budować własne sądy na temat świata.

1.2.3 Zagadnienie ugruntowania symboli w systemach SI

Ignorowanie problemów ucieleśnienia przez czysto obliczeniowe teorie umysłu nie spowodowało jednak ich zniknięcia, zatem teoretycy próbowali rozwiązać dylematy SI starając się połączyć systemy symboliczne ze środowiskiem. Steven Harnad przekonywał, że trudność leży w ugruntowaniu symboli, należy więc użyć hybrydowej kombinacji najmocniejszych stron teorii symbolicznych i koneksjonistycznych⁴⁰. Była to próba utrzymania obliczeniowej teorii poznania, wzmocnionej dzięki systemom robotycznym, które mogłyby połączyć wewnętrzne reprezentacje systemu ze światem, który mają reprezentować. Rozwiązanie problemu ugruntowania symboli stało się jednym z najważniejszych teoretycznych zadań w kontekście prac nad sztuczną inteligencją od lat

39 T. Ziemke, N.E. Sharkey, *A stroll through the worlds of robots and animals: Applying Jakob von Uexkull's theory of meaning to adaptive robots and artificial life*, „Semiotica”, t. 134, nr 1/4, 2001, ss. 701–746.

40 S. Harnad, *The symbol grounding problem*, „Physica D: Nonlinear Phenomena”, t. 42, nr 1, 1990, ss. 335–346.

90. XX wieku⁴¹. Starano się dokładnie określić, w jaki sposób odwzorować treść reprezentatywnego elementu w systemie symbolicznym oraz w jaki sposób dane sensoryczne wpływają na całokształt symbolicznych reprezentacji. Takie systemy muszą odnosić się bezpośrednio do świata rzeczywistego, by rozwijać umiejętności poznawcze a dzięki temu wykonywać inteligentne zadania. Tworzenie formalnego systemu symboli jest wynikiem podwójnej translacji danych: najpierw świat zewnętrznych symboli jest przekładany przez programistę na jego własne reprezentacje a później zostaje zamieniony na język formalny, implementowany w maszynie. Symbole sztucznego systemu nie są więc bezpośrednio związane ze światem, lecz pośrednio, poprzez umysł i język człowieka. Zwykle symbole tworzące system symboliczny nie przypominają ani nie są przyczynowo związane z ich odpowiednimi znaczeniami. Są tylko konwencją uzgodnioną przez użytkowników. Kluczowym następstwem tego jest, że taki system jest pozbawiany treści semantycznej zależnej od kontekstu lub relacji, w jakich symbol jest postrzegany przez interpretującego. Żywe inteligentne organizmy są ugruntowane bezpośrednio w świecie i tym samym różnią się znacząco od komputerów. Nie potrzebują tłumacza, by poznawać przedmioty i rozumieć relacje. Znaczenie symboli jest ugruntowane w ich zdolności do interakcji z realnym światem obiektów, zdarzeń i stanów rzeczy a symbole przez nie używane, są systematycznie interpretowane jako odnoszące się do czegoś. Czyli są interpretowane jako posiadające otrzymane tu i teraz znaczenie.

Wszystkie podejmowane po 1990 roku próby miały na celu ugruntowanie symboli poprzez zdolności sensomotoryczne sztucznych systemów wchodzących w interakcje z otoczeniem. Starano się opracować dane uzyskane z doświadczeń czuciowo-ruchowych oraz wynikających z nich reprezentacji powstających w procesie generowania symboli. Strategia polegała na próbach uchwycenia istotnych cech wspólnych dla zbiorów danych, wyodrębnienie ich z zestawów danych, rozpoznania abstrakcji jako treści kategoryalnych i pojęciowych reprezentacji a ostatecznie na użyciu tych reprezentacji w celu ugruntowania symboli. Niestety, wszystkie strategie były mocno osadzone w

41 A. Cangelosi, A. Greco, S. Harnad, *Symbol Grounding and the Symbolic Theft Hypothesis*, [w:] *Simulating the Evolution of Language*, A. Cangelosi, D. Parisi (red.), Springer, London 2002, ; S. Harnad, *Computation is just interpretable symbol manipulation; cognition isn't*, „Minds and Machines”, t. 4, nr 4, 1994, ss. 379–390; K. MacDorman, *Cognitive Robotics. Grounding symbols through sensorimotor integration.*, „Journal of the Robotics Society of Japan”, t. 17, nr 1, 1999, ss. 20–24; L. Steels, *The symbol grounding problem has been solved. so what's next*, „Symbols and Embodiment: Debates on Meaning and Cognition”, 2008, ss. 223–244.

podejściu syntaktycznym, więc nie przyniosły żadnego rozwiązania dla semantyki⁴². Jeśli systemy symboli manipulują symbolami, symbole te powinny reprezentować coś dla samego systemu, a nie tylko dla zewnętrznego obserwatora. Problem ugruntowania symboli jest obecny w systemach naturalnych i wciąż stanowi wyzwanie, nieznajdujące rozwiązania w sztucznych systemach.

Pod koniec lat 80. XX wieku wiele badań zostało porzuconych. Duże rozczarowanie przyniosły badania dotyczące komputerowego rozpoznawania obrazów i przetwarzania mowy. Okazało się, że możliwości człowieka dotyczące rozpoznawania, określania odległości i orientacji przestrzennej, znacząco przekraczają możliwości komputerów. Pomimo ogromnych inwestycji w systemy mowy, ich wydajność, dokładność i praktyczna użyteczność pozostawała poniżej oczekiwań. Starania, by modelować ludzką percepcję wzrokową i język poprzez klasyczne algorytmy obliczeniowe, okazały się bezużyteczne. Ponadto dowiedziono, że próba konceptualizacji ludzkich ekspertów jako maszyn symbolicznych do przetwarzania danych, nie była do końca trafna i nie prowadziła do powstania systemów zastępujących mentalną pracę człowieka⁴³. Zastosowanie wspomnianych wytycznych do budowania systemów robotycznych, wykreowało obraz robota jako poruszającego się układu z wbudowanym zaawansowanym komputerem z oprogramowaniem sztucznej inteligencji, na przykład ogromnym systemem eksperckim.

Roboty były projektowane i programowane bezpośrednio do wykonywania zadań. To często prowadziło do obliczeniowo kosztownych rozwiązań, które nie tylko nie skutkowały powstawaniem naturalnych zachowań, ale też procesy w systemie przebiegały zbyt wolno, by osiągnąć satysfakcjonujące wykonanie konkretnej funkcji (np. biegania, chwytania poruszających się przedmiotów). Cechą robotyki podejścia GOF AI był tradycyjny sposób projektowania przy użyciu podejścia odgórnego (*top-down*), które polegało na dekompozycji systemu w celu manipulacji jego elementami składowymi. W tym redukcjonistycznym podejściu już na początku szczegółowo planowano budowę systemu, określając podsystemy pierwszego poziomu. Poprawa wydajności odbywała się poprzez udoskonalenie poszczególnych modułów funkcjonalnych. Każdy podsystem był

42 M. Taddeo, L. Floridi, *Solving the symbol grounding problem: a critical review of fifteen years of research*, „Journal of Experimental & Theoretical Artificial Intelligence”, t. 17, nr 4, 2005, ss. 419–445.

43 P. McCorduck, *Machines who think: a personal inquiry into the history and prospects of artificial intelligence*, Natick, Mass. 2004, ss. 205–208.

bardzo szczegółowo przemyślany i wykonany, przez co można było zredukować działanie całego systemu do elementów podstawowych. Było to trudne ze względu na nieelastyczność różnych części, a zmiany w jednym module negatywnie wpływały na działanie innych, więc cały projekt musiał zostać ponownie rozważony na każdym etapie zmiany projektu. W podejściu *top-down* bardzo łatwo było manipulować modelem, lecz nie dawało ono możliwości weryfikacji ani wyjaśnienia mechanizmów, które model miał za zadanie symulować (np. procesów poznawczych). Łączenie wielu modułów pozwalało uzyskać zachowanie systemu będące wyłącznie odzwierciedleniem wizji projektanta, zatem u podstaw ograniczało możliwość wykształcenia zachowań inteligentnych.

Wykształcenie i przekonania filozoficzne prekursorów i twórców SI wykreowały silny paradygmat w tej dziedzinie oraz wywarły nacisk na badania. W tym ujęciu przejawy inteligencji takie jak powstawanie wiedzy, racjonalne wybory i rozwiązywanie problemów miały być efektem symbolicznych manipulacji. Dowodem na nieadekwatność założeń GOFAI jest fakt, że nieliczne zbudowane w zgodzie z metodami tego paradygmatu systemy często wykazywały braki, nieelastyczność i problemy z działaniem w czasie rzeczywistym. Systemy te nie były w stanie wziąć pod uwagę, że obiekty świata zewnętrznego zmieniają się, a system inteligentny powinien swoje działanie do tych zmian dopasować⁴⁴. Naturalny fenomen inteligencji został zredukowany do formalnego lub funkcjonalnego modelu. Stał się kolejnym programowalnym systemem mechanicznym (w sensie Turinga), który realizuje funkcje oraz działania same w sobie tak jak robi to maszyna Turinga jako komputer. Inteligencja była postrzegana jako abstrakcyjny ruch, forma logiczna wyabstrahowana z podłoża materialnego systemu, w którym się manifestuje. W opisywanym podejściu do SI wciąż opierano się na modelu komputacyjnym, który zakładał już na samym początku rozwoju tej dyscypliny, że umysł ludzki jest analogiem *software'u* działającego na fizycznym mózgu przypominającym *hardware*. Założenie to okazało się błędne w świetle faktu, że możliwości obliczeniowe szybko przekroczyły możliwości ludzkie, natomiast efekty działania systemów SI były dalekie od oczekiwań. W sztucznych systemach umiejętności takie jak logiczne wnioskowanie, skomplikowane obliczenia, wyszukiwanie danych są dużo szybsze i dokładniejsze, niż ludzkie. Klasyczne podejście GOFAI nie przyczyniło się znacząco do pogłębienia naszego zrozumienia percepcji, lokomocji, manipulacji, codziennej mowy i konwersacji, interakcji społecznych, zdrowego rozsądku, emocji etc.

44 P. Hayes, *What the frame problem is and isn't*, 1987.

Sposób w jaki postrzegano inteligencję człowieka postrzeganą obliczeniowo stworzył problem testowania inteligencji jako pewnej obiektywnej wartości. Niezależnie od tego, czy inteligencja dotyczy człowieka, zwierzęcia czy sztucznego systemu, jej badacze są przekonani, że może zostać wyznaczona podobnie jak fizyczne wartości: siła, waga, energia. Taka perspektywa doprowadziła do nadmiernego wartościowania testów mierzących iloraz inteligencji (testy IQ). Zwolennicy tych badań twierdzą, że iloraz inteligencji jest bezpośrednio związany z różnymi miarami sukcesu lub porażki w życiu⁴⁵. Zakładają, że zdolność człowieka do powtarzania długiej sekwencji cyfr, znajdowania ogólników dla szeregu pojęć, formowania wielościanów, budowania kształtów z części, porządkowania obrazów w określonym czasie i ciągu logicznym koreluje dobrze z innymi zdolnościami intelektualnymi. Niektóre problemy związane z testowaniem inteligencji są rzeczywistymi wyzwaniem dla programu sztucznej inteligencji, lecz przede wszystkim mierzą ilość informacji jaką człowiek może przetworzyć w określonym czasie. Taka liczba operacji jest trywialna nawet dla bardzo prostych komputerów, które osiągnęły wysokie wyniki w testach, lecz ich wysoki poziom IQ wcale nie wiąże się z ich zdolnością do odnajdywania się w swoim środowisku i do rozwiązywania problemów, które tam się pojawiają. Dodatkowo należy zauważyć, że inteligencja nie może być mierzona ani tym bardziej porównywana międzykulturowo. Wszelkie kultury wykazują inteligentne zachowanie odpowiednie dla różnych dziedzin życia, mimo różnic inteligencji poszczególnych członków⁴⁶. Mierzenie poziomu inteligencji wydaje się być zasadne tylko w obrębie konkretnej kultury, ponieważ ukazuje działanie badanych w specyficznej domenie i jest ściśle związane z obserwatorem lub kreatorem testu. Jeśli ten należy do innej kultury, nie będzie mógł zmierzyć potencjału badanych do inteligentnego działania w innych dziedzinach lub ich ogólnej inteligencji. Największy problem testów IQ to założenie, że możliwe jest rozróżnianie poziomów inteligencji, które okazuje się błędne, gdy zauważymy, że inteligentne zachowanie może dotyczyć różnych, często bardzo odmiennych, dziedzin. Wydaje się zatem, że jeszcze trudniej jest wnioskować o inteligencji bądź jej braku w sztucznym systemie, który nie posiada ani historycznych, ani

45 A.L. Duckworth i in., *Role of test motivation in intelligence testing*, „Proceedings of the National Academy of Sciences”, t. 108, nr 19, 2011, s. 7716–7720.

46 H.R. Maturana, G.D. Guilloff, *The quest for the intelligence of intelligence*, „Journal of Social and Biological Structures”, t. 3, nr 2, 1980, ss. 135–148.

kulturowych, ani społecznych uwarunkowań, do których można by się odnieść nawet dzięki daleko posuniętej analogii.

Humberto Maturana i Gloria Guilloff twierdzili, że należy wziąć pod uwagę to, jakiego zachowania oczekują ci, którzy chcą badać i mierzyć inteligencję. Inteligentne działanie musi być wynikiem historycznych powiązań strukturalnych ze środowiskiem i lub innymi organizmami. Każde zachowanie, które odnosi sukces w ramach sprzężenia strukturalnego musi zostać uznane za zachowanie inteligentne⁴⁷. Nie sposób się nie zgodzić z takim oświadczeniem. Tym bardziej, że pozwala ono spojrzeć na problem inteligencji niejako z zewnątrz. Inteligencja w takiej perspektywie, nie jest kompleksem logicznych zdolności umysłu, lecz obserwowalnym fenomenem, który jest zawsze związany jest z fizycznymi i społecznymi zasadami wymuszonymi przez specyficzne środowisko. Pomimo, że określenie środowisko często pojawia się w badaniach nad inteligencją, jego wpływ nie jest zazwyczaj brany pod uwagę. Tymczasem inteligentne działania nie są tylko wynikiem pewnych umiejętności umysłowych organizmu, lecz mają zdecydowanie teleologiczne podstawy. Zachowania inteligentne mają za zadanie doprowadzić do określonego celu. Organizm musi zapewnić sobie przetrwanie i w tym celu dostosowuje się do środowiska, w którym ma funkcjonować.

1.3 Ucieleśniona sztuczna inteligencja – programy w świecie rzeczywistym

Rola środowiska i ucieleśnionego podmiotu w dziedzinie sztucznej inteligencji zaczęła odgrywać ważną rolę w próbach zrozumienia jak kształtuje się inteligentne zachowanie. Zaczęto zdawać sobie sprawę z tego, że działanie podmiotu związane jest z danymi zmysłowymi, za pomocą których uzyskuje niezbędne informacje ze świata, a także wprowadza je w świat. Świat wydaje się w ten sposób nierozłączny z ciałem i jego sensomotorycznymi funkcjami, co stanowi ważny wymóg dla powstania indywidualnej inteligencji podmiotu. Dzięki przyjęciu tej perspektywy w badaniach SI nastąpił przełom. Olbrzymią rolę odegrała tu gałąź badań nad sztuczną inteligencją – robotyka. Trudności programu GOF AI doprowadziły do realizacji innego podejścia, uwzględniającego interakcję system-środowisko. W 1980 roku Rodney Brooks stwierdził, że nacisk na

⁴⁷*Ibid.*

logikę, rozwiązywanie problemów i rozumowanie, opierają się na ludzkiej introspekcji, czyli na tym, jak ludzie postrzegają siebie i własne procesy psychiczne, i dlatego prace nad rozwojem sztucznej inteligencji przybierają niewłaściwy obrót⁴⁸. Zamiast tego zaproponował, że należy zapomnieć o przetwarzaniu symboli, reprezentacji wewnętrznej oraz o rozwiązaniach związanych z wysokim poziomem rozumowania, a skupić się raczej na interakcji z realnym światem. „Inteligencja wymaga ciała”⁴⁹ stało się hasłem nowego paradygmatu ucieleśnionej inteligencji. Dzięki tej zmianie orientacji, charakter pytań badawczych również się zmienił: społeczność naukowa zainteresowała się poruszaniem się, manipulacją przedmiotami i przede wszystkim działaniem agenta – poznawczego podmiotu w zmieniającym się świecie⁵⁰. W konsekwencji, wielu badaczy SI na całym świecie rozpoczęło pracę nad robotami. Jednak nawet tego typu badania nie były w stanie automatycznie rozwiązać problemów. Wykonanie większości zadań w przestrzeni świata naturalnego przez roboty autonomiczne było wciąż uważane za niezadawalające. Chodzenie, bieganie, postrzeganie, manipulowanie obiektami, interakcja z innymi przekraczało możliwości obliczeniowe dostępnych komputerów⁵¹.

Koncepcja ucieleśnienia przeniknęła do nauk kognitywnych, w tym do dziedziny sztucznej inteligencji i została związana z koncepcją ucieleśnienia fizycznego, zakładającą, że sztuczny system inteligentny musi posiadać reprezentację w świecie fizycznym, do którego jest podłączony za pomocą zestawu czujników i silników. Ucieleśnienie opiera się na przekonaniu, że poznanie, myśl i inteligencja nie są izolowanymi, oddzielnymi funkcjami umysłowymi, ale codzienną działalnością żywego organizmu – jego „byciem w świecie”. Wzory zachowań rozwinęły się dzięki strukturalnemu sprzężeniu podczas

48 R.A. Brooks, *Intelligence without representation...*, op. cit.

49 R.A. Brooks, *Intelligence without reason*, [w:] *The artificial life route to artificial intelligence: Building embodied, situated agents*, Routledge 1995, ss. 25–81.

50 W literaturze anglojęzycznej funkcjonuje pojęcie „agent” określające robota, program komputerowy lub czynnik jakiegoś procesu. Pojęcie to ma za zadanie przedstawić sztuczny system odwzorowujący działanie żywych organizmów w badaniach nad inteligentnymi autonomicznymi agentami (*Intelligent Agent Technology*). W języku polskim słowo to ma szerokie i odmienne zastosowanie, trudno jest jednak znaleźć określenie, które oddawałoby znaczenie tego pojęcia. W tej pracy oznacza ono podmiot poznawczy (zarówno naturalny jak i sztuczny), działający w świecie, umiejący się do komunikować, zdolny do postrzegania i działania w swoim środowisku, podejmujący autonomiczne wybory w celu osiągnięcia założonych skutków działania.

51 H. Moravec, *When will computer hardware match the human brain*, „Journal of Evolution and Technology”, t. 1, nr 1, 1998, s. 10.

interakcji w środowisku. Ostatecznie uznano, że myślenie nie jest oderwaną refleksją, ale częścią podstawowego stosunku do świata - jednym z ciągłych, celowych działań. Wiedza zaś nie składa się z reprezentacji obiektywnych i niezależnych przedmiotów, lecz powstaje w trakcie aktywnego postrzegania i działania w świecie. Jest nieustannie budowana i reorganizowana, a przy tym wpływa na celowe działania. Wiedza opiera się na przeszłych doświadczeniach i historii strukturalnych sprzężeń. Naukowcy zgodzili się, że inteligencja zawsze przejawia się w zachowaniu - dlatego należy zrozumieć przede wszystkim to zachowanie. Inteligencja jest również wynikiem interakcji podmiotu ze środowiskiem. Jej badanie musi uwzględniać podstawy biologicznego postrzegania i działania w świecie, które są konieczne a często również wystarczające (jak to ma miejsce u prostych organizmów) do inteligentnego działania.

Jedną z najbardziej elementarnych zdolności każdego stworzenia jest kategoryzacja: umiejętność rozróżniania rzeczy w świecie rzeczywistym. Nawet proste organizmy odróżniają żywność i trucizny, sytuacje niebezpieczne od bezpiecznych, przedstawicieli tego samego gatunku od innych. Dlatego też sztuczne systemy, które nie są w stanie dokonać podstawowych rozróżnień, nie są zbyt użyteczne. Tworzenie takich kategorii jest bardzo bezpośrednio zdeterminowane przez ucieleśnienie, czyli morfologię i właściwości materialne ciała. Morfologia obejmuje kształt ciała, rodzaje kończyn i miejsca ich zamocowania, rodzaje czujników odpowiedzialnych za wrażenia sensoryczne i miejsca, w których znajdują się na ciele. Przez właściwości materiału rozumiemy, na przykład, odkształcalność końców palców i skóry lub elastyczność układu mięśniowo-ścięgowego. Podczas interakcji ze światem rzeczywistym ciało jest stymulowane w bardzo szczególny sposób, a stymulacja ta stanowi w pewnym sensie surowiec do pracy z mózgiem. Ten surowiec może zostać użyty do stworzenia kategorii, które opisują środowisko w naturalny sposób. Możliwe jest oczywiście konstruowanie bardzo abstrakcyjnych kategorii, ale nawet te są pod wpływem podstawowych pojęć, które można stworzyć. George Lakoff i Rafael Núñez argumentowali w książce *Where Mathematics Come From*⁵², że nawet abstrakcyjne pojęcia matematyczne, jak wyobrażenie liczby rzeczywistej lub zbioru mają swój początek w konkretnym ucieleśnieniu i nie można ich było skonstruować inaczej. Koncepcje matematyczne, według Lakoffa i Núñeza, opierają się na metaforach, które są z kolei oparte na ucieleśnieniu. Zatem nie tylko kategoryzacja

52 G. Lakoff, R. Núñez, *Where Mathematics Comes From: How the Embodied Mind Brings Mathematics into Being*, New York 2000.

jest ugruntowana w cielesności, ale także ogólne poznanie, w tym poznanie przestrzenne i społeczne, rozwiązywanie problemów i rozumowanie oraz język naturalny.

Chodzenie i poruszanie się stały się ważnymi obszarami badawczymi związanymi z niskopoziomową inteligencją czuciowo-ruchową. Była to olbrzymia zmiana w porównaniu z programowaniem gry w szachy, dowodzeniem twierdzeń i abstrakcyjnym rozwiązywaniem problemów. Innymi tematami, które zaczęto badać, były zachowania orientacyjne np.: znajdowanie drogi w tylko częściowo znanych i zmieniających się środowiskach i unikanie przeszkód. Również w tym miejscu badań znaczenie terminu sztuczna inteligencja zaczęło się zmieniać, a raczej przyjmować dwa znaczenia. Pierwsze znaczenie GOFAI obejmowało tradycyjne podejście algorytmiczne, drugie oznaczało podejście ucieleśnione EAI – *Embodied Artificial Intelligence*, zwane również *Novelle AI*. Paradygmat zakładał trzy cele: zrozumienie systemów biologicznych, wyodrębnienie abstrakcyjnych ogólnych zasad inteligentnego zachowania oraz zastosowanie tej wiedzy do budowy sztucznych systemów. W rezultacie współczesne, ucieleśnione podejście zaczęło przenosić się z laboratoriów informatycznych do laboratoriów robotyki i inżynierii lub biologii⁵³. Naukowcy zdali sobie sprawę, że naturalne systemy neuronowe są wyjątkowo skuteczne w kontrolowaniu interakcji z rzeczywistym światem. W ujęciu EAI nastąpiło odnowienie i znacznie silniejsze zainteresowanie neuronaukami, ponieważ dawały nadzieję na zrozumienie w jaki sposób organizmy poruszają się i manipulują obiektami dzięki kontrolowaniu ich przez naturalne sieci neuronowe. Ponadto ich działania przebiegały bardzo sprawnie przy niskim zużyciu energii. Uznano, że te rodzaje zachowań można osiągnąć dzięki wykorzystaniu dynamicznych właściwości sieci neuronowych⁵⁴. Było to całkowicie sprzeczne z tradycyjnym podejściem sztucznej inteligencji, w którym stosowano głównie na statyczne sieci sprzężenia zwrotnego.

Wraz z tą zmianą orientacji, charakter pytań badawczych również zaczął się zmieniać. Rodney Brooks i jego grupa w Laboratorium sztucznej inteligencji w *Massachusetts Institute of Technology* opracowali praktyczną krytykę paradygmatu GOFAI, czyli założenia, że inteligentne zadania muszą być wdrażane przez proces rozumowania w modelu obliczeń symbolicznych⁵⁵. Nowe techniki modelowania i zasady budowania systemów stopniowo zyskiwały wpływy i stały się znane jako

53 M.E. Martinez, *Future Bright: A Transforming Vision of Human Intelligence*, OUP USA 2013.

54 D.E. Rumelhart, J.L. McClelland, *Psychological and Biological Models*, MIT Press 1986, ss. 543–544.

55 R.A. Brooks, *Intelligence without representation...*, *op. cit.*

„badanie systemów reaktywnych” lub „badanie agentów”. Równoległe do tego ruchu i głęboko zainspirowane nim, powstawały nowe pojęcia procesów poznawczych związane z ucieleśnieniami, usytuowanymi układami⁵⁶.

Jednym ze źródeł, z których czerpano założenia nowego paradygmatu, była teoria *Umweltu* Jakoba von Uexküll⁵⁷, która została użyta w celu wykazania, że istnieje ściśle, dynamiczne połączenie ciała żywych istot z ich doświadczanym światem oraz że świat postrzegany przez zwierzę jest całkowicie odmienny od świata postrzeganego przez zoologa. Pozwoliło to uświadomić fakt, że sztuczny system będzie istniał w „postrzeganym” świecie, który różni się znacznie od świata projektanta, a mechanizmy intelektualne sztucznej inteligencji mogą być odmienne niż te, obserwowane u żywych organizmów. Wielu badaczy również podkreślało, że systemy poznawcze są nie tylko strukturalnie sprzężone ze swoim otoczeniem w teraźniejszości, ale że ich realizacja jest w rzeczywistości rezultatem lub odbiciem historii interakcji między agentem a środowiskiem, a w wielu przypadkach koadaptacji. Twórcy SI uznali, że systemy są ucieleśnione, jeśli są strukturalnie sprzężone ze swoim otoczeniem. Przyjęli również, że to niekoniecznie wymaga biologicznego ciała. Trafnie podsumował to Stan Franklin, twierdząc, iż systemy oprogramowania bez ciała (w sensie fizycznym) mogą być inteligentne, lecz muszą być ucieleśnione w usytuowanym poczuciu bycia autonomicznymi środkami strukturalnie powiązanymi z ich otoczeniem⁵⁸. Koncepcja sprzężenia strukturalnego wywodzi się z prac Humberto Maturany i Francesco Vareli na temat biologii poznania⁵⁹ i można ją podsumować następująco: system jest ucieleśniony w środowisku, jeśli między nimi istnieją połączenia, poprzez które mogą na siebie wzajemnie oddziaływać. Najbardziej znanym opisem ucieleśnienia jest jednak praca Georga Lakoffa i

56 por. J. Engelkamp, *Organisation and recall in verbal tasks and in subject-performed tasks*, „European Journal of Cognitive Psychology”, t. 8, nr 3, 1996, s. 257–274; A. Koriat, S. Pearlman-Avni, *Memory organization of action events and its relationship to memory performance*, „Journal of Experimental Psychology: General”, t. 132, nr 3, 2003, s. 435; M.C. Steffens, A. Buchner, K.F. Wender, *Quite ordinary retrieval cues may determine free recall of actions*, „Journal of Memory and Language”, t. 48, nr 2, 2003, s. 399–415.

57 R.A. Brooks, *Intelligence without representation...*, *op. cit.*

58 S. Franklin, *Autonomous agents as embodied AI*, „Cybernetics & Systems”, t. 28, nr 6, 1997, ss. 499–520.

59 H.R. Maturana, F.J. Varela, *The tree of knowledge: The biological roots of human understanding*, New Science Library/Shambhala Publications 1987.

Marka Johnsona, którzy twierdzili, że abstrakcyjne ludzkie pojęcia są oparte na metaforach opartych na cielesnym doświadczeniu i aktywności⁶⁰.

Brooks, jeden z pierwszych promotorów programu ucieleśnionej inteligencji, zaczął konstruować proste mobilne systemy, których lokomocja była inspirowana poruszaniem się owadów. Zaangażowanie w badania robotów w miejsce programów komputerowych pozwoliło na prowadzenie szerszych badań związanych z nowym paradygmatem. Był to kluczowy aspekt prac, które skupiały się na rzeczywistych fizycznych systemach. Komputery i roboty znacząco różnią się od siebie. Komputery mają w większości bardzo ograniczone rodzaje interakcji ze światem zewnętrznym: wejście odbywa się z reguły za pomocą klawiatury lub kliknięcia myszą, a wyjście odbywa się za pośrednictwem ekranu. Innymi słowy, możliwość komunikacji z otoczeniem jest bardzo słaba i ograniczona. Komputery również postępują według jasno zdefiniowanego schematu przetwarzania danych wejściowych, który ukształtowało myślenie o systemach inteligentnych i stało się wiodącą metaforą klasycznego podejścia kognitywistycznego. Natomiast roboty mają znacznie szerszy repertuar zmysłowo-motoryczny, który umożliwia ściśle sprzężenie ze światem zewnętrznym, a metafora komputerowa przetwarzania danych wejściowych nie może być już stosowana bezpośrednio.

W przekonaniu Brooksa dwa kamienie węgielne nowego podejścia do SI to usytuowanie i ucieleśnienie⁶¹, które są zgodne z antymentalistycznymi podstawami nowej robotyki. Usytuowanie oznacza, że robot jest fizycznie osadzony w świecie i dzięki temu świat fizyczny wpływa bezpośrednio na jego zachowanie. Taki system może percypować świat, zamiast pozostawać przy jego abstrakcyjnej reprezentacji. Ucieleśnienie zakłada posiadanie przez sztuczny system fizycznego ciała, co pozwala na bezpośrednie doświadczanie świata. Działania stają się częścią dynamicznej interakcji ze światem i otrzymują natychmiastowe informacje zwrotne na temat własnych wrażeń, co umożliwia przyczynowe postrzeganie ciała⁶². Połączenie usytuowania i ucieleśnienia pociąga za sobą immanentność świata, która uzasadnia zachowanie robotycznych systemów, których inteligencja pojawia się dzięki procesom oddolnym, dzięki dynamice fizycznego ugruntowania. Dla Brooksa w ten sposób zostaje spełnione zobowiązanie realizmu, który nakazuje budować kompletne inteligentne systemy, które prawdziwie percypują i działają

60 G. Lakoff, M. Johnson, *The metaphorical structure of the human conceptual system*, „Cognitive science”, t. 4, nr 2, 1980, ss. 195–208.

61 R.A. Brooks, *New approaches to robotics*, „Science”, t. 253, nr 5025, 1991, ss. 1227–1232.

62 *Ibid.*

w realnym świecie⁶³. Trudność z budową inteligentnych maszyn ma związek z odwzorowaniem procesów, które wiążą się z podstawowymi działaniami organizmu, takimi jak widzenie lub chodzenie. Nie są one wynikiem świadomej introspekcji i zamierzonej dekompozycji zadań, lecz „po prostu się dzieją”⁶⁴. Potencjał żywej istoty do poruszania się nie należy do poziomu epistemicznego podmiotu, posiadającego zdolność do autorefleksji i świadomości. Ta umiejętność jest sprawnością ciała.

W celu stworzenia inteligentnych maszyn, które poradzić by sobie mogły z powyższymi zdolnościami, Brooks opracował architekturę subsumpcyjną, która zmieniła tradycyjne podejście do sztucznej inteligencji. Była to struktura wielowarstwowa, oparta o działanie modułów behawioralnych w której poszczególne podsystemy rozkładały zachowanie na wiele zadań⁶⁵. Wszystkie moduły były prostymi podsystemami obliczeniowymi, ułożonymi w dyskretnie działające warstwy. Działały równolegle z innymi bez centralnego mechanizmu, który kontrolowałby cały system. Poszczególne moduły mogły odpowiadać na konkretne bodźce płynące ze środowiska. Możliwa była wymiana informacji pomiędzy podsystemami, lecz dane były ubogie i nie wpływały znacząco na wzajemną pracę. Chodziło raczej o podzielenie skomplikowanych zadań między kilka modułów. Umiejętność dekompozycji złożonego zadania miało być istotną cechą inteligentnego i samoorganizującego się układu. Działanie pojedynczych modułów, posiadających szczególne domeny interakcji bądź kompetencje, odpowiadałoby za autonomię. Każdy z podsystemów działał autonomicznie: wykrywał sygnał, przetwarzał go w trakcie obliczeń i powodował należyte działanie. Każdy z podsystemów został połączony z odbiornikiem, od którego otrzymywał bodziec, oraz nadajnikiem odpowiedzialnym za odpowiedź (zachowanie). Interfejs między warstwami pozwalał wyższemu poziomowi zastąpić inne, jeśli pojawiały się konflikty. Wielowarstwowa struktura mogła wzmacniać bądź osłabiać działanie modułów, przede wszystkim jednak konsolidowała efekty ich pracy. Dane były przetworzone i przesyłane kanałami o niskiej przepustowości, bez wpływu na działanie innych⁶⁶. Zatem całościowe zachowanie systemu wynikało z pracy więcej niż jednego podsystemu, mimo że te nie pracowały ani liniowo

63 R.A. Brooks, *Intelligence without representation...*, op. cit.

64 R.A. Brooks, *Flesh and machines: How robots will change us*, Vintage 2003, ss. 36–37.

65 R.A. Brooks, *New approaches to robotics...*, op. cit., s. 67.

66 R.A. Brooks, *A robust layered control system for a mobile robot*, „Robotics and Automation”, t. 2, nr 1, 1986, ss. 14–23.

ani hierarchicznie. W rezultacie miał powstać wydajny i elastyczny system sterujący zachowaniem robota w dynamicznym środowisku.

Brooksowi zależało na zbudowaniu systemu wykazującego emergentną funkcjonalność⁶⁷. Ze wzrostem intensywności interakcji ze środowiskiem system miał przejawiać nowe własności. Inteligentne funkcjonowanie miało być wynikiem działania prostych mechanizmów behawioralnych w interakcji ze środowiskiem. Otoczenie stawałoby się narzędziem wykorzystanym przez działający podmiot. Konstrukcja systemu pozwalała na bezpośrednie sprzężenie percepcji i działania. Maszyna nie tylko odbierała i przetwarzała sygnały, lecz także monitorowała wpływ własnej aktywności na swoje reakcje wewnętrzne oraz na środowisko. System mógł nieskończenie przetwarzać dane, ponieważ jego działanie i percepcyjna aktywność wciąż dostarczała mu nowych informacji o środowisku i o nim samym. W ucieleśnionej SI inteligencja była traktowana jako wartość dodatkowa, a nawet emergentna, która miała się pojawić w wyniku działania sensomotorycznego systemu. Poznanie miało powstać w oparciu o dużą liczbę równoległych, asynchronicznych, luźno powiązanych procesów, sprzężonych z wewnętrzną organizacją pętli sensomotorycznej agenta⁶⁸. Taka pętla byłaby odpowiednikiem semiotycznych procesów. Nowy model poznania wyłaniał się w wyniku odpowiedniej koordynacji czuciowo-ruchowej agenta, który zmieniał się w trakcie efektywnej interakcji z otoczeniem.

Taki model systemu miał pozwolić na rozwiązanie problemu użycia abstrakcyjnych reprezentacji symbolicznych. Reprezentacje zostały zastąpione przez modele zewnętrzne – innymi słowy, rozróżnienie między modelem świata a światem traciło rozróżnienie⁶⁹. Zadaniem systemu nie było tworzenie symbolicznej reprezentacji, lecz aktywne odnoszenie się do realnego świata. Problem reprezentacji odgrywał jedną z ważniejszych ról w badaniach nad tradycyjną sztuczną inteligencją⁷⁰. Tam sukces miało zapewnić

67 M.J. Mataric, R.A. Brooks, *Learning a distributed map representation based on navigation behaviors*, [w:] *Proceedings of 1990 USA Japan Symposium on Flexible Automation*, 1990, ss. 499–506.

68 R.A. Brooks, *A robust layered control system for a mobile robot...*, *op. cit.*

69 R.A. Brooks, *Intelligence without representation...*, *op. cit.*

70 GOFAI starała się oddać symbolicznie podstawowe pojęcia, by potem w systemie formalnym opisać związane z nimi procedury. Tymczasem koncepcje pojęć są szerokie i często arbitralne. Pojawiał się problem przekładu na język formalny nawet prostych zadań. Tymczasem małe dziecko może dokonać właściwej interpretacji i zaproponować plan działania w przypadku polecenia „otwórz pudełko z kredkami”. Jednak w przypadku

całkowite i jednoznaczne przedstawienie stanu świata poprzez wewnętrzne reprezentacje, które miałyby być wykorzystane do symulowania inteligentnych zachowań. Budowa inteligentnego systemu polegała na formalnym przetwarzaniu symboli, które łączyły się ze sobą za pośrednictwem reprezentacji. Tymczasem w ucieleśnionej SI, dzięki sprzężeniu percepcji i działania, zależność od reprezentacji miała zniknąć⁷¹. Dopiero obserwator mógłby przypisać reprezentację, której system sam w sobie nie miał – był tylko zbiorem konkurencyjnych zachowań. Warunkiem dostrzeżenia spójności procesu była wiedza obserwatora posiadającego doświadczenie ugruntowania symboli.

Ucieleśniona sztuczna inteligencja, ze względu na naturę bycia wcielonymi systemami w rzeczywistym świecie fizycznym, musiała zajmować się wieloma kwestiami, które były zupełnie obce w perspektywie kognitywistycznej. Fizyczne narządy sensoryczne działające w realnym świecie są złożone, a ich badanie wymaga współpracy wielu różnych obszarów. Implikacje tej zmiany perspektywy są dalekosiężne i trudno je przecenić. Ucieleśnione i usytuowane podejście do sztucznej inteligencji stało się realną alternatywą dla podejścia tradycyjnego. Zaproponowano konstrukcję systemów, które mogły zachowywać się stabilnie i elastycznie w zmiennych warunkach. EAI była pod wieloma względami dużym sukcesem. Zasady programowania ucieleśnionych agentów pokazały jak ważne jest uwzględnienie powiązanie systemu ze środowiskiem, w którym ma on funkcjonować. Ucieleśniony system sprawniej funkcjonuje, gdy odnosi się do otaczającego świata bądź niszy, w której działa. Pożądane działania docelowe łatwiej jest uzyskać, gdy agent wykorzystuje kontekst środowiskowy. Zachowanie systemu jest nieredukowalne do mechanizmów wewnętrznych, ponieważ funkcjonowanie systemu łączy w sobie trzy aspekty: przyjęcie szczególnej perspektywy (odniesienia), dostosowanie działania do złożoności środowiska oraz zachowanie wynikające z działania wielu mechanizmów. Wszystkie są wspierane przez ucieleśnienie i usytuowanie w świecie rzeczywistym⁷².

W przeciwieństwie do projektowanych w GOF AI systemów, gdzie skupiano się na pojedynczych komponentach, w EAI można uzyskać lepsze zrozumienie inteligencji, badając w jaki powstają zachowania adaptacyjne i jak zmienia się dynamika układu mózg-

systemów SI, programista musiał usunąć większości szczegółów w celu stworzenia prostego opisu w kategoriach pojęć atomowych, takich jak osoba, pudełko, kredki.

71 R.A. Brooks, *Intelligence without representation...*, *op. cit.*

72 R.D. Beer, *The Dynamics of Active Categorical Perception in an Evolved Model Agent*, „Adaptive Behavior”, t. 11, nr 4, 2003, ss. 209–243.

ciało-świat. W rzeczywistości problem kategoryzacji percepcyjnej staje się znacznie uproszczony poprzez nie-obliczeniowe wykorzystanie rzeczywistego świata. Przyjęto również zasadę taniego projektowania robotów, dzięki ograniczeniu się do docelowej niszy i wykorzystaniu tylko tych fizycznych zasobów, które znajdują tam zastosowanie. Przekłada się to na zachowanie pewnej ekologicznej równowagi, ponieważ systemy sensoryczne, motoryczne i kontrolne są dopasowane do złożoności środowiska a budowa morfologiczna agenta zastosowana w określonym celu ułatwia mu kontrolę zadania. Dodatkową wartością jest umiejętność uczenia agenta poprzez otrzymywanie informacji zwrotnej w odniesieniu do jego działań⁷³.

1.4 Fizyczne systemy symboliczne a problemy nadawania znaczenia

Należy zastanowić się w jaki sposób program badań SI powinien zastosować analizę podstaw intencjonalnego i autonomicznego działania oraz procesów powstawania znaczenia. Najlepszym wyjściem byłoby znalezienie współzależności między syntaktyczną i semantyczną strukturą, tak by mogły powstać intencjonalne zdolności systemu. Szczególną rolę odgrywa tu idea symboli oraz systemu symbolicznego, którego zadaniem jest reprezentowanie świata: szczególnych rzeczy, ich właściwości, planów, stanów, zdarzeń i sytuacji, zjawisk, cech percepcyjnych – czyli ogólnie znaków pojawiających się w świecie⁷⁴. Sztuczne systemy inteligentne robota muszą sobie radzić z danymi, które zostają do nich wprowadzone w procesie percepcji. Poznanie jest postrzegane jako proces manipulowania symbolami i przetwarzania informacji. Lecz między formalnym systemem symbolicznym a naturalnym systemem symbolicznym istnieje olbrzymia rozbieżność.

Według Alana Newella hipoteza systemu symbolicznego jest odpowiednia dla każdej aktywności symbolicznej w świecie fizycznym⁷⁵. Wynika z tego, że wszystkie istoty żywe korzystają z tego samego, uniwersalnego systemu symboli, który może zostać zapisany w sposób formalny. Fizyczny system symboliczny opisuje zbiór wielu systemów, które manipulują symbolami świata. Symbole są tu obiektami powiązanymi z sobą w sposób fizyczny. System symboliczny przetwarzający symbole ma pozwolić na

73 R. Pfeifer, G. Gómez, *Interacting with the real world: design principles for intelligent systems*, „Artificial Life and Robotics”, t. 9, nr 1, 2005, ss. 1–6.

74 J. Pelc, *Wstęp do semiotyki*, Warszawa 1984, s. 94.

75 A. Newell, *Physical symbol systems...*, *op. cit.*

inteligentne działanie. Newell i Simon twierdzili, że symbole fizycznego systemu symbolicznego są analogiczne do symboli występujących w świecie żywych organizmów⁷⁶. Newell był przekonany, że badania nad sztuczną inteligencją wniosą ogromny wkład do badań nad symbolami przetwarzanymi przez człowieka⁷⁷.

Krytycy hipotezy Newella zauważają jednak, że fizyczne systemy symboliczne nie używają rzeczywistych symboli, lecz zaledwie sygnałów, które mogą zostać przetworzone przy użyciu narzędzi matematyki. David Touretzky i Dean Pomerleau twierdzą, że symbol w fizycznym systemie symboli ma zbyt szerokie znaczenie i w tym ujęciu nawet proste automaty mogą być uważane za systemy symboliczne⁷⁸. Tymczasem hipoteza fizycznego systemu symbolicznego zakłada, że „X oznacza Y, jeśli X oznacza aspekty Y, tj. jeśli istnieją procesy symboliczne, które mogą przyjąć X jako dane wejściowe i zachowywać się tak, jakby miały dostęp do niektórych aspektów Y”⁷⁹. Jak zauważa w swoim artykule Jarosław Boruszewski pozostaje niewiadomym, w jaki sposób symboliczny proces obliczeniowy miałby być kształtowany przez dane wejściowe już na wejściu systemu. Relacja symbolizowania ma funkcję przechodnią (jeśli X symbolizuje Y oraz Y symbolizuje Z, to X symbolizuje Z) i ona właśnie ma zapewnić odniesienie fizycznego systemu symbolicznego do świata⁸⁰. Relacje pomiędzy systemem symbolicznym a jego odniesieniem przedmiotowym zazwyczaj nie są tak proste i bardzo często są wynikiem wysoce niepośredniego złożenia.

Wydaje się, że ta hipoteza zastosowana do działań symbolicznych człowieka i zwierząt stanowi poważną redukcję, a z pewnością różni się od pojęcia symbolu i symbolizowania w naukach humanistycznych. Symbolizować można nie tylko przedmioty, własności, stany rzeczy, ale również wartości. W przypadku żywych istot symbole nie są definiowane w rzeczywistym świecie, lecz go odzwierciedlają. Istnieje izomorfizm między symbolicznymi reprezentacjami a światem działania i percepcji, przy czym znaczenie symboli nie może zostać zredukowane do percepcji i działania⁸¹. Zdolność do postrzegania świata symbolicznie, a także do abstrakcyjnego myślenia, wyewoluowała z konkretnych

76 A. Newell, H.A. Simon, *Computer science as empirical inquiry...*, *op. cit.*

77 A. Newell, *Physical symbol systems...*, *op. cit.*

78 D.S. Touretzky, D.A. Pomerleau, *Reconstructing physical symbol systems*, „Cognitive Science”, t. 18, nr 2, 1994, ss. 345–353.

79 A. Newell, *Physical symbol systems...*, *op. cit.*, tłum. własne.

80 J. Boruszewski, *Fizyczne systemy symboliczne i ich magia*, „Filo-Sofija”, t. 17, nr 36, 2017.

81 R.N. Shepard, *Toward a universal law of generalization for psychological science*, „Science”, t. 237, nr 4820, 1987, ss. 1317–1323.

form postrzegania. Sposoby, w jakie żywe istoty modelują świat w swoim umyśle jest zgodny z konkretnymi uwarunkowaniami filogenetycznymi jak i ontogenetycznymi. Wszystkie zwierzęta uczą się i mają wspomnienia proceduralne. Na pewnym poziomie (z pewnością na poziomie naczelnym) zwierzęta są w stanie reprezentować sobie świat poprzez uogólnione zapisy doświadczeń⁸². Podczas ewolucji biologicznej ma miejsce imitacja, której celem jest przysposobienie nowych umiejętności. Mózgi żywych istot nie rozwijają jednak w jej trakcie nowych sposobów funkcjonowania, lecz używają wcześniej używanych struktur w zupełnie nowych celach. W trakcie ewolucji te same zmysłowe obszary mózgu żywych istot zostają wykorzystane do nowych zadań⁸³.

W odróżnieniu do działania fizycznego systemu symbolicznego zaproponowanego przez Newella, znaczenie symboli nie jest definiowane przez żywe organizmy jako odniesienie jednego symbolu do innego, lecz poprzez jego związek z wieloma innymi symbolami. Znaczenie symbolu zależy od relacji z innymi symbolami oraz od bardzo szerokiego kontekstu w jakim te symbole się pojawiają⁸⁴. Oczywiście takie podejście nie rozwiązuje problemu symbolizowania, ale skutecznie zmienia sposób podejścia. Znaczenie symboli nie jest sprowadzane do bezpośredniej funkcji wynikania, ani też do pojedynczych cech percepcyjnych lub działań, lecz pokazuje, że należy zbadać korespondencję między różnymi poziomami symbolicznymi. Składnia może ograniczać semantyczną interpretację, mimo, że syntaktyczna struktura zdania zostaje tak skonstruowana, by określać jego znaczenie. Trudno jednak wyobrazić sobie taki opis rzeczy, by zawierał w sobie wszelkie możliwe znaczenia przedmiotu lub sytuacji, których nieodłączną częścią są doświadczenia, często związane z wrażeniami lub emocjami, niemożliwymi do odtworzenia w sztucznym systemie (choćby znaczenie smaku lub węchu).

Widać więc, że pewne struktury symboliczne są zależne od innych struktur symbolicznych, które mogą być kształtowane w sposób zależny tylko od percepcji podmiotu poznającego⁸⁵. Znacząco utrudnia to uzyskanie symbolizowania semantycznego poprzez desygnację, która musiałaby składać się z olbrzymiej (a być może nieskończonej) ilości powiązań. Dodatkowo użycie symboli wymaga pewnej istniejącej uprzednio wiedzy

82 Por. J. Pelc, *Wstęp do semiotyki*, Warszawa 1984, s. 221–225.

83 Np. Sposób używania ludzkich nóg jest zakodowany w tym samym obszarze mózgu odpowiedzialnym za działania motoryczne, gdzie na wcześniejszych etapach ewolucji zapisywane były sposoby używania płetw.

84 R.N. Shepard, *Toward a universal law of generalization for psychological science...*, *op. cit.*

85 J. Pelc, *Wstęp do semiotyki...*, *op. cit.*, s. 79.

o nich samych oraz ich relacji. Fizyczne systemy symboliczne nie posiadają takiej wiedzy. Zwolennicy teorii Newella nie starają się używać symboli w sposób subiektywno-podmiotowy, lecz uniwersalny, co prowadzi do zredukowania jego funkcji: pominięcia jego zmiennego i praktycznego funkcjonowania. W fizycznym systemie dyskretne symbole nie mogą reprezentować rzeczywistego analogowego świata, który w dużej części składa się z niedefiniowalnych stanów. Dla sztucznego systemu świat jawi się zawsze jako pewien skończony i czytelny stan: jednostopniowa reprezentacja. Oczywiście systemy formalne są coraz bardziej zaawansowane i wyrafinowane, więc mogą tworzyć wiele interpretacji otrzymanych danych. Problem jednak pozostaje ten sam: nie tworzą one własnej semantycznej interpretacji, lecz przedstawiają to, co zostało na nie narzucone przez programistów. Niezależnie od tego, jak bardzo są systemem o możliwościach wielokrotnej interpretacji, brakuje im semantyki z „pierwszej ręki”, która świadczyć mogłaby nie tylko o ich własnej autonomii, ale i inteligencji.

W programie sztucznej inteligencji GOFAI zakłada się – zgodnie zresztą z koncepcją sygnału, że pomimo posiadania wysoce bogatego systemu znaczeń, w świecie ożywionym istnieje wspólny, podstawowy poziom symboliczny⁸⁶. Nie uzasadnia to jednak umniejszenia znaczenia systemów symbolicznych odpowiedzialnych za procesy czuciowo-motoryczne. Odrzucanie redukcjonizmu, który zaprzecza autonomii procesów symbolicznych oznacza, że należy dołożyć wszelkich starań, aby lepiej zrozumieć, jak wyglądają i jak są skoordynowane symboliczne i nie-symboliczne reprezentacje. W systemach koneksjonistycznych podjęto próbę zastosowania systemów subsymbolicznych, w których treści semantyczne mogłyby zostać lepiej odwzorowane dzięki rozproszonym reprezentacjom. Ale te systemy same w sobie również pozostały czysto obliczeniowe.

Mimo wszystkich zarzutów, formalny system symboliczny pozostaje wyjątkowo cennym narzędziem, ponieważ niemalże każda teoria może zostać opisana w kategoriach obliczeniowych. Większość teorii przyjętych przez człowieka – nawet teoria poznania – mają pewną organizację przyczynową, a obliczenia są wystarczająco elastyczne, aby tą organizację przedstawić niezależnie od tego, czy związki przyczynowe zachodzą pomiędzy reprezentacjami wysokiego poziomu, czy między niskopoziomowymi procesami neuronalnymi. Rozbudowane narzędzia matematyki pozwalają ukazać niemalże każdą teorię poprzez odpowiednią formę obliczeniową. Można je zatem wykorzystać dla nowego

86 T.A. Sebeok, *Contributions to the Doctrine of Signs*, University Press of America 1976, t. 4, s. 121.

rodzaju kognitywistyki poprzez stworzenie przestrzeni umożliwiającej pracę programistom i inżynierom.

1.5 Nowe problemy i możliwe kierunki rozwoju programu SI

Stworzenie ucieleśnionej sztucznej inteligencji jako ważnego programu badawczego pozwoliło przekroczyć wiele trudności GOF AI. Włączono różne podejścia i stworzono zunifikowaną teorię sztucznej inteligencji i kognitywistyki. Jednak takie próby nie były pozbawione problemów. Mimo, że ucieleśniona sztuczna inteligencja unika lub rozwiązuje wiele problemów napotykanym przez „dobrą staromodną sztuczną inteligencję” problem powstawania w pełni inteligentnych zachowań wcale nie znika. Opisany powyżej problem redukcyjnego fizycznego systemu symbolicznego jest wyjątkowo trudny do rozwiązania. Sztuczna inteligencja opiera się o działanie maszyn symbolicznych i nie wygląda na to, żeby miało się to w najbliższej przyszłości zmienić. Szczególna trudność dotyczy autonomiczności, intencjonalności oraz celowości systemów. Z perspektywy inżynierii nie mają one większego znaczenia, ponieważ systemy osiągają cele nałożone na nie przez projektantów, lecz dla badaczy SI i filozofów, którzy badają mechanizmy inteligencji lub starają się symulować naturalne procesy poznania, sztuczne systemy nadal pozostają mocno ograniczone i nie spełniają fundamentalnych założeń. Niewątpliwie w ucieleśnionej SI opracowano metody, które pozwalają formułować lepsze modele naturalnego poznania niż GOF AI. Nie ma jednak wątpliwości, że to podejście nie rozwiązało fundamentalnego problemu. Wciąż nie powstały maszyny, które można nazwać inteligentnymi w silnym znaczeniu SI.

Ucieleśnienie i usytuowanie sztucznego systemu wydają się nie wystarczać do stworzenia perspektywy, która będzie tworzyć znaczące odniesienie dla niego samego. Do dokładniejszej diagnozy obecnej ucieleśnionej sztucznej inteligencji niezbędne wydaje się drobiazgowo opracowanie biologicznych podstaw poznania, które są wciąż trywializowane w świetle złożoności czynników biologicznych zaangażowanych w konstytuowanie się inteligencji. Samo sprzężenie percepcji i działania jest niewątpliwie warunkiem koniecznym do powstania inteligentnego zachowania, ale bez wątpienia niewystarczającym. Wewnętrzny i zewnętrzny świat pozostaje pozbawiony znaczenia

innego niż zdefiniowane przez twórców. Wciąż nie można z pełną odpowiedzialnością stwierdzić, że sztuczny system realizuje własne, istotne dla niego samego działania⁸⁷.

Zatem powraca problem związany z działaniem umysłu i ze świadomością. Dreyfus twierdzi, że uzyskanie prawdziwie inteligentnych maszyn wymaga przewyciężenia trudności związanych z „heideggerowską SI”⁸⁸. Dreyfus słusznie podnosi problem, który związany jest z powstawaniem znaczenia w sztucznych systemach. Żywe organizmy nadają przedmiotom, sytuacjom, stanom rzeczy znaczenie, na które są szczególnie uwrażliwione, ponieważ pozwala im ono reagować na to, co jest dla nich ważne. Chodzi tu zarówno o rzeczy w świecie mogące zaspokoić potrzeby biologiczne, jak również – w przypadku wyższych zwierząt – psychiczne. Ten problem w sztucznej inteligencji nadal nie znajduje rozwiązania. Główna trudność polega na określeniu, czym miałyby być znaczenie dla sztucznego agenta, który dopasowywałby do niego swoje działanie autonomicznie zgodne z własnym systemem semantycznym.

Problem usytuowania symboli został częściowo rozwiązany poprzez osadzenie systemów w świecie rzeczywistym, co pozwoliło na bezpośrednie odbieranie przez układ fizyczny bodźców płynących ze środowiska. Mimo rezygnacji z tworzenia wewnętrznych reprezentacji symbolicznych w sztucznych systemach i opieraniu się na danych otrzymanych ze środowiska, żaden sztuczny system nie potrafi działać w ten sposób, by współistnieć ze światem, który powinien percypować jako zorganizowany pod względem potrzeb, zainteresowań i zdolności. Żywe umysły nie przekształcają po prostu bodźców w reakcje odruchowe, jak życzyłby sobie tego Brooks. Model Brooksa reprezentować może w najlepszym razie stany mózgu, ale pozbawione znaczenia, które nadawane są rzeczom w codziennym świecie żywych organizmów. Sztuczne sieci neuronowe symulują przetwarzanie danych w mózgu, ale nie zawierają symboli reprezentujących cechy świata. SSN posiadają z góry ustalone założenia oraz cele ich twórców. W rezultacie stanowią serie dyskretnych przejść, uczące się w wyniku przeszłych doświadczeń sukcesu lub porażki, które zostały uprzednio zdefiniowane, lecz nie są postrzegane jako znaczące. Model działania sieci pozostaje jedynie zewnętrznym opisem, pozbawionym jakiejkolwiek intencjonalności.

87 Tak jak mógłby ująć to Thomas Nagel. Nie ma nic, co ukazałoby „jak to jest być” takim agentem. Por. T. Nagel, *Jak to jest być nietoperzem*, [w:] *Pytania ostateczne*, A. Romaniuk (tłum.), Warszawa 1997, s. 203–219.

88 H.L. Dreyfus, *Why Heideggerian AI failed and how fixing it would require making it more Heideggerian*, „Artificial Intelligence”, t. 171, nr 18, 2007, ss. 1137–1160.

Wciąż aktualny pozostaje też wzmiankowany wcześniej problem ramowy. John McCarthy i Patrick Hayes, którzy są twórcami tego pojęcia, używają go w odniesieniu do szczególnego, wąsko pojmowanego problemu dotyczącego reprezentacji, który pojawia się tylko w przypadku niektórych strategii radzenia sobie z problemem dotyczącym systemów planowania w czasie rzeczywistym⁸⁹. Problem ten dotyczy sztucznej inteligencji zarówno z perspektywy filozoficznej i inżynierskiej, ponieważ ma odpowiedzieć na pytanie w jaki sposób sztuczny system powinien tworzyć model świata, który pozostaje zgodny ze zmieniającym się środowiskiem rzeczywistym i jak opisać wystarczająco skutki działań bez konieczności reprezentowania olbrzymiej ilości nieoczywistych efektów⁹⁰. Zatem, jeśli system posiada pewien model aktualnego stanu świata, lecz w świecie tym coś się zmienia, jak ma ustalić co należy zaktualizować, a co pozostawić niezmienione. W ucieleśnionych systemach próbowano rozwiązać ten problem, poprzez wykorzystywanie świata jako modelu z pominięciem wewnętrznych reprezentacji, co pozwoliłoby rozwiązać szybkie i ciągłe dezaktualizacje modelu. Ucieleśnione i usytuowane systemy reagowały na ustalone możliwe do wyizolowania cechy środowiska, ale nie na kontekst czy zmianę znaczenia. Agent mógł postrzegać rzeczy świata, które został uprzednio opisane a wszystkie możliwe stany i zjawiska zostały wcześniej określone. Dynamika interakcji robota i jego otoczenia została uznana za podstawowy wyznacznik jego inteligencji⁹¹, ale w żaden sposób nie odnosiła się do tego, jak reagować na nowe informacje, ani jak uprzednie doświadczenia mogą zmienić znaczenie.

Dodatkowy istotny problem filozoficzny wynika z antropocentrycznego postrzegania świata przez człowieka. Jedyne znaczenie, jakiego człowiek doświadcza i które jest dla niego aktualne, wiąże się z jego własnym systemem znaczeniowym, lecz ma tendencję do nieuzasadnionego przypisywania zachowań, motywacji, intencji czynnikom i rzeczom reszcie nie-ludzkiego świata⁹². Nie ma on dostępu do znaczeń tworzonych przez inne gatunki zwierząt ani tym bardziej przez sztuczne systemy. Jeśli zatem znaczenie ma powstać w sztucznym systemie będzie i musi się odnosić do ludzkiego znaczenia, aby

89 J. McCarthy, P.J. Hayes, *Some philosophical problems from the standpoint of artificial intelligence*, [w:] *Readings in artificial intelligence*, Elsevier 1981, ss. 431–450.

90 Więcej o problemie ramy i sposobach jego rozwiązania w M. Shanahan, *The Frame Problem*, [w:] *The Stanford Encyclopedia of Philosophy*, E. N. Zalta (red.), Metaphysics Research Lab, Stanford University 2016.

91 R.A. Brooks, *Intelligence without representation...*, *op. cit.*

92 N. Epley, A. Waytz, J.T. Cacioppo, *On seeing human: A three-factor theory of anthropomorphism.*, „Psychological Review”, t. 114, nr 4, 2007, ss. 864–886.

mogło zostać przez nas odpowiednio zinterpretowane. Jedynym rozwiązaniem jest zbudowanie szczególnego modelu ludzkiego sposobu osadzenia w świecie rzeczywistym, aby mógł on doświadczać tego, co jest istotne dla ludzkiego pierwowzoru. Wychodząc z podobnych założeń, Dreyfus dowodzi, że model komputerowy musiałby otrzymać szczegółowy opis ludzkiego ciała i motywacji, aby mógł działać inteligentnie i by świat stał się dla niego znaczący⁹³. Wydaje się jednak, że ta propozycja wciąż jest niewystarczająca, ponieważ brak własnego znaczenia dla systemów SI pozostaje aktualny. Ponadto wystarczająco dokładny model ciała żywego człowieka jest niemożliwy do stworzenia chociażby z powodu naturalnych ograniczeń sztucznych systemów. Z drugiej strony takie szczegółowe i nieukierunkowane podejście do modelowania jest zbędne, ponieważ nie pomaga ani zrozumieć fenomenu inteligencji żywego organizmu, ani tym bardziej stworzyć sztucznej inteligencji.

Jedną z najważniejszych kwestii związanych z inteligentnym działaniem jest autonomia agenta. Podmiot autonomiczny odczytuje znaki świata rzeczywistego i działa w nim tak, aby zwiększyć swoją skuteczność i osiągnąć swoje cele. Nie da się jednak przypisać takich umiejętności sztucznym systemom. Różnica między żywymi a sztucznymi agentami polega na tym, że te pierwsze posiadają raczej wewnętrzne ukierunkowanie na cel – same go generują. W przypadku sztucznych systemów cel zostaje zewnętrznie narzucony przez programistę i opisany jako zamierzony przez zewnętrznego obserwatora. Z perspektywy inżynierskiej nie ma nic złego w maszynach inteligentnych, które za cel stawiają sobie służyć człowiekowi bądź produkowanie jak największej ilości samochodów. Takie podejście pozwala też uniknąć obaw o przyszłość współistnienia człowieka i sztucznej inteligencji, poprzez wyeliminowanie możliwości „buntu” bądź „wyzwolenia się” sztucznych systemów. Z drugiej jednak strony nie ma w tym momencie powodu, by sądzić, że sztuczna inteligencja będzie podzielać motywacje człowieka. Mimo tego rozważając problem sztucznej inteligencji w silnym sensie, należy przyjąć, że musi posiadać takie własności jak planowanie i osiąganie własnych celów, chociaż mogą one być zupełnie odmienne od ludzkich.

Dzięki EAI możliwe stało się lepsze zrozumienie znaczenia umysłu i obliczeń oraz rozwój koncepcji ucieleśnionego systemu. Nauki kognitywne coraz wyraźniej podkreślają rolę działania, percepcji i doświadczenia w poznaniu ludzkim, w przeciwieństwie do

93 H.L. Dreyfus, *Why Heideggerian AI failed and how fixing it would require making it more Heideggerian...*, op. cit.

bezcielesnego wnioskowania i rozumowania w tradycyjnym podejściu. Zwolennicy zmiany założeń programu sztucznej inteligencji powinni starać się jak najlepiej wykorzystać empiryczny aspekt poznania. Następnym zadaniem SI jest przyjrzenie się procesom doświadczania i tworzenia się znaczenia. Teorie obliczeniowe bez wątpienia odgrywają kluczową rolę w badaniach, ale istnieją również alternatywne formy wyjaśnień, którym należy się wnikliwie przyjrzeć. Koniecznym jest, aby odkryć pojęcia, które będą niezbędne do zrozumienia inteligencji, a z których można skorzystać w projektowaniu inteligentnych systemów. Aby uchwycić te koncepcje różnych form doświadczenia należy użyć już takich dziedzin wiedzy, które mogą się mieć rewolucyjny wpływ na zrozumienie pewnych pojęć i rozwój wiedzy związanych z semantyką żywych i sztucznych systemów. Doskonałym punktem wyjścia wydają się być ramy koncepcyjne związane procesami semiotycznymi i biologicznymi żywych organizmów. Takie podejście daje szansę na krytyczną analizę fundamentalnych problemów leżących u podstaw ugruntowania symboli oraz kształtowania się znaczenia. Dokładne zbadanie sposobów użycia znaków stosowanych przez autonomiczne organizmy może pozwolić je zastosować w sztucznych systemach. Niezbędnym jest przyjęcie perspektywy, która pozwoli przyjrzeć się semiotyce sygnałów odbieranych przez żywych autonomicznych agentów, a których wzór być może da się wykorzystać w tworzeniu sztucznych systemów. Biosemiotyczna teoria odniesienia wykorzystuje etologiczną teorię interakcji między żywymi organizmami a środowiskiem, w którym bytują. Próbę zastosowania biosemiotyki podjął już wcześniej Brooks, lecz skupił się na problemie ucieleśnienia i usytuowania systemu, banalizując problem kształtowania się znaczenia. Tymczasem wiedza o procesach semiotycznych żywych organizmów może doprowadzić do zastosowania jej w sztucznych systemach, których największym problemem jest brak semantycznych rozwiązań.

Koncepcje inspirowane biologicznie raz na jakiś czas są przywoływane w badaniach sztucznej inteligencji. Trudno powiedzieć, dlaczego idee biosemiotyczne nie zostały szerzej wykorzystane. Można tylko przypuszczać, że trudności są spowodowane niedostateczną wiedzą na temat działania umysłu, świadomości a także biologicznych uwarunkowań, które kształtują znaczenie. W dziedzinie sztucznego życia (ALife, AL) badacze robotyki koncentrują się na problemach związanych ze znakami w sztucznych systemach. Dotychczasowe prace w SI. Które były biosemiotycznie inspirowane dotyczyły badania komunikację stygmergiczną, programowanie agentowe (*agent-based model*), które oferują dużą elastyczność sztucznych systemów, ich decentralizację, samoorganizację i uniwersalność lub modułową architekturę subsumcyjną. Równie ważne są prace nad

systemami, których działanie opiera się na używaniu znaków i specyficznego dla nich języka. W szczególności aktywną dziedziną badań są techniki mapowania obiektów w konstrukcjach wewnętrznych systemów w ich autonomicznym środowisku (*anchorage symbol*). Zastosowanie rozwiązań biosemiotyki pozwala na stworzenie nowego modelu poznawczego, który poza opisem działań i zachowań ma też potencjał twórczy.

W wyniku inspiracji biologicznej powstały algorytmy ewolucyjne lub adaptacyjne systemy robotyczne, które doskonale radzą sobie z nałożonymi zadaniami, jednak nie mają zastosowania w badaniach nad sztuczną inteligencją w wersji silnej. Brakuje tu wiedzy, która pozwoliłaby na odkrycie relacji organizmów ze środowiskiem, dzięki znakom znajdującym się już w ich polu semantycznym, to znaczy porzuceniu koncepcji separacji fizycznego symbolu i jego sygnału. Powstaje zatem pytanie, w jaki sposób takie znaki o charakterze czysto indykatywnym można wytworzyć w systemach sztucznej inteligencji? W rzeczywistości większość podejść do inteligentnych systemów przyjmuje naiwną definicję procesów semiotycznych, która zwykle odgrywa drugorzędną rolę w badaniach. Warto się przyjrzeć zestawowi narzędzi do analizy i konstrukcji systemów autonomicznych, którymi dysponuje etologia. Po dokładnym zbadaniu interakcji żywych organizmów z otoczeniem otrzymamy szczegółową analizę działania narządów zmysłowych i motorycznych, oraz ich roli w procesach kształtowania się znaczenia. Analiza ta pozwoli sformułować program badawczy w celu wyjaśnienia znaczenia znaków napotykanym przez niezależnych agentów. W tym celu należy przygotować silne teoretyczne podłoża niezbędne dla realizacji semiozy w sztucznych systemach. Zapewnić to może biologicznie inspirowany model semiotyczny, który obejmować będzie zarówno możliwości jak i ograniczenia symulowanej semiozy, opis podejścia do przeprowadzania eksperymentów ze sztucznymi systemami w syntetycznym kontekście etologicznym. Podejście to może zostać użyte do zbadania sposobu użycia symboli w sztucznym systemie, możliwości powstawania znaczenia i związanych z nim pojawianiem się zachowań inteligentnych.

2 Zagadnienie inteligencji w perspektywie biosemiotycznej

Poprzedni rozdział przedstawił główne problemy silnej sztucznej inteligencji, które można w skrócie podsumować jako brak zaangażowania sztucznych systemów w semiotyczne działania. Ponieważ większość krytyków programu SI skupia się na tym zagadnieniu, należy przyjrzeć się tej kwestii i poszukać odpowiedzi, które pozwoliłyby rozwikłać ten problem. Przede wszystkim, należy zwrócić uwagę na to, co sprawia, że właśnie semiotyczne działania mają być kluczowe dla powstawania inteligencji. Warto również zauważyć, że badania przyjmujące paradygmat ucieleśnionego umysłu w SI odniosły znacznie większy sukces niż klasyczne, symboliczne podejście. W tej perspektywie badania biosemiotyczne wydają się dawać interesujące wyniki związane z inteligentnym zachowaniem, które mogą okazać się strategiczne dla badań nad inteligentnymi maszynami.

Zagadnienia i problemy opisane w pierwszym rozdziale, zmuszają do przyjrzenia się biosemiotycznej koncepcji powstawania zachowań. Cechy poznawcze i emocjonalne człowieka nie są w biosemiotyce uważane za wyjątkowe, co nie pozwala na traktowanie gatunku *Homo sapiens* w sposób nadzwyczajny w badaniu zjawisk naturalnych. Zakłada się bowiem, że system mentalny człowieka pojawił się w wyniku procesu ewolucyjnego a informacji o jego rozwoju należy szukać w naturze prostszych zwierząt. Biosemiotyka sugeruje, że sposób rozwiązywania problemów, użycie języka, zastosowanie wiedzy oraz rozumowanie są łatwiejsze do wytłumaczenia, jeśli zrozumiemy istotę bycia i funkcjonowania prostszych żywych organizmów. Esencją tej istoty jest często zdolność działania w dynamicznym środowisku i postrzegania otoczenia w stopniu wystarczającym do utrzymania życia i reprodukcji. Istnienie wyrafinowanej ludzkiej inteligencji utrudnia krytyczną analizę inteligencji innych zwierząt. Antropocentryczna tendencja do badania życia zwierząt w ich środowisku powoduje poważne utrudnienie prac, ze względu na niezmiennie porównywanie ze człowiekiem i jego światem fenomenalnym. Stanowi też wyzwania, które muszą zostać skonfrontowane i rozwiązywane w każdym konkretnym przypadku. W badaniach nad inteligencją działanie ludzkie jest traktowane jako jedyny model, a ludzie stają się wyjątkowymi istotami. Jednak, tak jak każdy inny gatunek na świecie, gatunek ludzki jest wytworem ewolucji a świat, który zdołał stworzyć ludzki gatunek, nie może być pozbawiony śladu ludzkich zdolności.

2.1 Koncepcja semiozy w biologii

Biosemiotyka jest gałęzią semiotyki ogólnej, łączącą biologię i semiotykę, której głównym celem jest wykazanie, że to semioza jest zasadniczym aspektem życia. Jej dążeniem jest budowanie pomostu między biologią, filozofią, językoznawstwem i badaniami nad komunikacją. Główne wyzwanie stanowi pogłębienie wiedzy naukowej na temat informacji biologicznej, w przekonaniu, że kody organiczne i procesy interpretacji są podstawowymi elementami świata żywego⁹⁴. Thomas Sebeok, językoznawca, rozpoczął badania nad biosemiotyką w celu zbadania biologicznych korzeni ludzkiej semiozy, co miało mu pomóc w zrozumieniu przejście od świata organicznego do świata sygnałów, znaków, komunikacji i języka. Niektórzy z biologów zaczęli dostrzegać, jak wiele zjawisk i funkcji znajdujących się w centrum życia organicznego (kod genetyczny, metabolizm komórki, działania ekosystemu itd.) można zbadać jako mechanizmy semiotyczne⁹⁵. Oznacza to, że znaki i powstałe w wyniku ich przetwarzania znaczenie odgrywają fundamentalną rolę we wszystkich działaniach żywych systemów. Proces semiozy jest nieodzowną cechą wszystkich znanych form życia jako zdolność do gromadzenia, replikowania i przekazywania wiadomości oraz wydobywania z nich sensu. Studiowanie tych procesów komunikacji i powstałego w ich wyniku znaczenia można uznać za dyscyplinę nauki o życiu związaną zarówno z naturą jak i z kulturą⁹⁶.

Biosemiotycy respektują złożoność procesów życiowych jako zbiór danych dostarczonych przez nauki biologiczne, ale również badania behawioralne. Biosemiotyka zajmuje się procesami przetwarzania znaków w przyrodzie we wszystkich badanych wymiarach, w tym pojawieniem się w przyrodzie semiozy, która może przewidywać pojawienie się żywych komórek, naturalną historię znaków, aspekty semiozy w ontogenezie organizmów, w komunikacji roślinnej i zwierzęcej, funkcje znaku w układach odpornościowym i nerwowym a także semiotykę poznania i języka⁹⁷. Biosemiotyka może

94 M. Barbieri, *Life is "artifact-making"*, „Journal of Biosemiotics”, t. 1, nr 1, 2005, ss. 81–101.

95 M. Anderson i in., *A semiotic perspective on the sciences: Steps toward a new paradigm*, „Semiotica”, t. 52, nr 1/2, 1984, ss. 7–47.

96 T.A. Sebeok, *The semiotic self revisited*, [w:] *A sign is just a sign*, 1991, s. 22.

97 C. Emmeche, *Modeling life: A note on the semiotics of emergence and computation in artificial and natural living systems*, [w:] *Biosemiotics. The Semiotic Web*, T. A. Sebeok (red.), Berlin 1992, s. 78.

być postrzegana jako przyczynek do ogólnej teorii ewolucji, polegającej na syntezie różnych dyscyplin. Dziedzina proponuje nowe i ujednolicone podejścia do zjawiska życia: biorąc pod uwagę zarówno znaczenie i funkcje pojedynczego rybosomu jak i ekosystemu; początek życia i jego ostateczne znaczenie. W koncepcji biosemiotycznej podkreśla się, że sferę życia przenikają procesy przetwarzania znaków które skutkują powstaniem indywidualnego sensu. To, co wyczuwa organizm, oznacza dla niego zawsze coś istotnego: jedzenie, konieczność ucieczki, rozmnażanie płciowe. Jasper Hoffmeyer twierdzi, że semiotyczne wymagania stanowią decydujące wyzwanie dla sukcesu gatunku, a ewolucja organiczna świadczy o rozwoju coraz bardziej wyrafinowanych semiotycznych środków niezbędnych do przetrwania. Najbardziej widoczną cechą ewolucji organicznej nie jest tworzenie mnogości struktur morfologicznych, ale rozwój „wolności semiotycznej”, która oznacza znaczny wzrost „bogactwa lub głębi znaczenia”⁹⁸.

Równie istotna w dociekaniach biosemiotycznych jest rola obserwatora systemu, który stara się dotrzeć do natury różnych systemów, które może obserwować, opisywać i konceptualizować. To stanowisko zakłada możliwość przeprowadzenia eksperymentu naukowego, obserwacji, interpretacji, działania a także pomiaru różnych systemów. Kwestie te były wielokrotnie badane w obrębie tradycji filozofii nauki, lecz ich szczególne znaczenie w odniesieniu do systemów żyjących zostało opracowane przez biologów i naukowców zajmujących się systemami⁹⁹. Podstawowym problemem podczas rozważań, interpretacji czy obserwacji innych żywych systemów jest zagadnienie przypisywania cech antropomorficznych innym gatunkom. Podczas konstruowania teorii dotyczących semiotycznych aspektów żywych organizmów zachodzi tendencja projekcji konkretnych ludzkich umiejętności na przedmiot badań. Właściwości i funkcje wysokiego poziomu jak wnioskowanie racjonalne, świadoma intencjonalność lub zdolność językowa są często projektowane na opis zachowania zwierząt nawet gdy dostrzegamy, że zachowaniem systemu rządzą odmienne od ludzkich mechanizmy¹⁰⁰.

98 J. Hoffmeyer, *Signs of Meaning in the Universe*, Indiana University Press 1997, s. 61.

99 G. Kampis, *Self-modifying systems*, „Biology and Cognitive Science: A New Framework for Dynamics, Information, and Complexity”, 1991; H.H. Pattee, *Simulations, Realizations, and Theories of Life*, [w:] 1987, ss. 63–78; R. Rosen, *Fundamentals of measurement and representation of natural systems*, Elsevier Science Ltd 1978, t. 1; T. Uexküll, *Semiotics and the problem of the observer*, „Semiotica”, t. 48, nr 3/4, 1984, ss. 187–195.

100 C. Emmeche, J. Hoffmeyer, *From language to nature: The semiotic metaphor in biology*, „Semiotica”, t. 84, nr 1–2, 1991, ss. 1–42.

Celem biosemiotyki jest wyraźne określenie tych założeń, które są importowane do biologii przez koncepcje teleologiczne jak: funkcja, informacja, kod, sygnał, wskazówka itp. oraz zapewnienie im teoretycznego ugruntowania. Powszechne stosowanie tych terminów w biologii wskazuje na to, że takich pojęć nie można ani uniknąć ani całkowicie zastąpić zapisem logicznym, chemicznym lub matematycznym. Dlatego też biosemiotyka stawia sobie za zadanie utrwalenie takich pojęć w kontekście fizyczno-biologicznym; zdefiniowanie i powiązanie ich w celu unikania antropomorfizmów, które stają się zbędnymi ukrytymi założeniami. Biosemiotyka bardzo ostrożnie używa pojęć, które są stosowane zazwyczaj do opisu ewolucyjnie złożonego produktu procesów semiotycznych, jakim jest ludzka kultura. W celu uniknięcia błędów antropomorfizmu czy witalizmu, czynione są starania, by wyraźnie odróżnić, które z tych pojęć są odpowiednie do dokładnego poziomu badań¹⁰¹.

Historia semiotyki jest głęboko zakorzeniona przez strukturalizm w lingwistyce a biosemiotyka związana jest z podobnym ruchem strukturalistycznym w biologii teoretycznej. Różnorodne podejścia podkreślają właściwości relacji, znaczenia, całości i kontekstualności, dzięki czemu zaczynamy postrzegać żywe stworzenia nie tylko jako biernie poddane uniwersalnym prawom natury, ale także jako aktywne systemy produkcji znaków, mediacji znakowej i interpretacji znaków, które wykorzystują prawa fizyczne, aby żyć, a nawet problematyzować to życie. Podejście takie pozwala wyjaśnić istotę inteligentnych zachowań, ponieważ powstawanie życia pełnego znaczenia, intencjonalności, samoświadomości czy poczucia podmiotowości przestaje być niezrozumiałą tajemnicą. Należy tutaj wyjaśnić, dlaczego biosemiotyka jako badanie procesów znaku dostarcza silnego zestawu narzędzi koncepcyjnych do badania powstawania inteligencji w jej najbardziej podstawowych formach. Dlatego, na samym

101 Istnieje kilka specjalnych projektów mających na celu zwiększenie naukowego rygoru biosemiotyki. Jedną z takich inicjatyw jest projekt *Biosemiotic Glossary*, w którym różni autorzy publikują analizy koncepcji biosemiotycznych poprzez mapowanie i omawianie wykorzystania terminów używanych w teorii biosemiotycznej. Projekt ma też na celu pomoc w budowaniu wspólnej wiedzy na temat podstawowych narzędzi pojęciowych biosemiotyki oraz świadomości zróżnicowania używanych terminów. Patrz M. Tønnessen, *Biosemiosis: Questionnaire from Biosemiotics - and info about the biosemiotic glossary project*, [w:] <http://biosemiosis.blogspot.com/2013/11/questionnaire-from-biosemiotics-and.html> (4.12.2018).

początku trzeba zrozumieć, czym jest semioza¹⁰². Teza Sebeoka, że semioza jest tym, co odróżnia systemy ożywione od nieożywionych¹⁰³ jest kluczowym punktem wyjścia dla dziedziny biosemiotyki. Teza została już wcześniej sformułowana przez Friedricha Rothschilda, który twierdził, że „żywe systemy są od samego początku tworzone jako systemy znaków”¹⁰⁴. Ważnym i koniecznym założeniem biosemiotyki, jest istnienie znaczącej komunikacji u innych gatunków poza *Homo sapiens*.

Semioza oznacza proces tworzenia, odbierania i interpretowania znaków. Każda forma aktywności, zachowania lub procesu, który obejmuje znaki i w którym powstaje znaczenie, jest aspektem semiozy. Semioza jest procesem, który niesie znaczenie i w którym powstaje znaczenie. Proces ten pośredniczy w relacjach celowości i przyczynowości. Pozwala przezwyciężyć surowy dualizm umysłu i materii, poprzez badanie działania znaków zapewniając lepsze podejście do systemów żywych niż dychotomie właściwości psychicznych i fizycznych. Semioza to proces dynamiczny i mimo, że cel komunikacji znaku bywa osiągnięty, potencjalnie może trwać nieskończenie długo. Początkowo idea semiozy była rozwijana w celu powiązania języka z innymi systemami znakowymi, zarówno ludzkimi, jak i nieludzkimi. Język jest bez wątpienia prototypem teorii semiotycznych, lecz jego badanie wskazuje, że procesy semiozy można zastosować do innych systemów znaków¹⁰⁵. Istnieją również teorie opisujące systemy metasygnałów i traktujące język jako jeden z wielu kodów służących do komunikowania znaczenia¹⁰⁶. Przyjęcie tej perspektywy pozwala już wstępnie zdefiniować semiozę jako

102 Termin semioza został wprowadzony przez Charlesa Sandersa Peirce'a w celu opisanie procesu, który interpretuje znaki jako odnoszące się do obiektów. Peirce twierdził, że interpretacja znaku sama jest znakiem, który może być dalej interpretowany, a wynik tej interpretacji ponownie staje się znakiem. Semioza nie kończy się wraz z powstaniem interpretacji, lecz może trwać dalej. Patrz C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 5: Pragmatism and Pragmaticism*, Charlottesville, Va 1994, t. 5, par. 484; C.S. Peirce, *The essential Peirce: Selected philosophical writings*, N. Houser (red.), Indiana University Press 1998, t. 2, par. 411.

103 T.A. Sebeok, *Communication, Language and Speech: Evolutionary Considerations*, [w:] *Semiotic Theory and Practice: Proceedings of the Third International Congress of the IASS* Palermo, 1984, M. Herzfeld, L. Melazzo (red.), t. II, Berlin, Boston 1988, ss. 1083–1091.

104 F.S. Rothschild, *Laws of symbolic mediation in the dynamics of self and personality*, „Annals of the New York Academy of Sciences”, t. 96, nr 3, 1962, ss. 774–784.

105 W. Noth, *Semiotic machines*, „Cybernetics & Human Knowing”, t. 9, nr 1, 2002, ss. 5–21.

106 G. Bateson, *Information, codification, and metacommunication*, „Communication and Culture, Holt, Rinehart and Winston, New York”, 1966, ss. 412–426. Dla przykładu: podobnych dowodów dostarczył Ivan Pavlov w swoich eksperymentach, które pokazywały,

dowolną czynność odpowiadającą za komunikację znaczenia, dzięki której następuje ustanowienie relacji między znakami.

Dla ludzi, semioza jest zjawiskiem bardzo szerokim, dotyczącym m.in. systemów interakcji społecznych, których ważnym aspektem jest wymiana informacji. Istotnym przejawem jej badania powinno być sprawdzenie wszystkich warunków, które sprawiają, że przekazywanie, odbiór i interpretacja znaków są możliwe i skuteczne. Nie tylko więc przekaz językowy, lecz także sposoby myślenia, reakcje emocjonalne, przekonania, motywy i cele mają w mniejszym lub większym stopniu, charakter semiotyczny. Warto w tym miejscu zwrócić uwagę, że wszystkie wyżej opisane aspekty semiozy łączą się z warunkami powstawania inteligentnego zachowania. Przekazywane są nie tylko podstawowe informacje na temat obiektów i sytuacji, ale również zakodowane w przekazie stany emocjonalne, odczucia i przekonania. Przekaz znaków rozwija się dynamicznie się w procesie komunikowania. Odbiór znaków, interpretacja i produkcja oraz przekazywanie ich dalej nie jest ani chaotyczne, ani przypadkowe. Odpowiedź jest czymś więcej niż bodziec i reakcja. Poprawnie przetworzony znak sprawia, że w zaistniałych warunkach odpowiedź jest adekwatna. Spośród wielu możliwości interpretacji znaku, organizm wybiera optymalny dla siebie. W procesie semiozy następuje przekazywanie sygnałów w celu wywołania konkretnej reakcji, co w pewnym stopniu wymusza właściwe ustosunkowanie się do odbieranych sygnałów. Dzięki właściwej interpretacji znaków, mamy do czynienia z inteligentnymi odpowiedziami systemu na sytuacje, w których się znajduje, odpowiednio analizując je w kontekście środowiskowym.

W trakcie obserwacji zachowań żywych organizmów zauważono, że nawet najprostsze formy posługują się systemem znaków. Szczególnie widać to wśród zwierząt, które informują innych członków stada, gdzie znalazły pożywienie, ogłaszają innym osobnikom tego samego gatunku czas płodności, ostrzegają się przed zagrożeniem. Komunikaty przybierają różne formy: chemiczne, dźwiękowe, wizualne lub dotykowe, mogą być przekazywane jako pojedyncze sygnały lub jako ich kompleks¹⁰⁷. Sebeok sugeruje, że endosymbioza, samoodniesienie, funkcje receptorów, autopoeza i inne

że dzwonek dla psa jest bezpośrednią komunikacją informacji, ale kontekst komunikacji również przekazuje informacje, mimo, że żadna z tych form nie ma postaci werbalnej. Patrz. I.P. Pavlov, *Conditioned reflexes: An investigation of the physiological activity of the cerebral cortex*, G. V. Anrep (tłum.), Oxford University Press London 1928.

107 T.A. Sebeok, *Semiotics and ethology*, [w:] *Approaches to Animal Communication*, T. A. Sebeok, A. Ramsay (red.), Haga 1969, ss. 210–231.

właściwości wszystkich żywych systemów pokrywają się z definicją semiozy: coś jest żywe i komunikuje znaczenie¹⁰⁸. Nie przekazuje dyskretnych informacji, lecz nadaje im głębszy sens. Przekazywanie znaczenia przybiera formę, która może być opisana w kategoriach inteligentnego zachowania, ponieważ świadczy o odniesieniu się do okoliczności oraz właściwym przetworzeniu danych dotyczących sytuacji w jakiej system się znalazł.

Przebieg procesu semiozy wymaga również pewnego wartościowania znaków dostarczanych organizmowi. Do organizmów docierają olbrzymie ilości zmysłowych danych w pochodzących zarówno ze środowiska jak i z ich własnego organizmu. Gdyby wszystkie procesy semiozy przebiegały na jednakowym poziomie, musiałoby dojść do przeładowania sensorycznego, ponieważ mózg nie może sobie poradzić ze wszystkimi informacjami równocześnie. Zatem z perspektywy poznawczej tylko ważniejsze znaki będą traktowane priorytetowo. Musi więc istnieć klasyfikacja znaków pod względem ich znaczenia (choćby związane z przeżyciem). Znak zaczyna funkcjonować, kiedy zostanie odróżniony od szumu tła. Wtedy uruchamia się aktywność poznawcza, która interpretuje dane wejściowe i przekształca je w sensowne informacje. Kalevi Kull twierdzi, że istnieje powtarzalny porządek semiozy:

- a) podmiot poznawczy filtruje dane otoczenia i rozpoznaje znaki (na podstawie wcześniej istniejącego modelu w pamięci);
- b) w jego umyśle powstaje znaczenie – nowa struktura, lecz zachowany jest oczywisty izomorfizm między oryginalnym a nowym znakiem;
- c) znak jest interpretowany jako znaczący i dopasowany do istniejących wzorców i ich znaczeń przechowywanych w pamięci;
- d) wyniki udanej interpretacji są obserwowalną reakcją na postrzegane bodźce¹⁰⁹.

Na początku zachodzi więc proces rozpoznawania, który rozpoczyna proces semiozy i jest niezbędny dla innych procesów na innych poziomach. Znaczenie, które powstaje nie jest tylko przywołane z pamięci, lecz tworzone na nowo i porównywane z uprzednimi doświadczeniami. Inteligentna reakcja wymaga właściwego wybrania

108 T.A. Sebeok, *Animal. Biological and semiotic perspective*, [w:] *What is an Animal*, T. Ingold (red.), 1988, s. 72.

109 K. Kull, *On Semiosis, Umwelt, and Semiosphere*, „Semiotica”, t. 120, nr 3–4, 1998, ss. 299–310.

odpowiedniej interpretacji znaku, która może się jednak znacząco różnić od zapisów już istniejących.

Należy w tym miejscu zaznaczyć, że podmiot interpretujący znaki, sam również pozostaje wynikiem interpretacji – jego rozwój fizyczny, sposób działania pamięci są efektami uprzednich procesów semiozy. W tej perspektywie inteligentne działanie ukazuje się jako ciągła interpretacja znaczenia, na którą wpływ mają jej historyczne uwarunkowania. Według Jurija Łotmana „znak sam w sobie stanowi program dla procesu twórczego”¹¹⁰. Każdy znak, który ma zostać zinterpretowany w umyśle jest już uprzednio zinterpretowany i zmagazynowany w pamięci, a wynik interpretacji wpływa na proces przekazywania go jako znaku kolejnym odbiorcom. Ponieważ wszystkie składniki semiozy są wielokrotnie interpretowane, staje się możliwe tworzenie nowych, jak i zapominanie starych znaczeń. Świadczy to o kreatywności umysłu, jego możliwości tworzenia nowych związków i pewnego rodzaju „nadpisywania znaczenia” na uprzednio istniejące engramy – ślady pamięciowe. Zatem powstawanie znaczenia, nie jest wciąż na nowo powtarzanym procesem, lecz twórczym działaniem organizmu, wykorzystującym wcześniejsze informacje. W trakcie semiozy, pojawiają się związki między rzeczami, które nie wchodzą w interakcje ani nie wpływają na siebie nawzajem poprzez bezpośrednie fizyczne lub chemiczne procesy. Tego typu zjawiska semiotyczne nie należą do rzeczywistości fizycznej, lecz umysłowej. Są dowodem na istnienie zrozumienia oraz inteligentnego odniesienia się do rzeczywistości.

Relacja, która charakteryzuje semiozę właściwą wszystkim formom poznania, wyłania się z interpretacji znaków, będących wynikiem strukturalnego sprzężenia podmiotu i jego środowiska. Musi przebiegać w ten sposób, by zagwarantować spójność i trwałość świata organizmu. W rzeczywistości, struktury systemu życia jak i środowiska zmieniają się w wyniku wzajemnego oddziaływania. Sprzężenie, które powstaje w wyniku ich plastyczności, wytwarza autonomiczną i ściśle określoną jednostkę, która zostaje ukształtowana w trakcie swojej historii rozwoju. Poznanie jest łączone z ucieleśnieniem, osadzeniem i usytuowanym doświadczeniem¹¹¹, ponieważ zawsze obejmuje podmiot

110 Y. Lotman, *Universe of the Mind. A Semiotic Theory of Culture*, Indiana 1990, s. 101.

111 Por. A. D. Wilson, S. Golonka, *Ucieleśnienie poznania to nie to, co myślisz*, „Avant”, t. 5, 2014, ss. 21–56; G. Lakoff, M. Johnson, *Metaphors we live by*, University of Chicago press 2008; L.K. Miles, L.K. Nind, C.N. Macrae, *Moving Through Time*, „Psychological Science”, t. 21, nr 2, 2010, ss. 222–223.

określony zarówno fizyczną jak i umysłową architekturą, współdziałający ze światem, w którym jest zanurzony, tworząc dynamikę, która definiuje zdarzenia zachodzące w konkretnych kontekstach w przestrzeni i czasie. Jednostka taka jest zdolna do uczenia się na podstawie doświadczeń i do odpowiedniego dostosowywania swojego zachowania.

Semioza nie jest automatycznym procesem, lecz pozostaje związana z umiejętnościami żywego organizmu, który preferuje znaki posiadające dla niego szczególne znaczenie. Inteligentne zachowanie organizmu w trakcie semiozy można opisać jako umiejętne wykorzystanie takich funkcji jak: przekaz i tworzenie nowych informacji – szczególnie tych, które nie są do końca przewidywalne i nie muszą być zgodne z istniejącymi już wzorami, oraz zdolność do zachowania i wybiórczego odtwarzania informacji. W perspektywie semiotycznej inteligencja nie jawi się jako rozmyte pojęcie, lecz jako rozwinięta umiejętność używania znaków i właściwego wykorzystywania sygnałów płynących ze środowiska oraz porządkowania wiedzy. Howard Pattee i Robert Rosen podkreślają, że dzięki rozpatrywaniom procesów semiotycznych można dostrzec, w jaki sposób systemy postrzegają siebie nawzajem¹¹². Wydaje się, że jest to niemal niemożliwe zadanie dla czysto fizycznej nauki, lecz po przyjęciu ostrożnego opracowania nowych metod i zestawu narzędzi pojęciowych, można spróbować to uchwycić.

2.2 Zagadnienie znaczenia w pracach Jakoba von Uexküll

Jakob von Uexküll¹¹³ Jest uznawany za jednego z pionierów fizjologii behawioralnej i etologii oraz prekursora biocybernetyki. Swoją pracę poświęcił badaniu, w

112 R. Rosen, H.H. Pattee, R.L. Somorjai, *A symposium in theoretical biology*, „A question of physics: Conversations in physics and biology”, 1979, s. 87.

113 Jakob Johann von Uexküll urodził się 8 września 1864 roku w Keblaste (obecnie Mihkli) w Estonii, a zmarł na wyspie Capri w dniu 25 lipca 1944 roku. Studiował zoologię na Uniwersytecie w Dorpacie (teraz Tartu) w Estonii w latach 1884-189. Następnie pracował w Instytucie Fizjologii Uniwersytetu w Heidelbergu (w grupie Wilhelma Kühne, który był redaktorem naczelnym czasopisma biologicznego *Zeitschrift für Biologie*, oraz autorem pojęcia *enzymu*) i na stacji zoologicznej w Neapolu. W 1907 roku otrzymał tytuł doktora *honoris causa* Uniwersytetu w Heidelbergu za studia w dziedzinie fizjologii mięśni. Jeden z jego wyników stał się znany jako prawo Uexküll, które jest prawdopodobnie jednym z pierwszych sformułowań zasady negatywnego sprzężenia zwrotnego występującego w organizmie. Jego późniejsza praca poświęcona była badaniom *Umweltu*. Jest to pojęcie, dzięki

jaki sposób istoty żywe subiektywnie postrzegają swoje otoczenie i jak to postrzeganie determinuje ich zachowanie. W książce *Umwelt und Innenwelt der Tiere*¹¹⁴ wprowadził termin *Umwelt*, aby określić subiektywny świat organizmu i opracował specjalną metodę, którą nazwał *Umweltforschung*. Naukowym celem badań Uexküll'a było zbadanie zachowania organizmów żywych jako podmiotów w rodzinach, grupach i społecznościach. Szczególnie interesujący dla Uexküll'a był fakt, że znaki i znaczenie zajmują pierwszorzędne miejsce we wszystkich aspektach procesów życiowych. Jego pojęcie kręgu funkcjonalnego (*Funktionskreis*) można interpretować jako ogólny model procesów znakowych – semiozę. Natomiast jego dokładny opis pozwala zrozumieć w jaki sposób działanie każdego systemu jest naznaczone inteligencją.

Uexküll był jednym z biologów, który mocno akcentował zasadniczą rolę, jaką odgrywa fizyczna architektura organizmu w kształtowaniu form interakcji ze światem zewnętrznym. Uexküll starał się stworzyć nową koncepcję biologii, która nie miała się ograniczać do opisu struktur lub badania procesów chemicznych, fizycznych i mechanicznych, które w nich zachodzą. Naukowiec uważał, że w biologii brakuje teoretycznych podstaw, które pozwolą uporządkować wyniki badań i zrozumieć w jaki sposób żywe organizmy funkcjonują w świecie¹¹⁵. Odmienność perspektywy badawczej Uexküll'a polegała na wprowadzeniu zasady, uznającej rzeczywistość za „subiektywny wgląd” konkretnego organizmu¹¹⁶. W ten sposób zinterpretował poglądy Immanuela Kanta, przedstawiające obiekty jako zjawiska, które zawdzięczają swoją strukturę podmiotowi¹¹⁷. Biolog twierdził, że należy rozpatrzyć te struktury (które nazywał „tematami” człowieka i zwierzęcia), czyli zbadać rolę narządów zmysłów i oddziaływanie

któremu Uexküll jest najlepiej znany we współczesnej literaturze. W 1926 roku Uexküll założył *Institut für Umweltforschung* na Uniwersytecie w Hamburgu. W latach 1927-1939 spędzał letnie wakacje z rodziną na półwyspie Puhtu (zachodnie wybrzeże Estonii) w swojej letniej chatce (od 1949 roku jest to Stacja Biologiczna Puhtu Instytutu Zoologii i Botaniki w Tartu w Estonii), gdzie badał zachowanie organizmów żywych. Uexküll uważał się za zwolennika Johanna Müllera i Karla Ernsta von Baera. Jego filozoficzne poglądy opierały się na pracach Immanuela Kanta. Uexküll napisał jedną z pierwszych monografii z biologii teoretycznej (*Theoretische Biologie*, 1920), która wniosła niezwykle wkład w dziedziny obejmujące porównawczą fizjologię bezkręgowców, psychologię porównawczą a także filozofię biologii.

114 J. von Uexküll, *Umwelt und Innenwelt der Tiere*, Berlin 1921.

115 J. von Uexküll, *Theoretische biologie*, Paetel 1920, s. 7.

116 *Ibid.*, s. 9.

117 *Ibid.*, s. 8.

ośrodkowego układu nerwowego na konstruowanie doświadczanej rzeczywistości. Kolejnym celem nowej biologii miałoby być zbadanie związku zwierząt z doświadczanymi przedmiotami świata. Uexküll interesowało zatem badanie subiektywnej rzeczywistości ludzi i zwierząt: odmienne doświadczanie świata przez różne gatunki oraz wpływ zmysłowych odczuć na sposoby funkcjonowania żywych organizmów.

Uexküll zwrócił uwagę, że z powodu fizycznej natury i specyficznego aparatu zmysłowego, każdy organizm konstruuje swój szczególny obraz świata. Uważał, że dostęp do różnych światów organizmów można osiągnąć jedynie poprzez badanie ich specyficznej organizacji fizycznej. Według Uexküll'a dostęp ten wymagał zanurzenia się w strukturze anatomicznej badanego organizmu, która odpowiada za sposób, w jaki oddziałuje on ze światem zewnętrznym. Miało to stanowić najlepszy z możliwych sposobów określenia zakresu działania tej struktury i umożliwić opis istnienia i funkcjonowania badanego organizmu:

Rozpoczynamy spacer w słoneczny dzień, idąc przez kwiecistą łąkę, gdzie brzęczą chrząszcze i fruwią motyle, a my budujemy wokół każdego z zamieszkujących ją zwierząt mydlaną bańkę, która reprezentuje wszystkie tylko jemu dostępne cechy *Umweltu*. Jeśli sami wkraczamy w jedną z tych baniek, środowisko, które dotychczas istniało wokół, zostaje całkowicie zmienione. Wiele cech kolorowej łąki znika całkowicie, inne tracą znaczenie lub pojawiają się w nowych związkach. W każdej bańce mydlanej powstaje się nowy świat¹¹⁸.

Przedstawiona przez Uexküll'a metafora bańki mydlanej pokazuje, że poznanie jest procesem subiektywnym, który odbywa się w swoistym gatunkowi obszarze. Za dynamikę poznania odpowiada proces semiotyczny, który wiąże podmiot poznawczy z konkretnym kontekstem świata. Bańka opisuje szczególną symboliczną sferę, która zawiera w sobie zakres możliwych interakcji, które są zgodne ze sposobem oddziaływania dopuszczalnym przez fizyczną architekturę form organizmu. Rodzaj postrzegania kształtuje też otoczenie – *Umwelt*¹¹⁹, znaczący świat. W tym *Umwelcie* swoiste cechy środowiskowe stają się istotne

118 *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten.*, Rohwohlt Verlag, Hamburg 1934, s. 22. tłum. własne.

119 Pojęcie *Umweltu* tłumaczone jest czasami w języku polskim jako *wokół-swiat* lub środowisko. W języku niemieckim termin ten odnosi się do środowiska lub otoczenia naturalnego. Jednak w przypadku użycia tego terminu przez Uexküll'a, *Umwelt* jest wewnętrznym postrzeganiem rzeczywistości i należy go odróżnić od otoczenia – *Umgebung*. Otoczenie organizmów zmusza do traktowania ich jako obiektów umieszczonych w

i zindywidualizowane oraz są postrzegane jako jakościowo odrębne i mają przypisane konkretne znaczenia. W tym przekonaniu dotyczącym subiektywności doświadczanego świata Uexkülla popierał Ernst Cassirer twierdząc, że wszystkie żywe istoty posiadają swój krąg działania, który jest dla nich zarówno ograniczeniem jak i punktem widzenia otwartym na ten konkretny świat¹²⁰.

Systemy naturalne są wstępnie wyposażone w informacje genetyczne, które sterują organizmem przez różne cykle rozwojowe w różnych kontekstach środowiskowych. Stanowią one rodzaj wcielonej przedempirycznej wiedzy i odziedziczonej pamięci, która pozwala systemowi zidentyfikować i przypisać znaczenie konkretnym środowiskowym sygnałom w celu wytworzenia odpowiedniego zachowania będącego odpowiedzią zaspokajającą wymagania jego *Innenweltu*¹²¹. Pojęcie *Innenwelt* odnosi się do wewnętrznego świata i kontrastuje z *Umweltem*, wskazując na doświadczenie pochodzące z wnętrza organizmu. Jeżeli *Umweltowi* odpowiada szczególne postrzeganie świata, *Innenwelt* można zdefiniować przez stany wewnętrzne, które charakteryzują fizyczne stany jednostki w danym czasie. Koncepcja *Innenweltu*, wydaje się mieć systemową naturę i jest niezbędna, by wyjaśnić, dlaczego poszczególne cechy środowiskowe uzyskują większe znaczenie w porównaniu z innymi. Znaczenie jest kształtowane nie tylko przez *Umwelt*, ale również potrzeby organizmu odzwierciedlone w stanach jego *Innenweltu*.

Teoria Uexkülla zmusza do zastanowienia się, jak jest opisywana rzeczywistość świata, a także co to znaczy być zwierzęciem. Nie tylko mnoży światy w różnorodnych środowiskach, ale także stara się odrzucić rozumienie zwierzęcia jako bezdusznej maszyny, bezmyślnego lub beznamietnego obiektu. Wszystkie zwierzęta tworzą swój własny *Umwelt* jako skutek swoich percepcji, działań i relacji. Po raz kolejny zauważyć tu można biologiczną interpretację Kanta, kiedy Uexküll stwierdza, że świat zawdzięcza swoje istnienie wewnętrznej organizacji podmiotowej organizmu, który zamienia cechy sensoryczne w formę przestrzenną¹²². Wnioski Uexkülla, które miały na celu wprowadzenie nowego spojrzenia na zwierzęta i środowisko były bliskie powstania

środowisku, *Umwelt* zaś jest projekcją świata wytworzoną wewnątrz nich przez nie same. *Umwelt* to świat doświadczany przez indywidualny organizm.

120 E. Cassirer, *The philosophy of symbolic forms. The metaphysics of symbolic forms*, J. M. Krois, D. P. Verene (red.), Yale 1996, t. 4.

121 J. von Uexküll, *Umwelt und Innenwelt der Tiere...*, op. cit., s. 46.

122 J. von Uexküll, *Theoretische biologie...*, op. cit., s. 12.

opracowania ontologii zwierzęcia na podstawie etologicznych obserwacji¹²³. Uexküll wprowadził również nowy sposób myślenia o rzeczywistości jako takiej. Nie był pierwszym, który postulował, że rzeczywistość jest czymś więcej niż jedynie światem fizycznym, ale jednym z pierwszych biologów, którzy naprawdę podkreślali subiektywne doświadczenie zwierzęcia.

Rozwój inteligencji zwierzęcia można dostrzec obserwując jego zachowanie w jego własnym środowisku. Uexküll proponował zrozumienie zachowania każdego żywego systemu jako zgodnego z jego własnymi percepcjami i działaniami. Postulował, by nie traktować zwierząt jak obiektów, ale jako podmioty, których podstawową aktywnością jest postrzeganie i działanie. Wszystko, co postrzega podmiot, staje się jego światem percepcyjnym (*Merkwelt*) a wszystko, co robi, jego światem działania (*Wirkwelt*)¹²⁴. Uexküll skupił się na subiektywnym zwierzęciu i na jego indywidualnym otoczeniu, w którym powstają nowe związki i relacje, które mogą nie być dostrzeżone przez człowieka – badacza. Zwierzę w swoim środowisku ma do czynienia z wieloma obiektami, z którymi może wejść w interakcje. Część postrzeganych obiektów staje się nośnikami znaczenia, gdy wejdą w relację z podmiotem¹²⁵. Przedmioty są doświadczane przez wzgląd na ich funkcjonalne znaczenia dla podmiotu. Oznacza to, że tylko te aspekty środowiska zewnętrznego, które są w jakiś sposób istotne dla podmiotu są przetwarzane przez niego w celu ich wykorzystania do działania. Badanie jego percepcji nie może być wyizolowane od innych funkcji organizmu. Musi być traktowane jako faza działania związana z motoryką i aktywnością intelektualną. Żywe organizmy wykazują aktywność, o ile pojawia się dla nich jakieś znaczenie, które wynika z funkcjonalnych relacji z innymi obiektami, niezależnie od tego, czy są to relacje przestrzenne, czasowe, przyczynowości lub celowości. Uexküll w ten sposób starał się podkreślić, jak ważne jest nabywanie i

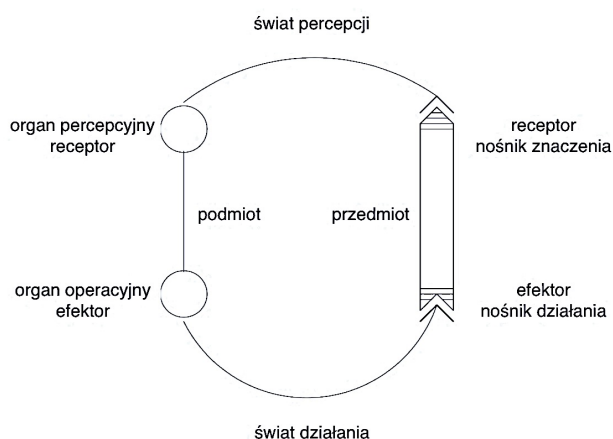
123 Etologia stała się powszechną koncepcją, dopiero gdy spopularyzował ją Konrad Lorenz, ale biologia Uexkülla, jak zauważył sam Lorenz, była pionierskim badaniem zachowania zwierząt i to Uexküll powinien być nazwany ojcem etologii.

124 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten*, op. cit., ss. 26–28.

125 *Ibid.*, ss. 105–107. Uexküll przedstawia przykład kamienia, który, zostaje rzucony w szczekającego psa. Kamień jako obiekt fizyczny nie zmienia się, a jednak zachodzi poważna zmiana jego znaczenia. Dopóki leżał na drodze, był wsparciem dla stóp i nie zaprzętał uwagi psa. Zmienił się w nośnik znaczenia, gdy tylko wszedł w relację z psem/podmiotem: stał się pociskiem i został związany z uczuciem bólu. Jego znaczenie zostało nadane przez podmiot. Przykład ten pokazuje znaczenie i prymat relacji funkcjonalnych.

powstawanie znaczenia, które jest niezbędne dla bycia podmiotem poznawczym posiadającym sensomotoryczną aktywność ciała, pozwalając mu na czynne i cielesne odnoszenie się do zewnętrznego środowiska.

Aby w pełni rozwinąć istotę teorii Uexküll'a należy przedstawić koncepcję kręgu funkcjonalnego. Pojęcie kręgu funkcjonalnego Uexküll wprowadził jako ważne narzędzie konceptualne, które podkreśla rolę podmiotu w jego interakcji ze środowiskiem i oferuje narzędzie do opisu sensomotorycznej aktywności. Zachowania w perspektywie badań Uexküll'a nie są zwykłymi ruchami czy tropizmami, jak życzyłby sobie jego przeciwnik Jacques Loeb¹²⁶, ale składają się z percepcji (*Merken*) i działania (*Wirken*). Zachowanie nie jest wynikiem mechanicznie regulowanych odruchów, ale w znacznym stopniu zorganizowane w związku z budową ciała podmiotu i jego zmysłowymi zdolnościami¹²⁷. Krąg funkcjonalny opisuje podstawową strukturę interakcji między organizmami zwierząt z przedmiotami pojawiającymi się w otaczającym ich świecie¹²⁸:



Rys. 1 Model kręgu funkcjonalnego wg Jakoba von Uexküll'a

126 Jacques Loeb rozwinął teorię zachowania zwierząt opartą na pojęciu tropizmu - mimowolnego, wymuszonego ruchu. Przedstawiał reakcje zwierząt na bodziec jako bezpośrednią i automatyczną odpowiedź. Uważał, że reakcja behawioralna jest wymuszona przez bodziec i nie wymaga wyjaśnień w kategoriach rzekomej świadomości zwierzęcia. Patrz. J. Loeb, *Concerning the theory of tropisms*, „Journal of Experimental Zoology”, t. 4, nr 1, 1907, ss. 151–156.

127 J. von Uexküll, *The Theory of Meaning*, „Semiotica”, t. 42, nr 1, 2009, s. 26.

128 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten.*, Rohwohlt Verlag, Hamburg 1934, s. 27.

Interakcje te składają się głównie z dwóch działań: neutralny obiekt ze środowiska zostaje ujęty jako nośnik znaczenia przez organ postrzegający lub komórkę postrzegającą, a następnie jest modyfikowany przez narząd efektorowy i wykorzystywany w celu udzielenia odpowiedzi (reakcji)¹²⁹. Każdy organizm aktywnie tworzy swój *Umwelt* poprzez powtarzające się interakcje ze światem. Jednocześnie obserwuje świat i go zmienia. Krąg funkcjonalny ma przedstawiać sposób interakcji organizmu jako podmiotu i przedmiotu jego działania. Pod tym pojęciem Uexküll opisał sposób odczytywania, przetwarzania i odpowiadania na sygnały płynące ze środowiska. W każdym procesie, w którym bierze udział organizm, wiodącą rolę odgrywa nośnik znaczenia¹³⁰, bądź jako nośnik bodźca percepcyjnego, bądź jako nośnik bodźca efektorowego. Według Uexküll'a to w kręgu funkcjonalnym dane sensoryczne są przekształcane w konkretne znaczenie, a sposób ich przetworzenia i wykorzystania wywołuje reakcję podmiotu. Znaczenie tu jest rozumiane jako struktura, łącząca postrzeganie i działanie. Kręgi funkcjonalne doskonale prezentują działanie biosemiozy, która jest podstawą formowania się procesów poznawczych. W kręgu funkcjonalnym dostrzec można, w jaki sposób następuje sprzężenie zwrotne organizmu i jego środowiska, które umożliwia rozumne działanie w *Umwelcie*. Obiekt i podmiot w kręgu łączą się ze sobą i tworzą całość. Dzięki plastyczności struktury znaczeniowej powstaje uformowana poprzez doświadczenie jednostka, która potrafi zidentyfikować znaki i nadać sens szczególnym wskazówkom środowiskowym, produkując zachowanie, które zewnętrzny obserwator postrzega jako adekwatne i inteligentne, a dla organizmu jest zgodne z jego własnymi potrzebami i wymaganiami.

Aby opisać model kręgu funkcjonalnego i zilustrować swoją koncepcję środowiska i umieszczenie organizmu w jego świecie Uexküll posłużył się przykładem kleszcza. W ten sposób opisał każdy z jego kręgów funkcjonalnych, które opisują interakcję podmiotu i obiektu, przez nośniki działania i nośniki znaczenia:

kleszcz pozostaje nieruchomy nieruchomo na gałęzi, dopóki nie pojawi się nośnik znaczenia; ssak przechodzący pod gałęzią, na której znajduje się kleszcz, emituje wytwarzany przez jego gruczoły, percepcyjny znak dla kleszcza – kwas masłowy, który zostaje odebrany przez jego receptor i staje się nośnikiem znaczenia dla kleszcza.

129 T. von Uexküll, *The Sign Theory of Jakob von Uexküll*, [w:] *Classics of Semiotics*, M. Krampen i in. (red.), Springer US 1987, s. 170.

130 J. von Uexküll, *The Theory of Meaning...*, *op. cit.*

Postrzeżony znak (percepcja kwasu masłowego) są przekształcane przez organy kleszcza w efektorowe nośniki znaczenia, które powodują określone oddziaływanie. Zostają przesłane do organów operacyjnych owada (do odnóży), wyzwalają efekt efektorowy i powodują, że kleszcz spada na ssaka;

a) percepcja włosów ssaka sprawia, że kleszcz porusza się po ciele ssaka w poszukiwaniu jego skóry;

b) uczucie ciepła skóry wywołuje reakcję kleszcza, której wynikiem jest wessanie się i picie krwi ssaka¹³¹.

Innymi słowy, podmiot używa swojego układu sensomotorycznego w celu efektywnego zaspokojenia swoich potrzeb. Wszystkie cechy podmiotu są ze sobą powiązane strukturalnie: cechy o znaczeniu percepcyjnym wpływają na cechy o charakterze operacyjnym i odwrotnie. Stan podmiotu zmienia się kilkakrotnie, powodując, że przystosowuje się do środowiska i sytuacji w najlepszy z możliwych sposobów. Obserwujemy działanie, które nosi w sobie znamiona inteligencji. Nawet w przypadku w tak prostego organizmu, jakim jest kleszcz, zauważyć można, że użycie narządów zmysłów i narządów ruchu jest w stanie przetworzyć powstające znaczenie w ten sposób, by prowadziło do osiągnięcia najlepszego skutku.

Traktowanie zmysłów i organów ruchowych zwierzęcia, jakby były częściami maszyny, jest powodem, dla którego ignoruje się ich rzeczywiste funkcje i działanie. Doznania i wola działania nie są pozorem, lecz wbudowanym w narządy zmysłów rodzajem postrzegania, a narządy poza działaniem mechanicznym, ukazują operatora, który jest wbudowany w te narządy. Wszystko, co postrzega i wszystko, co robi są wynikiem subiektywnego postrzegania świata, w którym działa w zgodzie z własnymi potrzebami i które tworzą razem zamkniętą całość, będącą jego *Umweltem*. Prosty cykl czynnościowy pokazuje, że zarówno sygnały receptorów, jak i efektorów są przejawami działania podmiotu a przedmiot jest jedynie nosicielem tych cech dla obiektu. Każdy podmiot żyje w świecie złożonym jedynie z subiektywnych rzeczywistości, którą niezwykle trudno byłoby odtworzyć u przedstawiciela innego gatunku. Obserwując zachowania zwierząt nasuwa się wniosek, że działają inteligentnie, niezależnie od tego, jak proste są przejawy takiego działania. Inteligencja ta nie może zostać opisana przy pomocy

131 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten.*, op. cit., ss. 28–29.

jednostronnych ludzkich definicji, ponieważ powstaje w wyniku subiektywnego gatunkowo postrzegania i działania.

Model kręgu funkcjonalnego zawiera wszystkie elementy, które są częścią procesu znaczenia, a ich połączenie przedstawia proces semiozy: organizm jest podmiotem interpretującym sygnały środowiskowe, których rolę odgrywają znaki, a stan biologiczny organizmu determinuje zachowanie. Z drugiej strony, przedmiot interpretowany jest trudniejszy do zidentyfikowania, ponieważ dla organizmu niekoniecznie istnieje jako abstrakcyjna jednostka, lecz może mieć tymczasowe istnienie jako różne obiekty semiotyczne o odmiennym znaczeniu, np. abstrakcyjny ssak ma trzy różne znaczenia dla kleszcza. Dlatego też wydaje się, że koncepcja kręgu funkcjonalnego może być ważnym narzędziem konceptualnym dla każdej teorii inteligencji. Łączy ona działanie i percepcję oraz umożliwia zrozumienie powodu i celu działania. Przedmioty w rzeczywistości mogą być powiązane z tymi samymi lub różnymi przedmiotami za pomocą kilku kręgów funkcjonalnych. Im bardziej skomplikowany organizm, tym działanie i ilość kręgów się problematyzuje. Mimo tego, każda aktywność organizmu, która składa się z percepcji i działania, odciska jego znaczenie na przedmiocie bez znaczenia i czyni go nośnikiem znaczenia dla podmiotu w danym *Umwelcie*¹³².

Zatem żeby zbadać umysł jakiegoś zwierzęcia, należy pamiętać, że jest tylko fragmentem wyizolowanym ze środowiska, które ludzki badacz postrzega jako świat człowieka. Dlatego też pierwszym zadaniem w rozpatrywaniu inteligentnego działania jakiegokolwiek systemu jest rozpoznanie sygnałów percepcyjnych spośród wszystkich bodźców w jego otoczeniu i zbudowanie z nich specyficznego świata, w którym pewne przedmioty i sytuacje stają się powodem, dla którego podejmuje działanie. Ilość obiektów, które organizm może odróżnić w swoim własnym świecie, jest równa liczbie funkcji, jakie może wykonać. Wraz ze wzrostem liczby funkcji, które może wykonać, rośnie ilość obiektów, które wypełniają rodzajowy *Umwelt*. Każde nowe doświadczenie wiąże się z ponownym dostosowaniem do nowych sytuacji. W ten sposób powstają nowe obrazy percepcyjne, posiadające nowe tony funkcyjne¹³³.

Koncepcja tonów funkcyjnych była niezwykle ważna dla Uexküll'a, który starał się wykazać, że każde zwierzę przypisuje sens przedmiotom, które napotyka, a które są istotne w konstruowaniu jego subiektywnego wszechświata. Po przyjęciu, że środowisko

132 *Ibid.*, s. 110.

133 *Ibid.*, s. 69.

organizmu jest zamkniętą jednostką percepcyjnego świata podmiotu, zbudowanym dzięki postrzeganiu i działaniu, należy się jednak zastanowić, w jaki sposób i dlaczego żywe organizmy wchodzi w interakcje z przedmiotami. Aby wyjaśnić to pojęcie Uexküll posłużył się przykładem człowieka, który pierwszy doświadcza obiektu jakim jest drabina i wchodzi po niej. Na początku dostrzega tylko pręty i puste przestrzenie między nimi i zastanawia się, w jaki sposób ich użyć. Następnie obraz pochodzący z jego receptorów zmysłowych zostaje uzupełniony przez efektorowe wyobrażenie działania jego ciała. W wyniku czego, przedmiot zyskuje znaczenie, które pojawia się jako funkcyjny ton – nowy efektorowy atrybut¹³⁴. Początkowo obcy i neutralny przedmiot przekształcił się w nośnik znaczenia. Podmiot wszedł z nim w relację i przekształcił go, w wyniku czego powstało inteligentne zachowanie, które umożliwiło podmiotowi właściwe wykorzystanie tego przedmiotu. Podmiot w pewnym sensie odcisnął znaczenie na przedmiocie, z którym został związany za pomocą kręgów funkcjonalnych.

Uexküll zauważył, że obiekty poznania są zawsze przekształcane w sygnały percepcyjne lub percepcyjne obrazy i wyposażone w tony funkcyjne. Sprawia to, że stają się rzeczywistymi przedmiotami *Umweltu*, mimo, że w samym bodźcu percepcyjnym nie pojawia się znaczenie¹³⁵. Przedmioty nie zyskują wartości, dopóki nie zostaną przekształcone i nie staną się nośnikami znaczenia. Znaczenie przedmiotu może się dzięki temu zmieniać i przybierać różne cechy i tony funkcyjne. Nie posiadają niezmiennych funkcji, tak jak obiekty, które tego znaczenia nie zyskały. Drabina jako przedmiot, z którym człowiek nie wszedł jeszcze w żadną relację, pozostaje pozbawiona znaczenia. Jest napotkaną nieznaną konstrukcją postrzeganą jako pręty i pusta przestrzeń pomiędzy nimi. Gdy staje się nośnikiem znaczenia, może służyć do wspięcia się na wysoki dach, lub jako kładka nad strumieniem, lub przyrząd do ćwiczeń fizycznych albo też materiał, który spalony w ognisku zapewni ciepło w chłodną noc. Wszystko zależy od kreatywności umysłu, która powstaje dzięki nałożeniu się postrzegania obiektu i wyobrazonego działania, zależnego od zaistniałej potrzeby. Znaczące przedmioty mogą przybierać różne cechy, mimo, że zmiany nie zależą od samego obiektu, ponieważ jego właściwości się nie zmieniają. Z punktu widzenia Uexkülla, właściwości przedmiotu są w rzeczywistości percepcją, która naznaczona zostaje znaczeniem przez podmiot wchodzący z nim w relację. We własnym subiektywnym wszechświecie podmiotu, przedmioty początkowo są

134 *Ibid.*, s. 67.

135 *Ibid.*, s. 110.

neutralne, od podmiotu zależy, czy staną się znaczące. Obiekty stają się nośnikami znaczeń, poprzez nadanie im tonu funkcyjnego, który zależy od nastroju i potrzeb podmiotu.

Po przyjęciu koncepcji nadawania tonów funkcyjnych przedmiotom, widać że działanie zwierząt nie musi być zawsze ukierunkowane na spełnienie podstawowych potrzeb i wypełnianie instynktownych zachowań. Zwierzęta mogą podobnie jak człowiek wykazywać aktywność, która jest daleka od biologicznych wymuszeń. Uexküll doskonale zdawał sobie sprawę z tego, że percepcja i działanie zależą również od doświadczenia podmiotu poznawczego. Dzięki doświadczeniu, które podmiot posiadał już wcześniej, mianowicie wyobrażeniu „wspinania się po czymś”, zdobył nową umiejętność i nowy sposób interpretowania napotkanego przedmiotu. Nowe doświadczenia dają podstawę do tworzenia uogólnień, które stają się punktem wyjścia dla tworzenia nowych poziomów ogólności. Zatem proces uczenia się, przyjmowania nowych wzorów, szukania alternatywnych rozwiązań może rozpocząć się ponownie na nowym poziomie. Po rozpoczęciu tego procesu nie może go zatrzymać, ponieważ nowo powstałe relacje generują nowe doświadczenia. Tą sieć interakcji, która obejmuje żywe systemy, Jesper Hoffmeyer nazywa semetyczną interakcją (*semethic interaction*), czyli taką, która odnosi się do interakcji, w których nawyki przyswojone przez gatunek lub osobnika stają się używane (i interpretowane) przez osobniki tego samego gatunku, a tym samym wywołują nowe nawyki u tego gatunku¹³⁶. Oznacza to, że każda regularność, która rozwija się w żywym systemie (na jakimkolwiek poziomie) ma tendencję do tworzenia nowych interpretacji i budowania zespołu nowych doświadczeń.

Koncepcja powstawania znaczenia, którą zaprezentował Uexküll daje możliwość wykorzystania biosemiotyki do badania inteligencji. Biosemiotyka, poprzez umieszczenie interpretacji w centrum uwagi, pokazuje, że semioza jest nieuniknioną cechą życia i twierdzi, że proces nadawania znaczenia jest podstawową formą inteligencji. Uexküll udowodnił, przede wszystkim, że koncepcje biosemiotyki, mimo że wychodzą od fizycznych atrybutów przedmiotów (ich percepcji i działania względem nich), nie mogą być zależne od ich konkretnej fizycznej realizacji. Podobnie, zdaniem Peirce’a¹³⁷

136 J. Hoffmeyer, *The Unfolding Semiosphere*, [w:] *Evolutionary Systems*, Springer, Dordrecht 1998, ss. 281–293.

137 C.S. Peirce, *The essential Peirce: Selected philosophical writings*, N. Houser (red.), Indiana University Press 1998, t. 2.

komunikacja semiotyczna obejmuje znak, przedmiot, do którego odnosi się znak, oraz interpretatora. Aby jednak coś mogło być znakiem, musi być zrozumiane. Znaki muszą być interpretowane względem siebie w kontekście, w przeciwnym razie mogą nawet nie być uznane za znaki. Podstawową zasadą semiotyki Peirce'a jest to, że indeksowe i ikoniczne znaki, a zwłaszcza znaki symboliczne, nie mają żadnego znaczenia w izolacji¹³⁸. Biosemiotyka bardzo poważnie podchodzi do zagadnień związanych z koncepcją znaku, starając się przedstawić nie tylko intuicyjne podejście do idei kształtowania się stanów mentalnych, lecz również biologiczne źródła tego procesu. Rozważając teorię Uexküll'a, która leży u podstaw wszystkich badań biosemiotycznych, nie można nie zauważyć, że interakcje semiotyczne na każdym poziomie mają wpływ na procesy umysłowe.

2.3 Ucieleśnienie i usytuowanie

Ciało jest niezbędnym czynnikiem umożliwiającym poznanie lub myślenie, przez co jego istnienie jest warunkiem koniecznym wszelkiego rodzaju inteligencji. Szczególnie ucieleśnienie procesów mentalnych jest powtarzającym się tematem w naukach kognitywnych i neurobiologii¹³⁹. Wydaje się, że ciało jest potrzebne nawet do takich funkcji, jak myślenie abstrakcyjne i matematyczne¹⁴⁰. Sposób wykonania zadań ma wpływ na powodzenie działania. Operowanie przedmiotami wymaga znacznie mniej skomplikowanej kontroli, jeśli uwzględnimy sposób w jaki odkształcalna tkanka palców ułatwia chwycenie twardych przedmiotów. Powodem, dla którego zadanie staje się łatwiejsze, jest to, że część kontroli neuronalnej jest w rzeczywistości przejmowana przez morfologiczne i materialne właściwości ręki. Dzięki sensorom ciała, sygnały czuciowe są dostarczane do mózgu i stanowią dobrą podstawę do działania. Na przykład: chwytanie drobnych przedmiotów opuszkami palców jest proste, ponieważ jest tam znacznie więcej czujników dotykowych niż w przegubach palców lub w śródręczu. Wykorzystanie

138 *Ibid.*

139 A. Clark, *Being there: Putting brain, body, and world together again...*, *op. cit.*; A. Damasio, *Descartes' error: emotion, reason, and the human brain*, New York 1994; J.N. Deely, *Basics of semiotics*, Indiana University Press 1990; G. Lakoff, M. Johnson, *Philosophy in the Flesh*, New York 1999; F.J. Varela, E. Thompson, E. Rosch, *The embodied mind*, MIT Press 1993.

140 G. Lakoff, R. Núñez, *Where Mathematics Comes From: How the Embodied Mind Brings Mathematics into Being...*, *op. cit.*

materialnych właściwości układu mięśniowo-ścięgnistego pozwala na wykonanie bardzo szybkich ruchów, nawet jeśli układ nerwowy jest zbyt wolny, aby kontrolować wszystkie szczegóły ruchu¹⁴¹. Jednym z aspektów działania mózgu jest propriocepcja, czyli świadomość ruchu i pozycji ciała. Nawet najprostsze ruchy wymagają ciągłego sprzężenia zwrotnego od narządów proprioceptywnych w ciele, mierząc naprężenia mięśni i przemieszczenia warstw komórkowych, w tym również orientację grawitacyjną. Maxine Sheets-Johnstone sugeruje, że zmysł proprioceptywny służy jako cielesna świadomość, zauważając, że każde stworzenie, które może się samo poruszać, odczuwa swój ruch i swój bezruch¹⁴².

Nauki poznawcze przyjmują, że powstawanie pojęć jest związane z konkretnymi programami motorycznymi, które są motywowane przez percepcję i działanie w kontekście doświadczonego zjawiska. Badania przeprowadzone przez Eleanor Rosch dowodzą, że podstawowe pojęcia w umyśle człowieka odnoszą się do rodzajów rzeczy lub działań, z którymi mamy doświadczenie motoryczne i z których możemy tworzyć schematyczne reprezentacje obrazów¹⁴³: stoły, ściany, rowery, budynki, mówienie, chodzenie, spanie itd. Sensoryczno-motoryczna znajomość świata determinuje podstawowe pojęcia. Podstawowe koncepcje mają rdzeń, zależny od podstawowych funkcji organizmu, ale też duża część jest powiązana z percepcją i praktyką działania opracowaną w danym środowisku. Ucieleśnione gesty stają się mentalnymi schematami używanymi w percepcji i rozumowaniu: część-całość, centrum-peryferie, cel-ścieżka, proste-zakrzywione i bliskie-dalekie, cykl, przyległość, ruch, równowaga. Struktury tych gestów wskazują, w jaki sposób ciało jest zorientowane na cel. Kolejną istotną konsekwencją jest to, że te podstawowe doświadczenia cielesne stanowią punkt wyjścia dla aktywności umysłowej. W wyniku doświadczania stanów i rzeczy świata powstaje abstrakcyjna myśl. George Lakoff i Mark Johnson skonstruowali przełomową teorię, czyniąc metaforę cielesną ważnym narzędziem poznawczym, dającym początek

141 Na przykład, gdy stopa w biegu uderza o ziemię, elastyczne rozciągnięcie i odrzut kostki zostaje przejęty przez sprężysty układ mięśniowo-ścięgnisty i nie musi być kontrolowany przez układ nerwowy.

142 M. Sheets-Johnstone, *Consciousness: A natural history*, „Journal of Consciousness Studies”, t. 5, nr 3, 1998, ss. 260–294.

143 E. Rosch, *Principles of categorization*, „Concepts: Core Readings”, t. 189, 1999, ss. 27-48.

metaforom strukturalnym, z których każda opiera się na metaforycznych wyrażeniach¹⁴⁴. Wyobrażenia okazuje się być ważnym narzędziem dzięki użyciu metafor, ale także w budowaniu rozwiniętych modeli pojęciowych w eksperymentach myślowych: wyidealizowanych modeli poznawczych zbudowanych z podstawowych pojęć, schematów i odwzorowań między nimi.

Biosemiotycy również próbują badać ucieleśnioną naturę sfery mentalnej. Stwierdzenie, że procesy mentalne są ucieleśnione zakłada, że naturalnej historii ich powstawania i rozwoju nie można oddzielić od ucieleśnionego życia. Życie mentalne opiera się na intencjonalności cielesnej, przejawiającej się w cyklach percepcji – działania, a więc ostatecznie w biosemiozie. Układ nerwowy rozwinął się jako narzędzie ruchu. Jego zadaniem jest ułatwienie komunikacji między komórkami w różnych częściach się organizmu. Dzięki różnym mechanizmom komórkowym, dane wejściowe mogą odpowiadać za rozwój dopasowania funkcjonalnego między ośrodkowym układem nerwowym a wymaganiami reszty ciała. Odczucia zmysłowe mogą również wpływać na rozbieżne wyniki rozwojowe i powstawanie nowych wzorów w mózgu. Neuronauki ukazują liczne mechanizmy, które mogą wyjaśnić efekty działania systemu sensorycznego w rozwoju behawioralnym¹⁴⁵. Istnieje związek między budową ciała i centralnego układu nerwowego. Wpływ ma na to wiele mechanizmów m.in.: odpowiedzi troficzne mięśni, procesy wywołane aktywnością komórek czuciowych i motorycznych oraz zmiany hormonalne powstałe nawet w odległych częściach ciała. Poprzez fizyczną interakcję żywego systemu z otoczeniem generowane są informacyjne sygnały w różnych kanałach zmysłowych. Odczucie poruszania się jest wywołane m.in. dzięki widzeniu zmian w otoczeniu skorelowanym z naprężeniem mięśni. Obiekty znajdujące się bliżej wydają się poruszać szybciej niż te znajdujące się dalej, co dostarcza organizmowi informacji o odległości. Istnieje więc związek pomiędzy aktywnością neuronową mózgu, morfologią ciała (kształtem i jego właściwościami materialnymi) i interakcją ze środowiskiem, która jest wykorzystana do osiągnięcia określonych zadań. Różnorodne procesy komórkowe

144 Np. Strukturalna, konceptualna metafora „Wiedza to widzenie” znana w wielu językach wywołuje długą serię różnych wyrażen, takich jak „oświecenie”, „Czy nie widzisz, co wyjaśniam?”, „Przyjrzyj się temu problemowi”. Patrz G. Lakoff, M. Johnson, *Metaphors we live by*, University of Chicago Press 2008.

145 D. Purves, *Body and brain: a trophic theory of neural connections*, Cambridge 1988, ss. 19–20.

prowadzą do przetrwania odpowiednich neuronów i wpływają na różnicowanie się form systemu nerwowego.

Merleau-Ponty uznał, że pierwotną świadomością nie jest „myślę, że”, ale „ja mogę”¹⁴⁶. Zapoczątkowanie ruchu współlistnieje z motywacją do jego wykonania i musi zawierać się w zakresie możliwości poruszania się specyficznym dla gatunku. Jest to podstawa potencjału „ja mogę”, która jest podporządkowana proprioceptywnym stanom i możliwościom działania podmiotu. Oznacza to, że ciało jako nośnik doznań, pozwala odczuć cechy zewnętrzne środowiska, a także same wrażenia ruchowe i dotykowe płynące z wnętrza ciała. Umożliwia to poczucie obiektu, ale przede wszystkim potencjalne możliwości ciała i „ja mogę” zamienić się mogą w „ja robię”. Merleau-Ponty zwraca uwagę, że zachowanie organizmu w *Umwelcie* jest bardziej podstawowe niż świadomość, która jest tylko jedną ze specjalnych form tego zachowania¹⁴⁷. Powiązanie między organizmem a otoczeniem jest podstawowym warunkiem powstania świadomego funkcjonowania. Czuć i ruch (percepcja i działanie) umożliwiają organizmowi aktywne odkrywanie i rozumienie świata. Narządy zmysłów w kręgach funkcjonalnych umożliwiają organizmowi zwiększenie precyzji działania. Oznacza to, że zwierzę może odróżnić swoją własną pozycję przestrzenną dzięki właściwemu neuronowemu systemowi propriocepcji, który ułatwia kontrolę zwrotną zachowania względem właściwego świata percepcyjnego i świata zachowań. Dopiero gdy ciało jest samo w sobie postrzegane, świat percepcyjny staje się możliwy jako pojmowany, reprezentowany świat umysłu dający możliwość i wybór dla inteligentnego działania w nim.

W tym kontekście interesująca jest interpretacja koncepcji Edmunda Husserla przez Merleau-Ponty’ego. Problemem Husserla jest, jak twierdzi Merleau-Ponty, znaleźć miejsce dla natury w filozofii refleksji¹⁴⁸. Natura interpretowana jako składająca się z czystych rzeczy, jest korelatem czystej świadomości, umożliwiającej postrzeganie środowiska życia, *Lebenswelt*. Merleau-Ponty przedstawia „świat ciała” Husserla jako silnik sensoryczny *Umweltu*. Tłumaczy to jako sposobność postrzegania przedmiotów nie w sposób oderwany, ale biorący pod uwagę możliwości motoryczne ciała podmiotu: „Obiekt wydaje

146 M. Merleau-Ponty, *Fenomenologia percepcji*, M. Kowalska, J. Migasiński (tłum.), Fundacja Aletheia 2001, s. 156.

147 *Ibid.*, ss. 239, 393.

148 M. Merleau-Ponty, *Nature: Course Notes from the Collège de France*, Northwestern University Press 2003, s. 72.

mi się funkcją ruchu mojego ciała”¹⁴⁹. Ciało jest odpowiednie dla bycia w świecie, a jednocześnie dla percepcji. Świat staje się niejako częścią ciała podmiotu, dając mu możliwość orientacji, nie tylko w czasoprzestrzeni, ale we wszystkich skalach normatywnych. Fenomenologiczna interpretacja teorii Uexküll’a przez Merleau-Ponty’ego ukazuje *Umwelt* jako aktualności egzystencji. Zachowanie w takim *Umwelcie* nie może być zrozumiane z chwili na chwilę, ale tylko jako znacząca całość istnienia w czasie. U wyższych zwierząt kręgi funkcjonalne poddają symbole pod interpretację, nie dążąc wyłącznie do osiągnięcia pierwotnych celów. Takie podejście wydaje się być kluczowe dla zrozumienia sposobu kształtowania się inteligentnych zachowań. Symbole wskazują na aktualną sytuację i stan świata, na przyszłe postrzeganie, pozwalają zwierzęciu ustosunkować się do zdarzeń i rozważyć możliwości działania. Działanie za pomocą symboli pozwala organizmom wykonywać niewrodzone skomplikowane działania.

Krąg funkcjonalny Uexküll’a stanowi podstawowy plan ciała sensoryczno-motorycznego. W tej teorii ciało jest pojmowane jako urządzenie semiotyczne, które jest w stanie postrzegać otoczenie za pomocą znaków i działać poprzez znaki¹⁵⁰. Uexküll dopuścił istnienie nieokreślonych neutralnych obiektów w *Umwelcie*¹⁵¹, które okazały się niezbędne dla opisu wzrostu złożoności *Umweltu*. Postrzeżenie obiektów neutralnych jest warunkiem wstępnym uczenia się, ponieważ umieszczenie ich w działaniu już istniejących kręgów funkcjonalnych powoduje konieczność rozwinięcia kręgu. Zdolność adaptacji zakłada percepcję obiektów neutralnych, które początkowo nie są istotne, lecz powodują zachwianie w dopasowaniu organizmu i środowiska. Rozwój działania kręgu funkcjonalnego jest wywołany przez rosnącą złożoność semiotyczną, która wywołuje konieczność jego aktualizacji. Sytuacja wzrastającej presji środowiskowej powoduje osiąganie optymalnego napięcia w uprzednio zdefiniowanym kręgu funkcjonalnym, poprzez dopasowanie go do pojawiających się zmian.

Uczeń Uexküll’a Konrad Lorenz również rozważał ten problem przyglądając się pojęciu instynktu. Według Lorenza instynkt jest wrodzoną serią działań, która wymaga pewnego wyzwalacza pojawiającego się w środowisku powodującego aktualizację

149 *Ibid.*, s. 74.

150 M. Heidegger, *Die Grundbegriffe der Metaphysik: Welt, Endlichkeit, Einsamkeit*, Frankfurt am Main 1983, s. 261.

151 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten & Bedeutungslehre...*, op. cit., ss. 92–93.

organizmu¹⁵². Wydawać by się mogło, że jest to czysto mechanicystyczna idea, jednak wyobrażenie bezcelowości i imprintingu wskazuje na coś zupełnie odmiennego. Instynkt u Lorenza jest bezcelowy i w wielu przypadkach niekompletny, przez co wymaga wypełnienia ze środowiska. Dopełnienie się instynktu jest możliwe dzięki konstrukcji *Umweltu*¹⁵³. Jak twierdzi Merleau-Ponty: bezcelowość i otwartość instynktu umożliwia jego wyobrażeniową reinterpretację jako podstawę do formowania się symboli¹⁵⁴. Instynktowne działania mogą zostać odcięte od ich możliwego celu i traktowane jako symbol zupełnie różnych zjawisk w komunikacji zwierzęcej. Zarówno Uexküll jak i Lorenz uznają, że istnieje głębokie wzajemne powiązanie żywych organizmów z ich otoczeniem.

2.4 Autopoeza i autonomia

Ucieleśnienie łączy się z pojęciami autopoezy i autonomii oraz umiejętnością utrzymania homeostazy ciała. Zdolność zachowywania stałych parametrów ciała oraz skoordynowanie go ze środowiskiem, jest warunkiem autonomicznego utrzymania samoorganizacji oraz wynika z aktywnego i stałego rekompensowania uszkodzeń. Żywe organizmy kształtują się odśrodkowo od jednej komórki, która najpierw rozwija się w gastrulę, a później w bardziej skomplikowane narządy¹⁵⁵. Stanowi to podstawowe kryterium autonomii Uexkülla¹⁵⁶. Te dociekania interpretują autonomię jako właściwość wyższego poziomu, powstającą w wyniku samorozwoju i adaptacji żywego organizmu. Autonomiczny organizm jest częścią pewnego środowiska, w którym funkcjonuje. Wykorzystuje umiejętność odczuwania tego, co dzieje się w środowisku zewnętrznym i wykorzystuje to w celu zaspokajania własnych potrzeb. Umie też wykorzystać nabyte uprzednio doświadczenia. Między organizmem i jego środowiskiem następuje sprzężenie zwrotne, które umożliwia powstanie autonomicznej, uformowanej doświadczeniami

152 K. Lorenz, *The Foundations of Ethology*, Springer Science & Business Media 1981, s. 147.

153 M. Merleau-Ponty, *Nature...*, *op. cit.*, ss. 193–194.

154 *Ibid.*, s. 195.

155 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten & Bedeutungslehre...*, *op. cit.*, ss. 116–117.

156 N.E. Sharkey, T. Ziemke, *Mechanistic versus phenomenal embodiment: Can robot embodiment lead to strong AI?*, „Cognitive Systems Research”, t. 2, nr 4, 2001, ss. 251–262.

jednostki. Organizm potrafi identyfikować znaki i nadawać sens wskazówkom środowiskowym, co umożliwia mu aktywność będącą adekwatną odpowiedzią na zmiany, zgodną z jego własnymi wymaganiami.

Humberto Maturana i Francisco Varela odróżnili organizację systemu od jego struktury¹⁵⁷. Organizacja, podobnie jak pojęcie planu budowy Uexküll'a (*Bauplan*), oznacza te relacje, które muszą zachodzić między składowymi systemu, aby mógł on należeć do określonej klasy. System autopoietyczny to szczególny rodzaj układu homeostatycznego, dla którego podstawą jest utrzymanie stałych wartości układu i jego organizacji. Natomiast struktura systemu oznacza „składowe i relacje, które faktycznie tworzą określoną całość i sprawiają, że jego organizacja jest rzeczywista”¹⁵⁸. Zatem struktura systemu autopoietycznego to konkretne urzeczywistnienie faktycznych składowych i faktycznych relacji pomiędzy nimi. Jego organizację konstytuują relacje pomiędzy tymi składowymi, co definiuje go jako całość określonego rodzaju. Relacje te to sieć procesów, które poprzez budowę, transformacje i niszczenie wytwarzają własne części składowe. Działania te stale odnawiają i realizują procesy, które te składowe wytwarzają. Żywy organizm jest rzeczywiście autonomiczny w tym sensie, że jest systemem autopoietycznym, jednością, którego składniki i aktywność warunkują samorealizację jako sieć procesów organizmu w przestrzeni, w której egzystuje¹⁵⁹.

Sprzężenie percepcji z działaniem ma dla systemów autonomicznych kluczowe znaczenie. Dzięki ucieleśnieniu otrzymywane informacje są natychmiast zwracane i pozwalają kontrolować funkcjonowanie. Sygnały wejściowe i wyjściowe odgrywają istotną rolę, ponieważ niosą informacje wpływające na całość funkcjonowania systemu. Interpretacja następuje w oparciu o zgromadzoną wiedzę¹⁶⁰. Podstawą autonomii jest samoorganizacja, umożliwiająca organizmowi uczenie się i dopasowywanie działania wewnętrznych mechanizmów do aktywności w świecie zewnętrznym. Zwierzęta zdobywają tę zdolność w wyniku adaptowania się do środowiska. Samoorganizacja jest prywatnym sposobem wykorzystania znaków i tworzenia reprezentacji. Dlatego znaczenie powstałe w wyniku interpretacji znaków też pozostaje prywatne. Interpretacja zmienia stan

157 H.R. Maturana, F.J. Varela, *The tree of knowledge...*, op. cit.

158 H.R. Maturana, *Autopoiesis and Cognition: The Realization of the Living*, Springer Science & Business Media 1980.

159 *Ibid.*, s. 78.

160 S. Harnad, *Computation is just interpretable symbol manipulation; cognition isn't...*, op. cit.

zarówno kręgu funkcjonalnego jak i umysłu zwierzęcia. Żywy organizm, w zależności od kontekstu, stanu i wiedzy, interpretuje znaki percypowanego świata. Sygnał pochwycony przez receptory może być wzmacniany dzięki wiedzy o sygnale efektorowym – wyniku uprzednio wykonanego działania. Sprzężenie pozwala dostosować zachowanie do środowiska, zwiększając pozytywne lub minimalizując negatywne wzmocnienia. Wzmocnienia są tu bodźcami zwiększającymi szansę odpowiedzi systemu. Krąg funkcjonalny, w obrębie, którego przebiega to działanie, zmienia się dynamicznie dostosowując do sygnałów, które otrzymuje.

Autonomia w silnym sensie biologicznym zakłada takie własności istot żywych jak umiejętność przetrwania, regeneracji uszkodzeń i reprodukcji. To znacząco różni je od sztucznych systemów i jest podnoszone przez krytyków programu sztucznej inteligencji. W trakcie ewolucji żywych organizmów wykształciły się procesy przetwarzania znaków, kręgi funkcjonalne i *Umwelt*. Wchodzenie w typowe dla gatunku semiotyczne interakcje z otoczeniem jest wynikiem wielu przekształceń i tysięcy lat kształtowania się tych relacji. Żywe organizmy są funkcjonalnie doskonałymi projektami. Ich spójne i celowe zachowanie da się często wyjaśnić dość prostymi potrzebami interakcji ze środowiskiem. Obserwacja organizmu pozwala zauważyć, że złożoność jego zachowania zależy od złożoności sygnałów w środowisku, w którym żyje. Zachowanie żywej istoty wykazuje spójność, mimo wielu mechanizmów działania, tworzących funkcjonalne kręgi. Współpracując z innymi lub będąc aktywnym podmiotem w środowisku przejawia inteligentne zachowanie, wywołane przez interakcję wszystkich części. Może wykorzystać swoją pamięć i dostosować zachowanie do zmieniającego się środowiska. Samoorganizacja pozwala układowi reagować na bodźce płynące ze środowiska w sposób, który nie jest zaprogramowaną czy mechaniczną reakcją.

2.5 Semiotyczne rusztowania w semiosferze

Dzięki autonomicznej aktywności organizmu, zachodzą działania, które sprawdzają się w wypadku niepełnej informacji. Autonomicznie funkcjonujący system może uzupełnić braki wiedzy, wykorzystując dane zewnętrzne. Pozwala mu na to sieć semiotycznych interakcji, za pomocą których poszczególne komórki, organizmy lub populacje kontrolują

swoje działania¹⁶¹. Rusztowania semiotyczne (*semiotic scaffolding*) zapewniają organizmom precyzyjne działania poprzez zapewnienie wydajnej interakcji z kluczowymi znakami pojawiającymi się w dynamicznych sytuacjach (np. polowanie lub łączenie się w pary). Termin ten jest rozumiany w sensie ogólnym jako byt lub proces, który wspiera inny, podstawowy proces, a tym samym wzmacnia stabilność, funkcjonowanie lub zakres możliwości procesów poznawczych¹⁶². Wzory działania wyłaniają się jako inteligentne zachowanie systemu, jako wynik różnych autonomicznych działań zależnych od środowiska, w miejsce scentralizowanego systemu podejmowania decyzji opartego o wewnętrznie reprezentowane kierunki działań lub cele. Horst Hendriks-Jansen stwierdził, że interaktywne zachowanie może być wyjaśnione w kategoriach generatywnych, ponieważ nie ma żadnych wewnętrznych reguł pozwalających przewidzieć rodzaje zachowań, które się pojawiają¹⁶³. Andy Clark natomiast zasugerował, że inteligentne stworzenia nie będą przechowywać ani przetwarzać informacji w kosztowny sposób, kiedy będą mogły w tym celu użyć struktur środowiska¹⁶⁴.

Znaczenie rusztowania zostało już podkreślone przez Lwa Wygockiego, który opisywał rozwój dziecka jako zdobywającego doświadczenia przy wsparciu zewnętrznych struktur¹⁶⁵. Nowe umiejętności mogą być przenoszone społecznie z opiekuna na potomka poprzez mimikrę lub naśladownictwo lub za pomocą rusztowania, gdy bardziej sprawny dorosły manipuluje interakcjami podopiecznego ze środowiskiem w celu rozwijania nowatorskich umiejętności. Użycie rusztowania polega na zmniejszeniu dystansu, zaznaczeniu najważniejszych cech zadania, zmniejszeniu liczby stopni wykonywanego planu i umożliwieniu doznania końca lub rezultatu działania, zanim zyska niezależną zdolność poznawczą lub fizyczną do poszukiwania i osiągnięcia celu. Koncepcja

161 K. Kull, *Evolution, choice, and scaffolding: semiosis is changing its own building*, „Biosemiotics”, t. 8, nr 2, 2015, ss. 223–234.

162 F. Stjernfelt, C. Emmeche, K. Kull, *Reading Hoffmeyer, rethinking biology*, Tartu University Press 2002, s. 29.

163 Pojęcie generatywności opisuje z system autonomiczny, który kreuje i wykorzystuje nowe unikatowe zachowań bez oparcia na zewnętrznych danych tego systemu. Por. H. Hendriks-Jansen, *Catching ourselves in the act: Situated activity, interactive emergence, evolution, and human thought*, Cambridge 1996, s. 9.

164 A. Clark, *Being there: Putting brain, body, and world together again...*, *op. cit.*, s. 46.

165 Może to być np.; wsparcie fizyczne, podczas nauki chodzenia lub pływania albo językowe podczas nauki mowy. L.S. Vygotsky, *Thought and language*, „Annals of Dyslexia”, t. 14, nr 1, 1964, ss. 97–98.

rusztowania opisuje sposób, w jaki żywe organizmy używają zewnętrznych struktur w celu uproszczenia zadań. Przykładem potężnego rusztowania semiotycznego człowieka jest język naturalny i technologia informacyjna. Ludzka wiedza może być zapisana w książkach i przekazana. Człowiek może opierać się na tym, co zostało już ustalone i zapisane: pomysły w jednym tekście polegają (bezpośrednio lub pośrednio) na innych tekstach. Sieć World Wide Web, wypełniona tekstem, obrazami, dźwiękiem i wideo, sprawiła, że sieć wiedzy i idei stała się bardziej przejrzysta i łatwiej dostępna¹⁶⁶. Jest to przykład zasady „taniego projektu”, który pozwala usunąć pojęcie centralnego modułu, który posiada wszystkie informacje dostępne w dowolnym miejscu w samym systemie. Problem z centralnym modelem polega na tym, że jest on całkowicie niepraktyczny w złożonych i nieprzewidywalnych zdarzeniach w dynamicznym środowisku.

Rusztowania są stawiane w celu wzniesienia budynku i w znacznym stopniu ograniczają i określają sposób budowania wieżowca, tak samo semiotyczna kontrola działań biologicznych ogranicza i określa czas i sposób, w jaki zachodzi aktywność. Konceptualizacja i analiza semiotycznych mechanizmów rusztowania działających na różnych poziomach w systemach naturalnych stanowi rdzeń badań biosemiotyki. Semiotyczne rusztowania mogą przybierać różne formy, ale ich podstawową właściwością pozostaje skupienie zachowania organizmu na ograniczonym repertuarze możliwości lub w kierowaniu zachowaniem tak, by zrealizować sekwencję działań. Receptor komórki jest dostrojony do otwarcia się wtedy i tylko wtedy, gdy trafi na odpowiedni bodziec, np. oko płaza powstaje w wyniku chemicznych interakcji między nowo powstałym pęcherzykiem optycznym i zarodkową warstwą ektodermy. Induktor chemiczny wytwarzany przez pęcherzyk optyczny jest używany jako zewnętrzne rusztowanie tego działania¹⁶⁷. Ten przykład pokazuje, jak ważnym aspektem rozwoju jest zdolność pojedynczych komórek do zmiany ich wewnętrznych ustawień pod wpływem czynników zewnętrznych lub nowych wskazówek molekularnych.

Działania oparte na rusztowaniach stają się bardziej skomplikowane, gdy cykl życiowy organizmów staje się bardziej złożony lub gdy zwierzęta zaczynają się angażować się w złożone procesy społeczne. Mechanizmy rusztowań semiotycznych zależą od

166 Tematyka ta została szeroko opisana w: A. Clark, *Natural-born cyborgs. Minds, technologies, and the future of human intelligence*, Oxford University Press 2004.

167 K. Kull, *Catalysis and scaffolding in semiosis*, [w:] *The Catalyzing Mind*, Springer 2014, ss. 111–121.

aktów interpretacji, które zawsze zakładają błąd. Zależą też od zdolności przewidywania i przygotowania do ważnych sytuacji i zdarzeń w cyklu życia danych zwierząt. Genowe rusztowania działają poprzez kontrolę i montaż białek, które są następnie przekształcane w schematy odzwierciedlające potrzeby organizmu¹⁶⁸. Mechanizmy te działają wystarczająco dobrze, dopóki repertuar zachowań zwierząt jest ograniczony do instynktownie uruchamianych reakcji na dające się przewidzieć zdarzenia. Jednak zwierzęta o dużych mózgach, jak ptaki i ssaki, poza instynktownymi odruchami, są zależne również od procesów uczenia się. Takie procesy wspomagane są genetycznie zapewnionymi preferencjami behawioralnymi, lecz umiejętność uczenia się musi być nieodłącznym elementem elastyczności zachowania, a tym samym przeniesienia kontroli behawioralnej z poziomu genomu do poziomu mózgu. To wprowadza potrzebę używania kolejnych mechanizmów rusztowania.

Żywe organizmy mogą wykorzystywać rusztowania, aby odnosić się do bezpośredniego otoczenia, do symbolizowania, rozumowania, używania schematów. Zwierzę czerpie korzyści ze zdolności pozyskanych w trakcie filogenezy lub poprzez ontogenezę, by ustalić prawidłowości w otoczeniu i odpowiednio ukierunkować działania. Sieweczka obrożna, która udaje, że ma złamane skrzydło, by odwrócić uwagę łasicy od gniazda, korzysta z wpływu na rusztowanie semiotyczne drapieżnika, który daje się zwieść fałszywym znakiem. Ptak oszukuje ssaka, ponieważ genetycznie lub ontogenetycznie zdobył wiedzę o tym, jak drapieżnik zinterpretuje pozorną relację. Może się jednak zdarzyć, że łasica nie da się zwieść, co dowodzi, że jej reakcje nie są ściśle deterministyczne. Drapieżnik może błędnie interpretować znak lub nie. Oznacza to również, że istotą znaku jest jego relacyjny charakter wywoływania świadomości czegoś, co nie jest tym samo w sobie. Fakt ten implikuje pełną peirceańską triadę znaku, który może powstać w dowolnym systemie zdolnym do autonomicznej aktywności antycypacyjnej.

Rusztowania semiotyczne są związane z interakcją wnętrza i otoczenia organizmu. Występują w semiosferze żywych organizmów, które polegają na komunikacji z otaczającym je światem: dźwiękami, zapachami, ruchem, kolorem, kształtem, polem elektrycznym, falami wszelkiego rodzaju, chemicznymi sygnałami, dotykiem itd., i zapewniają im aktywność poznawczą. Semiosfera jest wynikiem i warunkiem rozwoju

168 J. Hoffmeyer, *The semiome: From genetic to semiotic scaffolding*, „Semiotica”, t. 2014, nr 198, 2014, ss. 11–31.

kultury a przez analogię z biosferą Władimira Wiernadskiego, jest odczytywana jako całość organiczna żywej natury, a także miejsce kontynuacji życia¹⁶⁹. Koncepcja rusztowań semiotycznych semiosfery nadaje wymiar semiotyczny procesom życiowym, podkreślając ich przynależność do wspólnego świata aktywności znakowej. W swojej teorii Uexküll przedstawił procesy budowania własnych niezależnych relacji ze światem jako niezbędny warunek autonomii działającego systemu. Gatunki mają ograniczony dostęp do tej semiosfery, ponieważ zdolność gatunku do interpretacji potencjalnych sygnałów w jego otoczeniu ewoluowała, aby dopasować się do określonej niszy ekologicznej. Nisza ekologiczna obejmuje wszystkie informacje i wskazówki, które muszą być poprawnie interpretowane przez organizm. Liczba cech świata, które mogą w niektórych sytuacjach stać się istotnymi wskazówkami dotyczącymi zachowań organizmu, jest nieskończona lub większa niż liczba cech, z którymi organizm faktycznie oddziałuje fizycznie. W ten sposób ptak, aby zapewnić sobie najkorzystniejsze działanie, musi zauważać możliwość pożywienia się lub schronienia, ale także wykorzystywać rusztowania semiotyczne będące wzorami dźwięków, kierunków i prędkości wiatru, różnicami temperatury powietrza lub wiatru, zmianami natężenia i długości fal światła itd. W środowisku zwierzęcia w ciągu jego życia następuje wiele zmian. Dlatego aspekt semiosfery i rusztowań semiotycznych wskazuje na ogromną ilość możliwości działań bądź adaptacji, której automatyczna reaktywność nie jest w stanie w pełni wytłumaczyć. Organizmy nie mogą reagować tylko biernie lub instynktownie na stany i zdarzenia. Zamiast tego postrzegają, interpretują i działają w środowisku w sposób, który twórczo i nieprzewidywalnie zmienia wszystkie ewolucyjne i selektywne ustawienia.

169 Y. Lotman, *Universe of the Mind. A Semiotic Theory of Culture...*, op. cit., ss. 12–126.

Łotman, który wprowadził pojęcie semiosfery, użył jej jako analogii z koncepcją biosfery Władimira Wiernadskiego. Mimo, że w pismach Łotmana semiosfera jest pojęciem pierwotnie związanym z kulturą, przedstawia on semiozę jako podstawowy mechanizm działania łączący całą przestrzeń semiotyczną. Należy dodać, że koncepcja biosfery Wiernadskiego nie przetrwała w tym sensie i jest używana jako określenie ekosystemu obejmującego ziemię i zamieszkujące ją organizmy.

2.6 Podmiot jako „semiotyczne ja”

Dzięki Uexküllowi zysaliśmy podstawowe narzędzia do badania procesów semiozy zachodzących w świecie organicznym. Teoria stworzona przez biologa doskonale tłumaczy najbardziej podstawowe mechanizmy inteligencji, jednak celem tej pracy jest ukazanie jej wyższych form, szczególnie tych, które są kwestionowane przez oponentów sztucznej inteligencji. Należy się zatem przyjrzyć bardziej skomplikowanym przejawom inteligencji. Istotnym wkładem biosemiotyki do SI byłoby dostarczenie wyjaśnień naukowych, które uwzględniałyby także kognitywno-semiotyczne mechanizmy działania i wyjaśnienie, które uwzględniałyby istnienie stanów mentalnych. Pozwoliłoby to zmienić metaforę umysłu w pełnowymiarową, objaśniającą koncepcję naukową. Wyjaśnienia te muszą zatem obejmować: rozpoznawanie, wybór działania, pamięć, relacje między kodami, kategoryzowanie i komunikację. Szczególnie ważne pytanie dotyczy powstawania podmiotowości i intencjonalności jako ogólnych właściwości procesów semiocentrycznych¹⁷⁰.

Sposób biosemiotycznego traktowania podmiotu, a także jego podejścia do życia, wynikają z zasady, że wybór działania zależy od podmiotowości jednostki, co pozwala traktować go jako indywidualną tendencję organizmu. Dlatego relacja podmiotu wobec świata, którego efektem poznania przypisywana jest niepowtarzalność, czyli subiektywizm, związana ze świadomością, wolą, osobowością, rzeczywistością i prawdą. Podmiotowość to jakość lub stan, który zakłada istnienie subiektywności, która nie jest wolna od uprzedzeń, interpretacji, uczuć i wyobrażeń jednostki¹⁷¹. Podmiot oznacza taką jednostkę, która posiada świadome doświadczenia, takie jak perspektywy, uczucia, przekonania i pragnienia. Niektóre informacje, idee, sytuacje lub rzeczy fizyczne są uważane za prawdziwe tylko z perspektywy podmiotu lub podmiotów. Podmiot również działa lub sprawuje kontrolę nad innymi obiektami¹⁷². W związku z tym wydaje się, że inteligencja i podmiotowość istnieją w systemie wzajemnych odniesień i nie można ich

170 K. Kull, P. Torop, *Biotranslation: Translation between umwelten*, „Readings in Zoosemiotics”, t. 8, 2011, ss. 411–425.

171 M. Czarnocka, *Podmiotowość w neokantyzmie*, „Filozofia i Nauka. Studia filozoficzne i interdyscyplinarne”, t. 3, 2015, ss. 119–139.

172 A. Strazzoni, *Subjectivity and individuality: Two strands in early modern philosophy: Introduction*, „Societate si Politica”, t. 9, nr 1, 2015, ss. 5–9.

rozważać oddzielnie. Jest to specyficznie semiotyczny wymiar funkcjonowania, które ma szansę stać się wartościową nowością podejścia biosemiotycznego w programie badań SI.

W programie sztucznej inteligencji charakter pojęć takich jak podmiotowość, ja, jaźń, samoświadomość itd. nie odgrywały większej roli. O ile w ogóle poświęcano im uwagę, były traktowane w sposób metaforyczny, co musiało przynieść skutek w postaci ostrej krytyki. Już na wstępie, brak definicji i zrozumienia tych pojęć powodował niemożność ich adekwatnej implementacji w sztucznych systemach a tym samym dodatkowo utrudniało badania. Pojęcia ucieleśnienia i usytuowania uznano zatem za wystarczające, by opisać interakcje sztucznego systemu ze światem rzeczywistym. Inteligentne funkcjonowanie żywych organizmów jest nierozdzielnie związane z powyższymi pojęciami. Podmiotowość jest używana jako wyjaśnienie tego, co informuje, kształtuje i wpływa na osądy. Jest zbiorem percepcji, doświadczeń, oczekiwań, osobistego lub kulturowego zrozumienia i przekonań charakterystycznych dla podmiotu. Pierwsze próby wprowadzenia pojęcia podmiotowości w dyskursie biologicznym wiązały się z koniecznością przezwyciężenia nieadekwatności obiektywizmu i eksternalizmu, nieodłącznie związanego z tradycyjnym paradygmatem nauk przyrodniczych¹⁷³.

Biosemiotyczne pojęcie „ja” jest złożoną konstrukcją uwzględniającą elementy biologii i historyczny rozwój procesów pojęciowych¹⁷⁴. System biologiczny jest konstrukcją uwzględniającą rozpoznawanie czynników i procesów. Już na najniższym stopniu organicznym, komórka (zazwyczaj postrzegana jako podstawowa jednostka organiczna) podtrzymuje swój metabolizm poprzez pozyskiwanie energii i substancji pobieranych z otoczenia. Musi zatem odróżniać środki zawierające użyteczną energię od szkodliwych, które mogą zniszczyć funkcjonowanie komórki. Podstawowa granica między komórką a nie-komórką, lub własnym ja i nie-ja, jest granicą tworzoną przez membranę komórki, która jest podstawą procesów rozpoznawania. Istnieje kod lub system kategoryzacji, w którym komórka rozpoznaje substraty jako użyteczne, szkodliwe bądź

173 F.S. Rothschild, *Creation and Evolution: A Biosemiotic Approach*, Transaction Publishers 2000; J. von Uexküll, *Theoretische biologie*, Berlin 1928; J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten & Bedeutungslehre...*, *op. cit.*

174 D. Favareau, *Beyond self and other: On the neurosemiotic emergence of intersubjectivity*, „Sign Systems Studies”, t. 30, nr 1, 2002, ss. 57–99.

neutralne¹⁷⁵. Ta granica między komórką a otoczeniem jest informacyjną różnicą dla organizmu, a skomplikowany układ molekuł działa jako semiotyczne mechanizmy wspierające pracę, aby utrzymać różnicę między wewnątrz a zewnątrz. Z historycznej perspektywy inteligencja jest natomiast rozwiniętą formą adaptacji, która pozwala podmiotowi podejmować aktywność poznawczą i zdobywać wiedzę o środowisku poprzez dostosowanie się. Zakłada powstanie równowagi między poznaniem i przyswajaniem informacji środowiskowych i samoregulacją umysłu. Istotą inteligencji jest aktywność podmiotu przejawiająca się w procesie strukturalizacji mechanizmów poznawczych, które widoczne są na najniższym poziomie komórkowym.

Tematyka podmiotowości była wielokrotnie poruszana przez Uexküllą. Podmiot to istota, która jest związana z otoczeniem, nie tylko poprzez ciało i jego zmysły, lecz również dzięki percepcji i aktywności, poprzez ruchy w przestrzeni i w czasie. Każdy żywy organizm ma swoją subiektywną przestrzeń i własny subiektywny czas. Wszelkie zachowanie zwierząt Uexküll tłumaczył procesami w subiektywnym świecie podmiotu¹⁷⁶. Świat, jeśli ma charakter informacyjny, nie może być neutralny, bezstronny i wolny od jakiegokolwiek perspektywy. Świat taki, jaki dla niego istnieje – co sugeruje pojęcie *Umwelt* – zawiera jedynie znaczące relacje lub znaki, które wykorzystuje dla własnych celów. Podmiotowość oznacza tu odniesienie do indywidualnego i wcielonego punktu widzenia. Taka strategia badawcza miała stanowić alternatywę dla podejścia mechanistycznego w biologii¹⁷⁷. Propozycja Uexküllą skłaniała do spojrzenia na to, co jest istotne dla zwierzęcia, jednak może być użyta w celu przyjrzenia się działaniom, poprzez które zwierzę modeluje swój własny świat i siebie samego jako indywidualium.

Teoria kręgu funkcjonalnego Uexküllą określa przesłanki i podstawę do potwierdzenia podmiotowego charakteru kształtowania się inteligentnych zachowań. Znaki percepcyjne są zunifikowane i rzutowane z powrotem na świat zewnętrzny jako cechy obiektów, które są traktowane jako percepcyjne wskazówki¹⁷⁸. Przedmioty świata mogą istnieć, ponieważ podmiot przypisuje im cechy, które mogą spełniać pewne funkcje lub działania i służą jego potrzebom. W ramach teorii biosemiotycznej podmiotowość jest związana się z relacjami z

175 C. Emmeche, K. Kull, *Towards a semiotic biology: Life is the action of signs*, World Scientific 2011, s. 16.

176 J. von Uexküll, *Theoretische biologie...*, op. cit., s. 9.

177 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten & Bedeutungslehre...*, op. cit.

178 J. von Uexküll, *Theoretische biologie...*, op. cit.

obiektami, które są nośnikami znaczeń. Nie może istnieć pojedynczy i ostateczny model postrzeganego i działającego obiektu, ponieważ bodziec percepcyjny może być zastąpiony przez kilka funkcjonalnych obrazów (efektorowych sygnałów) w zależności od nastroju i potrzeby podmiotu. Istotą inteligencji jest wzrost złożoności struktur poznawczych. Sposób opanowania środowiska przez żywy organizm jest zależny od zwierzęcia, co wskazuje na pewien stopień autonomii interpretacji zwierzęcej aktywności. Rozwój zachowań inteligentnych musi opierać się na rozwoju procesów organizacji wiedzy, która jest integrowana w struktury poznawcze. Każda jednostka dysponuje zestawem zachowań, które pozwalają radzić sobie z trudnymi sytuacjami. Wpływ na to ma wiele cech podmiotu, lecz przede wszystkim jego temperament i osobowość. Niezwykle ważna jest również adaptacja podmiotu, która pozwala mu dostosować się do środowiska i sprostać jego oczekiwaniom. Inteligencja w perspektywie biosemiotycznej jest mocno rozwiniętą postacią adaptacji, podczas której podmiot sam rozpoczyna aktywność poznawczą i zdobywa wiedzę o swoim świecie w wyniku przystosowywania się do środowiska.

Sebeok zaczął rozwijać semiotykę podmiotu jako „semiotyczne ja”¹⁷⁹. Pojęcie „semiotycznego ja” w rozumieniu Sebeoka, jest rezultatem semiotycznej aktywności zwierzęcia lub człowieka, unikalną konfiguracją znaków, która nie jest sztywno ustalona, chociaż zależy od fizycznej rzeczywistości. Szczegółowe zrozumienie tego problemu może zapewnić podkreślenie semiotycznej natury procesów, składających się na tworzenie podmiotowości:

- a) podmiot jest wynikiem procesów semiotycznych. Wynika to z modelowania opartego na różnych kodach dostępnych w żywym organizmie;
- b) jest zakorzeniony w biologicznych procesach semiotycznych. Na najbardziej podstawowym poziomie wywodzi się z funkcjonowania układu immunologicznego i innych procesów, dzięki którym organizm-ja odróżnia się od otoczenia;
- c) wyłania się ze złożonych lub zbiorowych procesów na różnych poziomach organizacyjnych;
- d) jest bytem dynamicznym, który pozwala na rozwój jaźni, rozpoznanie siebie, a nawet wirtualną projekcję jaźni, która uwzględnia perspektywę środowiskową;
- e) ma warstwową strukturę, która łączy procesy biologiczne, kulturowe (językowe, narracyjne) i semiotyczne co umożliwia ich analizę w tych samych ramach koncepcyjnych.

179 T.A. Sebeok, *The semiotic self*, [w:] *A Sign is Just a Sign*, Bloomington 1979, ss. 263–267;
T.A. Sebeok, *The semiotic self revisited...*, *op. cit.*

Wydaje się, że ta koncepcja ma potencjał do opracowania szerszego opisu podmiotu semiotycznego i może okazać się istotna dla badania inteligencji. Pozwala rozwijać perspektywę środowiskową, podkreślając relacyjny charakter przedmiotu i jego otoczenia. Jest także odpowiednia do opisanego jednostkowej inteligencji, która jest właściwa jednostce w grupie pojedynczych osobników. Należy więc założyć, że „semiotyczne ja” przeżywa rzeczywiście subiektywność lub doświadczenie pierwszoosobowe. Biosemiotyczne rozszerzenie pojęcia podmiotowości na wszystkie żywe systemy opiera się na myśleniu o nich w kategoriach umiejętności podmiotu, szczególnie kompetencji semiotycznej i zdolności rozumienia otoczenia. Biosemiotyka stara się usunąć antropomorficzną tradycję rozważania różnic między ludzkim a nieludzkim podmiotem. Podmiotowość jest rozumiana tutaj jako wynik procesów semiotycznych, która powstaje w wyniku działania opartego na znaczeniach powstałych w żywym organizmie. W poszukiwaniu podstaw inteligencji zwierząt, podmiotowość musi zostać mocno podkreślona jako sposób na samodoskonalenie rozumienia świata w miejsce biernego przyswajania informacji.

2.7 Intencjonalność jako funkcja biologiczna

Biosemiotyka, zajmuje się wyjaśnianiem pojawiania się relacji pomiędzy biologicznym ukierunkowaniem na cel i semiotycznym badaniem intencjonalności (*aboutness*)¹⁸⁰. Intencjonalność jest niezwykle ważnym pojęciem zarówno dla badań nad inteligencją i sztuczną inteligencją, ponieważ wskazuje na wyjątkowy status ontologiczny stanów mentalnych. Dla badających inteligencję niezwykle ważnym pytaniem jest, jak to możliwe, że stany mentalne zawsze są „o czymś”. Intencjonalność odnosi się do celowego działania, które ma na uwadze cel lub zamiar a także do stanów psychicznych (myśli, wierzeń, pragnień) skierowanych na jakiś przedmiot lub stan rzeczy¹⁸¹. Dlatego też

180 Intencjonalność może być z jednej strony postrzegana jako relacja lub jako bezpośrednio do czegoś lub też jako właściwość „bycia o czymś” istniejącym lub nieistniejącym; jakość lub fakt odnoszący się do czegoś lub o czymś; stan psychiczny, symbol, reprezentacja itp. *Aboutness*, [w:] *Oxford Dictionaries / English*.

181 Intencjonalność jest pojęciem określanym przez *Stanford Encyclopedia of Philosophy* jako moc umysłu, który przedstawia lub broni własności i stanów rzeczy. Termin pochodzi od średniowiecznej filozofii scholastycznej, lecz został przypomniany przez Franza Brentana i

podczas badania inteligencji zrozumienie intencjonalności jest istotne, ponieważ pomaga wyjaśnić, w jaki sposób żywe istoty rozumieją, o czym myślą.

Istnienie stanów mentalnych, skierowanych na konkretne obiekty, pozwala na refleksję nad różnymi aspektami tych przedmiotów, rozumowanie, opisywanie ich, a nawet dokonywanie wiarygodnych prognoz na ich temat. Zrozumienie czym są doświadczane obiekty: np. czym jest kamień, a czym żółw, pozwala opisać różnice między nimi a nawet przewidzieć ich zachowanie. Kant twierdził, że jednostka musi dysponować szeregiem pojęć, które może wykorzystać do racjonalnej oceny charakteru danego przedmiotu. Musi zrozumieć chociażby przyczynowość, by pojąć, że przedmiot popchnięty, poruszy się. Bez istnienia tych pojęć, nie byłoby prawdą stwierdzenie, że jednostka pojmuje przedmiot¹⁸². Jeśli intencjonalne stany są zgodne z ich obiektami, mamy pewne wyjaśnienie, w jaki sposób takie zrozumienie jest możliwe, ponieważ forma przedmiotu, do którego skierowany jest stan zamierzony, powinna być dostępna dla istoty myślącej. Rola, którą odgrywa intencjonalność w działaniu umysłu doprowadziła Franza Brentana do uznania intencjonalności za „znak umysłu”, który jest koniecznym i wystarczającym warunkiem stanów mentalnych¹⁸³. Wiele zjawisk jak dźwięki, obrazy, drogowskazy lub słowa zdają się ją przejawiać. Niemniej jednak intencjonalność tych zjawisk wydaje się pochodzić z intencjonalności umysłu, który je wytwarza. Dźwięk pozostaje dźwiękiem, obraz pozostaje przedmiotem, jeśli nie nadano mu znaczenia lub intencji.

Husserl podkreślał, że przedmioty myśli mają szczególny charakter. Po pierwsze, muszą być powiązane z innymi koncepcjami i pomysłami w umyśle myślącego w spójny sposób. Jeśli wyobrażenia o obiektach, które są napotymane w doświadczeniu, zbyt mocno kolidują z ze zrozumieniem tego, jak działa świat, to te idee rozpadną się¹⁸⁴. Husserl proponował, by badać naturę ograniczeń, które charakter umysłu stawia obiektom myślenia poprzez metodę, którą nazwał redukcją fenomenologiczną, polegającą na odkryciu warunków świadomości poprzez refleksję nad charakterem doświadczenia. Podejście to ma wiele wspólnego z transcendentnym idealizmem Kanta,

przejęty przez Edmunda Husserla. Często przywoływany przez filozofów umysłu jak John Searle, Jerry Fodor, Hilary Putnam Por.P. Jacob, *Intentionality*, [w:] *The Stanford Encyclopedia of Philosophy*, E. N. Zalta (red.), Metaphysics Research Lab, Stanford University 2014.

182 I. Kant, *Krytyka czystego rozumu*, R. Ingarden (tłum.), Warszawa 1957.

183 F. Brentano, *Psychology from an empirical standpoint*, Londyn 1874, ss. 88–89.

184 E. Husserl, *Idee czystej fenomenologii i fenomenologicznej filozofii*, D. Gierulanka (tłum.), Państwowe Wydawnictwo Naukowe 1975, s. 453.

ponieważ w obu przypadkach musimy rozpoznać, że natura umysłu może narzucać bardzo specyficzny charakter przedmiotom, jakie spotykamy w doświadczeniu. Zakładamy, że nasze doświadczenia narzucają nam fakty o świecie zewnętrznym. Idea, że natura umysłów nakłada ograniczenia na sposób, w jaki doświadczamy świata, jest w rzeczywistości twierdzeniem, które jest coraz powszechniej akceptowane, a fenomenologia stała się obszarem szczególnego zainteresowania wyłaniającej się dziedziny kognitywistyki¹⁸⁵, która jest mocno zaangażowana w badanie umysłu i inteligencji.

W dziedzinie biosemiotyki, badania wskazują na fakt, że działanie ciała i działanie umysłu nie są całkowicie oddzielnymi kategoriami, lecz są połączone poprzez intencjonalność zwierzęcia. Doświadczenia zmysłowe mają na celu monitorowanie zmian zachodzących w środowisku. Stanowią one reprezentacje, których głównym zadaniem jest udostępnianie podmiotowi tymczasowych informacji o środowisku. Wrażenia natomiast są reakcjami na zmiany środowiskowe. Ich wartość informacyjna jest ograniczona czasowo, lecz wystarcza organizmowi dla skutecznego działania w oparciu o te dane. Sensoryczne reprezentacje są wywoływane przez bodźce, natomiast myślenie i rozumowanie a także marzenia i wyobrażenia są formami przetwarzania informacji. Te dwa ostatnie występują również w przypadku braku bezpośredniego lub pośredniego związku przyczynowego z przedmiotami, z którymi podmiot wchodzi w interakcje. Zachodzą więc procesy reprezentacyjne, które nie wynikają bezpośrednio z tego, o czym są te procesy.

Biosemiotyka wprowadza intencjonalność do badań nad naturą, mimo iż to pojęcie pochodzi ze świata ludzkiej kultury i wydaje się, że nie ma zastosowania lub nie jest podatne na tłumaczenie na świat natury. Hoffmeyer sugeruje jednak, że żywe istoty przekształciły się w „systemy intencjonalne” ponieważ stworzyły możliwość połączenia danych zmysłowych z samoodniesieniem¹⁸⁶. Konsekwencją tego rozumowania jest, że intencjonalność u ludzi jest cechą biologiczną, ciągłością, a nie cechą emergentną, która pojawiła się wraz z powstaniem gatunku *Homo sapiens*. To, że taki proces ostatecznie wytworzył inteligencję, jest uderzającym faktem, który nie jest łatwy do wyjaśnienia pod nieobecność naturalnej teorii intencjonalności. Biosemiotyka, stawiając interpretację w centrum uwagi, podkreśla, że semioza jako nieodłączna cecha życia była podstawową formą intencjonalności i inteligencji. Widać tu ambicje poszerzenia teorii ewolucyjnej,

185 Zob. F.J. Varela, E. Thompson, E. Rosch, *The embodied mind*, MIT Press 1993.

186 J. Hoffmeyer, *Biosemiotics: an examination into the signs of life and the life of signs*, Scranton 2008, s. 31.

kładącą wyraźny nacisk na wpływ procesów umysłowych w najszerszym możliwym tego słowa znaczeniu, jako obejmującą interakcje semiotyczne, nawet na poziomie komórkowym. W tym celu należy wyjaśnić ewolucję, która tłumaczy doświadczenie intencjonalności i subiektywności wśród żywych stworzeń. Nie można zaprzeczyć, że te doświadczenia istnieją, ponieważ stanowią część życia każdego zwierzęcia. Postrzeganie tego faktu jako iluzji z punktu widzenia biologii prowadzi donikąd. Dlatego należy spróbować wytłumaczyć, jak doszło do powstania intencjonalności. Przede wszystkim jednak, jaka jest rola semiozy w tym nowym ewolucyjnym tłumaczeniu. W kontekście ewolucyjnym semioza ustala związek między organizmem, jego doświadczeniem świata i samym światem, ponieważ te relacje są dokładnie tym, czemu powinniśmy się przyglądać próbując zrozumieć naturalną ewolucję organizmów.

Możliwe jest wprowadzenie rozłączności między zachowaniem a aktywnością umysłową: np. zwierzęta mogą śnić, co oznacza, że stany mentalne mogą czasami być odłączone od działania cielesnego. W perspektywie biosemiotycznej nie jest to niczym niezwykłym, ponieważ mentalna ludzka intencjonalność wyrosła z cielesnej intencjonalności¹⁸⁷. Jest to parafrazą twierdzenia Uexküll'a, przedstawiającym każdą aktywność, która odciska znaczenie na przedmiocie i tym samym czyni go podmiotowym nośnikiem znaczeń¹⁸⁸. Aktywność mentalna jest tylko wyjątkowo wyrafinowanym rozszerzeniem tradycyjnych zachowań zwierzęcych. Wynika to z tego, że niekoniecznie musimy działać z dwiema zupełnie odmiennymi kategoriami, takimi jak elastyczność fenotypowa prostych organizmów oraz uczenie się zwierząt o bardziej rozwiniętym układzie nerwowym. W tej perspektywie, uczenie się jest szczególnie wyrafinowanym rodzajem elastyczności fenotypowej¹⁸⁹. Biosemiotyka może być postrzegana jako próba operacjonalizacji tej zaniedbanej części wglądu w ucieleśnioną naturę sfery mentalnej. Życie mentalne opiera się na intencjonalności cielesnej, przejawiającej się w percepcji i działaniu, a więc ostatecznie w semiozie.

Przypomnijmy, że semioza jest procesem triadycznym, w którym znak prowokuje do formowania się interpretatora, który stoi w relacji do przedmiotu i który w jakiś sposób

187 J. Hoffmeyer, *Signs of Meaning in the Universe...*, op. cit.

188 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten & Bedeutungslehre...*, op. cit., s. 110.

189 J. Hoffmeyer, K. Kull, *Baldwin and biosemiotics: What intelligence is for*, [w:] *Evolution and learning: The Baldwin effect reconsidered*, London 2003, ss. 253–272.

odzwierciedla relację, w której znajduje się sam znak¹⁹⁰. W klasycznym przykładzie dym może być postrzegany jako znak, który prowokuje formowanie się interpretatora (procesów neuronalnych) w umyśle ludzkiego obserwatora, który odnosi się do ognia w sposób odzwierciedlający to, jak sam dym odnosi się do ognia. Podobnie zwiększenie się poziomu adrenaliny może stać się sygnałem, który prowokuje tworzenie się interpretantu w komórce wątroby (kaskada procesów enzymatycznych ostatecznie uwalniających wolne molekuly glukozy), która odnosi się do rzeczywistej sytuacji stresowej w ten sam sposób, w jaki sama produkcja hormonów wiąże się z sytuacją stresową¹⁹¹. Obserwator ludzki posiada pewną wolność interpretacyjną, która wydaje się być czymś więcej odpowiedzi komórkowa wywołana adrenaliną, ponieważ wie, że dym, który obserwuje, jest nieprawdziwy, gdyż właśnie ogląda program informacyjny w telewizji, ale komórka wątrobowa musi reagować na adrenalinę. Nie istnieje tu żaden tajemniczy czynnik, który indukuje element wolności. Oczywiście można przypuszczać, że żywa komórka, która znajduje się w jakiejś sytuacji, wyzwala mechanizmy, które składają się na kontrolę komórkową, lecz rozpoznawanie komórek jest mediowane przez te same białka G, a różne białka G mogą czasami być wykorzystywane przez ten sam receptor. Umożliwia to komórce zmianę jej odpowiedzi na dany sygnał¹⁹². Zdolność receptorów, białek G i efektorów do interakcji z więcej niż jednym gatunkiem molekuly wewnątrz komórki oznacza, że komórka może dokonywać różnych wyborów.

Znak nie musi być związany z kontekstem komunikacyjnym. Większość procesów zachodzących w organizmie jest niezamierzona, bo znak nie został stworzony w celu interpretacji. Nie chcemy, by drżały nam ręce, gdy się denerwujemy, lecz nie sprawujemy nad tym pełnej kontroli, tymczasem obserwator właściwie interpretuje wynik. Pewne gatunki nabyły zdolność interpretowania pewnych wzorców zachowań innych gatunków, aby wyzwolić u nich nowe rodzaje zachowań. Motyl *Maculinea arion* składa jaja w tymianku. Larwy spadają na ziemię, gdzie wytwarzają mieszaninę lotnych chemikaliów, która naśladuje zapach larw czerwonych mrówek *Myrmica sabuleti*. Mrówki biorą te larwy za własne i przenoszą je do gniazda. Larwy motyla żerują na jajach i larwach mrówek aż

190 S. Brier, *The paradigm of Peircean biosemiotics*, „Signs-International Journal of Semiotics”, t. 2, 2008, ss. 30–81.

191 J. Hoffmeyer, K. Kull, *Baldwin and biosemiotics: What intelligence is for...*, op. cit.

192 *Ibid.*

do przepoczwarczenia. Motyle składają więc jaja w miejscach obfitujących w tymianek¹⁹³, co każe przypuszczać, że jakiś parametr związany z tymiankiem jest interpretowany jako znak składania jaj. W przyrodzie każda regularność i nawyk są znakiem dla jakiegoś innego organizmu lub gatunku¹⁹⁴. Ta zasada pokazuje ważny aspekt semiotyczny rozwoju intencjonalności, ponieważ ze względu na mechanizm oddziaływań semiotycznych wszystkie gatunki działają w sieci relacji semiotycznych. Na poziomie indywidualnym, jak i na poziomie ekosystemów, wszystkie wzorce interakcji są kontrolowane poprzez relacje semiotyczne.

Intencjonalność można zatem wyjaśnić naturalistycznie, jeśli jest postrzegana jako związana z semiozą. Należy akceptować fakt, że ludzka zdolność semiotyczna jest udoskonalaniem potencjału biosemiotycznego, który rozwija się od 4 miliardów lat¹⁹⁵. Potrzeby wszystkich żywych istot powinny być postrzegane jako stopień w rozwoju gatunków o zwiększonej wolności semiotycznej. W ewolucji organicznej nastąpiło wzmocnienie zdolności semiotycznych poprzez pojawianie się istot ludzkich, które posiadały szczególną zdolność do zrozumienia symbolicznego odniesienia. Różnica między człowiekiem a innymi zwierzętami jest ogromna, przede wszystkim dlatego, że ludzie znają różnicę między znakami i rzeczami, podczas gdy zwierzęta nie. Ludzka intencjonalność nie jest jednak unikalna na świecie, ale musi być rozumiana jako wysoce rozwinięta postać na polu semiotyki przyrody. Wraz z narodzinami istoty ludzkiej intencjonalność osiągnęła poziom progowy, na którym współdziała ze środowiskiem społecznym i kulturowym. Pozostaje jednak uzależniona od semiotycznego rusztowania.

2.7.1 Działania intencjonalne jako czynnik ewolucyjny

Filozofia interpretuje działania w kategoriach intencjonalności, które wyjaśnia przez pojęcie przyczynowości związane ze stanami mentalnymi i zdarzeniami, w których

193 S.F. Gilbert, D. Epel, *Ecological developmental biology: integrating epigenetics, medicine, and evolution*, Sunderland 2009, s. 86.

194 W celu znalezienia więcej przykładów tych nawyków, zamienianych w znaki patrz J. Hoffmeyer, *Signs of Meaning in the Universe...*, op. cit.

195 T. Deacon, *The symbolic species*, New York 1997; A. Sebeok Thomas, „Tell me, where is fancy bred?” *The biosemiotic self*, [w:] *Biosemiotics: The Semiotic Web 1991*, J. Umiker-Sebeok (red.), Berlin 1992.

bierze udział podmiot działający (agent). Głównym problemem w dyskursie na temat intencjonalności jest to, że uczestnicy często nie precyzują czy używają tego terminu do implikowania takich pojęć jak intencjonalność (*agency*)¹⁹⁶. Skoro agent jest istotą intencjonalną, to intencjonalność oznacza sprawowanie lub manifestowanie tej zdolności. Istnieją alternatywne koncepcje intencjonalności, które wskazują, że pojęcie to może być aplikowane do systemów, których działania mogą być wyjaśnione bez odwoływania się do przyczynowo-skutkowych stanów i zdarzeń mentalnych¹⁹⁷. Dyskusje na temat natury intencjonalności trwają również w dziedzinie biosemiotyki. Intencjonalność jest tam traktowana jako główny atrybut życia i jako podstawa przypisania systemu do żywych organizmów¹⁹⁸. Może być określana jako kontrolowane selektywne działanie, związane z samorefleksyjnością i ukierunkowywaniem na przyszłość, a także jako istota samopodtrzymania systemu będąca celem każdej aktywności semiotycznej¹⁹⁹. Należy wspomnieć, że mimo iż Uexküll nie użył tego terminu przypisuje się mu go retrospektywnie²⁰⁰. Przypisanie to jest jednak zasadne, ponieważ opanowanie przez

196 Mimo, że termin *agencja* istnieje w języku polskim, nie zyskał tego samego pola semantycznego, co *agency* w języku angielskim. Powiązany termin Uexkülla to *Wirkwelt*, który jest interpretowany jako świat operacyjny i komplementarna część świata percepcyjnego. Jako całość te dwa światy tworzą *Umwelt* podmiotu. Por. T. von Uexküll, *Glossary, „Semiotica”*, t. 42, nr 1, 2009, ss. 83–87. Olbrzymia większość publikacji dotyczących tego zagadnienia w biosemiotyce powstaje w języku angielskim, zatem termin *agency* jest powszechnie używany. Dodatkowym uzasadnieniem na podtrzymanie tego zabiegu, jest użycie pojęcia *agent*, które występuje powszechnie zarówno w naukach przyrodniczych, inżynieryjno-technicznych jak i społecznych (tu też wymiennie jako *aktor*). Mimo tego tłumaczenie terminu *agency* jako *agencja* w języku polskim nie wydaje się być korzystnym zabiegiem, przez wzgląd na fakt, iż posiada już ono bardzo mocno sprecyzowane znaczenia, nie wiążące się w żaden sposób z pojęciem *agency*. W naukach społecznych funkcjonuje tłumaczenie pojęcia *agency* jako *sprawczość*. Termin ten nie wyczerpuje jednak znaczenia tego słowa, ponieważ określa zdolność do oddziaływania na inne jednostki bądź obiekty, co może być odnoszone nawet do przedmiotów. Thomas Eriksen używa w zamian terminu *intencjonalność*, by odróżnić sprawczość celową podmiotu od sprawczości biernej, instynktownej i zaprogramowanej. Pozwalam więc sobie pozostać przy pojęciu używanym przez antropologię społeczną.

197 O filozoficznym pojęciu *agency*: M. Schlosser, *Agency*, [w:] *The Stanford Encyclopedia of Philosophy*, E. N. Zalta (red.), Metaphysics Research Lab, Stanford University 2015.

198 K. Kull i in., *Theses on biosemiotics: Prolegomena to a theoretical biology*, [w:] *Towards a Semiotic Biology: Life is the Action of Signs*, World Scientific 2011, ss. 25–41.

199 A.A. Sharov, *Pragmatics and biosemiotics.*, „Sign Systems Studies”, t. 30, nr 1, 2002.

200 C. Emmeche, *The emergence of signs of living feeling: Reverberations from the first Gatherings in Biosemiotics*, „Σημειωτική - Sign Systems Studies”, t. 29, nr 1, 2001, ss. 369–376.

zwierzę jego środowiska jest rozumiane w pismach Uexküll'a jako czynna interpretacja, która zakłada pewien stopień autonomii i niemechanistycznej przyczynowości. Intencjonalność tam nie jest kwestią spontaniczności przyjętej w granicach związków przyczynowo-skutkowych, ponieważ dla Uexküll'a te mogą być osiągnięte tylko na powierzchni organizmu, gdzie zewnętrzne bodźce są przekształcane w impulsy nerwowe. Mechanistyczna przyczynowość kończy się w chwili, w której bodźce są podejmowane selektywnie²⁰¹. Nie jest też urzeczywistnieniem związków przyczynowo-skutkowych, zgodnie z którymi przetwarzanie informacji jest uporządkowane.

Intencjonalność jest zdolnością systemu jednostek do generowania zachowań ukierunkowanych na cel i nie można jej wytłumaczyć doborem naturalnym. Wynika to z faktu, że dobór naturalny nie może, jak twierdził Karol Darwin, działać bez wysiłku i bez walki by zapewnić przetrwanie gatunkowi²⁰². Wszystkie organizmy wykazują się intencjonalnością: nieustannie dążą do znalezienia pożywienia, pary, schronienia, ucieczki przed drapieżnikami itd. Większość biologów wydaje się opisywać intencjonalny aspekt życia teorią konkurencyjnej strategii genów²⁰³. Hoffmeyer twierdzi jednak, że geny jako elementy cząsteczek nie mogą posiadać strategii i uważa, że neo-darwinizm wyjaśnia ich funkcje w oparciu o przemycanie homunkulusa, który i tak pełni rolę intencjonalności ukrytej w głębokiej strukturze teorii²⁰⁴. W wielu przypadkach można interpretować agencję jako prekursor doboru naturalnego, który może przyjmować szlaki, które zostały dla niego częściowo przygotowane, a nie kierowane przypadkami korzystnych mutacji²⁰⁵. Rozważania te stanowią poważne wyzwanie dla koncepcji ewolucji inteligentnych zachowań.

201 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten & Bedeutungslehre...*, op. cit., s. 127.

202 C. Darwin, *The Origin of Species*, Alachua 2009, s. 97,108,115.

203 J. Hoffmeyer, *Surfaces inside surfaces. On the origin of agency and life*, „Cybernetics & Human Knowing”, t. 5, nr 1, 1998, ss. 33–42.

204 *Ibid.* Koncepcja homunkulusa (małego człowieka) była pierwotnie używana w odniesieniu do badań alchemików, którzy dosłownie traktowali koncepcję dualizmu. Termin upowszechnił się, gdy przez pierwsze mikroskopy odkryto „żywe zwierzęta” w spermie ludzi i innych zwierząt. Wówczas pojawił się pomysł, że mały człowiek, homunkulus, leży uformowany w głowie plemnika, co tłumaczy ludzką ontogenezę: po połączeniu z komórką jajową homunkulus rozwija się do dorosłości i manipuluje ciałem człowieka.

205 T. Maran, K. Kleisner, *Towards an evolutionary biosemiotics: semiotic selection and semiotic co-option*, „Biosemiotics”, t. 3, nr 2, 2010, ss. 189–200.

Biosemiotyka utrzymuje, że utrwalanie genetyczne jako mechanizm stabilizacji nowych cech lub zachowań, może być opisane za pomocą kręgów funkcjonalnych. Przedstawienie rozwoju intencjonalności dostarcza koncepcja cyklu funkcjonalnego, a zwłaszcza idea zamiany sygnałów percepcyjnych na działania efektorów. Intencjonalność spełnia rolę medium przypisującego znaczenie dla informacji wejściowej i kodowanie jej w kategoriach sygnałów, w wyniku czego powstaje znak percepcyjny²⁰⁶. Na przykład każdy lotny związek, feromon, może być wychwytywany i wykorzystywany jako nośnik w procesie znaczenia. Dzięki niemu pewne zachowanie jest uwalniane w organizmie, bez względu na to, jaka jest dokładna budowa chemiczna tego feromonu²⁰⁷. Ten brak bezpośredniego zaangażowania znaku w biochemię i fizjologię usuwa ograniczenia mechanizmów przyczynowych, które w innym przypadku musiałyby być przestrzegane jako takie. W późniejszych etapach ewolucji takie działanie semiotyczne prawdopodobnie miało istotny wpływ na rzeczywisty wynik selektywnych rozwiązań. Zamiast opisywać zachowania organizmu jako odziedziczone instynkty, biosemiotyka sugeruje nowe subtelne rodzaje działań intencjonalnych. Biosemiotyka akceptuje istnienie przyczynowości, ponieważ proces powstawania znaczenia opiera się na relacjach między znakami. Uznaje jednak, że przyczynowość ta jest kierowana, nie wywoływana. Rozwój organizmu może być analizowany jako sfera relacji łączących sygnały środowiskowe i sensomotoryczne moduły z wzorcami kontekstowymi zależnymi od pamięci i intencji.

Wprowadzenie tej koncepcji ma znaczący wpływ na zrozumienie powstania inteligentnych zachowań. Nowe rozumienie dynamiki złożonych systemów dało nam obraz świata, w którym oddolne procesy wchodzą w interakcje z procesami odgórnymi w skomplikowany sposób²⁰⁸. Oznacza to, że nie wszystko można przewidzieć i że złożone systemy, takie jak systemy żywe, mogą ewoluować, tworząc struktury wyższego rzędu, które mogą mieć wpływ przyczynowy na działania na niższym poziomie. Pojawienie się systemów życiowych implikuje pojawienie się intencjonalności jako integralnej części związanej z dynamiką takich systemów. Jeśli chcemy wyjaśnić inteligencję, musimy włączyć koncepcje intencjonalności do badań. W przypadku braku tego narzędzia

206 S. Rafieian, *A biosemiotic approach to the problem of structure and agency*, „Bisemiotics”, t. 5(1), 2012, ss. 83–93.

207 A.A. Sharov, *Biosemiotics: A functional-evolutionary approach to the analysis of the sense of information*, [w:] *Biosemiotics: The semiotic web*, 1991, ss. 345–373.

208 R.B. Laughlin, *A Different Universe: Reinventing Physics from the Bottom Down*, New York 2005.

wyjaśniającego, nauki poznawcze nie pozbędą się ukrytych homunkulusów, które pojawiają się pod mechaniczną strukturą organizmu. Biosemiotyka jest niezbędna, aby jasno przedstawić te różnorodne założenia importowane do biologii za pomocą takich niezanalizowanych koncepcji teleologicznych, jak funkcja, adaptacja, informacja, kod, sygnał, znak itd. oraz dostarczyć teoretycznego uzasadnienia dla tych pojęć.

Koncepcja intencjonalności jako kierującej ewolucją, uwzględnia kreowanie przyszłości, w miejsce biernego oczekiwania na selektywne zmiany. Intencjonalność musi być częścią procesu przewidywania i uczenia się, która radzi sobie ze zmianami zachodzącymi w środowisku. W biologii proces uczenia się można mierzyć zmianami w częstotliwości pojawiania się genów w kolejnych pokoleniach. W tym uproszczonym modelu każde pokolenie stanowi zestaw genotypów wyrażających się jako fenotypy, które mają znaczenie w aktach przetrwania i reprodukcji. Powstały zestaw genotypów w następnym pokoleniu odzwierciedla zatem niepowodzenia i sukcesy poprzedniego zestawu. Agenci są zdefiniowani jako systemy, które są w stanie kontrolować swoje zachowanie w celu zwiększenia własnej wartości. Swoboda działania w ramach intencjonalności opiera się na rozróżnieniu znaków, które opisują stan lub etap, lub są interpretowane jako pamięć. Działanie, percepcja i relacje powstają w wyniku działania intencjonalności w trakcie ewolucji i uczenia się, by wspierać użyteczne i elastyczne zachowania. Inteligentne zachowanie można wyjaśnić, przewidzieć i zmodyfikować za pomocą zasady optymalności, zgodnie z którą organizmy wybierają te działania, które mają zwiększyć ich wartość.

2.7.2 Koncepcja nowej przyczynowości biologicznej

Równie ważnym elementem dyskursu biosemiotycznego jest ugruntowanie nowej koncepcji przyczynowości związanej z możliwością nadawania działaniu ostatecznych priorytetów – już od rozwoju genomu, poprzez tendencje rozwojowe i ewolucyjne²⁰⁹. Wdrożenie koncepcji podmiotowości jest próbą wprowadzenia i uzasadnienia innej niż mechaniczna przyczynowości, która opiera się na sformułowaniu wyraźnej biosemiotycznej koncepcji znaczenia. Przypisanie podmiotowości żywym systemom i

209 F. Stjernfelt, *Tractatus Hoffmeyerensis: Biosemiotics as expressed in 22 basic hypotheses*, „Sign Systems Studies”, t. 30, nr 1, 2002, ss. 337–345.

afirmacja własnej przyczynowości stanowi sposób, w jaki biosemiotyka może przedstawić teoretyczne ramy dla zrozumienia prostych inteligentnych systemów, odmiennie od wyobrażenia, że są po prostu zorganizowanymi komórkami i narządami²¹⁰. Organizmy, komórki i tkanki, z których są zbudowane, są nie tylko obiektami, ale także podmiotami w tym sensie, że są to czynniki semiotyczne zdolne do interakcji z otoczeniem w inteligentny sposób. Inteligencja oparta na percepcji przyczyn i skutków pozwala postrzegać świat poprzez ciągi zdarzeń, logiczne myślenie i konstruktywne rozwiązywanie problemów. Część biosemiotyków, jak Hoffmeyer poszukuje wyjaśnień podmiotowości w badaniu dynamiki semiozy żywych systemów²¹¹ lub postrzega ją jako zależną od zdolności zmysłowych organizmu²¹². Inni definiują to pojęcie czerpiąc z koncepcji relacji znakowych Peirce'a²¹³. W badaniach pojawiają się również próby zastosowania heurystycznego potencjału semiotycznej teorii kultury Łotmana w odniesieniu do biosemiotyki²¹⁴. Wśród tekstów poświęconych temu zagadnieniu znaleźć można argumenty za istnieniem ostatecznej przyczyny²¹⁵, przekonanie, że obiektywna przyczynowość jest najbardziej odpowiednia dla wyjaśnienia działania znaków²¹⁶ lub też opisu działania złożonych systemów życia jako opierających się na formalnej i ostatecznej przyczynowości²¹⁷. Formalny ma oznaczać przyczynowość w dół (*top-down causation*) od struktury organizmu do najmniejszych jego elementów. Ta przyczynowość polega na ograniczeniu ich działania i nadaniu im funkcjonalnego znaczenia w odniesieniu do całego metabolizmu. Z drugiej strony przyczynowość wskazuje na tendencję do nabywania przyzwyczajęń i tworzenia interpretacji na temat przyszłych wydarzeń.

Hoffmeyer przyjmuje termin „ostateczna przyczyna”, którą rozumie (podobnie jak Peirce) nie jako celową, świadomie pojmowaną przyczynę końcową, ważną jedynie w

210 F. Stjernfelt, C. Emmeche, K. Kull, *Reading Hoffmeyer, rethinking biology...*, *op. cit.*, s. 7.

211 J. Hoffmeyer, *Semiotic Scaffolding Of Living Systems*, [w:] *Introduction to Biosemiotics*, M. Barbieri (red.), Springer Netherlands 2008, ss. 149–166.

212 F. Stjernfelt, *Tractatus Hoffmeyerensis...*, *op. cit.*

213 S. Brier, *The paradigm of Peircean biosemiotics...*, *op. cit.*; T.A. Sebeok, *The semiotic self revisited...*, *op. cit.*

214 K. Kull, *Towards biosemiotics with Yuri Lotman*, „Semiotica”, t. 127, nr 1–4, 1999, ss. 115–132.

215 J. Hoffmeyer, *Semiotic Scaffolding Of Living Systems...*, *op. cit.*

216 J. Deely, *Four ages of understanding*, University of Toronto Press 2001.

217 M. Barbieri, *Biosemiotics: a new understanding of life*, „Naturwissenschaften”, t. 95, nr 7, 2008, ss. 577–599.

kontekście ludzkiej woli, ale jako ogólną zasadę przyczynowości. Jest to odrzucenie ostatecznej przyczyny w tradycyjnym paradygmacie wyjaśnienia życia.

[...] należy wnioskować, że życie i ostateczna przyczyna są - przynajmniej potencjalnie - nieodłącznie związane z podstawową fizyką naszego wszechświata, a nie zaprzeczać ostatecznej przyczynowości od razu, powinniśmy dokonać rozróżnienia między akceptacją i nieakceptacją ostatecznej przyczynowości²¹⁸.

Istnieje potrzeba uzasadnienia nowej przyczynowości niezbędnej do wyjaśnienia poznania i stanów mentalnych. Procesy znakowe są rodzajem obiektywnej, a nie ostatecznej przyczynowości. Dzięki nim związek między zjawiskami jest dostrzegany, a wydarzenia zostają połączone. Obiektywna przyczynowość wydaje się odpowiednia, ponieważ określa zarówno aktywność życiową, jak i przypadkowe interakcje na poziomie natury nieorganicznej. Ta przyczynowość jest odpowiednia do ugruntowania zachowania znaków, które prowadzi ostatecznie do koncepcji jaźni. Relacyjny charakter znaku rodzi możliwość nowej przyczynowości. To także otwiera nową możliwość wyjaśnienia fenomenalnego życia: wyjaśnienie w kategoriach modeli mechanistycznych lub informacyjnych pozostawia nam całkowicie niemożliwy do rozwiązania problem przyczynowości, model semiotyczny wskazuje na to, że kładziemy nacisk na zjawiska relacyjne, które w zasadzie są niezależne od wagi powiązanych ze sobą elementów²¹⁹.

Interpretant jest częścią procesu relacyjnego, w którym postrzegany nośnik znaku staje się powiązany z obiektem, w taki sposób, że naśladuje jego własny stosunek do tego samego przedmiotu. Interpretant jest tworzony jako reakcja w odpowiedzi na wydarzenie, ale jest także wynikiem historii organizmu, co sugeruje, że te poprzednie doświadczenia wpływają na proces interpretacji od najwcześniejszych etapów. Stąd bierze się niezależność interpretatorów od materialności znaku oraz ich nieograniczona różnorodność ze względu na kontekst i potrzeby organizmu. Stosunki przestrzenne między przedmiotami w środowisku postrzegane są w uniwersalnej przestrzeni, która jest niezależna od doznań sensorycznych. Autopojetyczny charakter interpretacji, w trakcie którego jeden interpretator zostaje wygaszony przez inny pokazuje, że proces znakowy nie wywołuje odpowiedzi w tradycyjnym sensie skutecznej przyczynowości, ale przewiduje dopuszczenie przyczynowości semiotycznej lub poddanie przedmiotu pod interpretację w

218 J. Hoffmeyer, *Semiotic Scaffolding Of Living Systems...*, *op. cit.*, s. 98. tłum. własne.

219 *Ibid.*, s. 103.

kontekście lokalnym²²⁰. Pozwala to na wyodrębnienie się podmiotu ze środowiska i pojawienie się stanów mentalnych uwarunkowanych zdarzeniami zewnętrznymi.

Potwierdzenie ostatecznej przyczyny, wskazuje na potrzebę wyeliminowania zasady mechaniki sensorycznej, która jest traktowana jako wyjaśnienie, w jaki sposób zewnętrzny świat wpływa na umysł. Zamiast tego mechanikę sensoryczną należy zastąpić semiotyką sensoryczną. Zamiast „ciała mechanicznego” nauki poznawcze powinny przyjąć „ciało semiotyczne” jako punkt wyjścia²²¹. Dzięki temu można uznać wrażenia za interakcję między pamięcią, impulsami czuciowymi i aktywnością motoryczną. Co więcej, przyjęcie takiej przyczynowości prowadzi do postawienia jednego z twierdzeń biosemiotyki, zgodnie z którym intencjonalność fenomenalnej świadomości jest wyrafinowanym przejawem intencjonalności cielesnej, której historia łączy się z naturalną historią znaczenia, pozwalającej na osiągnięcie pożądanego stanów rzeczy.

2.8 Biosemiotyczna koncepcja inteligencji

Inteligencja jest jedną z wyrafinowanych działalności, która jest tradycyjnie postrzegana jako rozszerzenie biologicznego zachowania zwierząt. Współczesne badania uznają jednak, że zwierzęta posiadają inteligencję²²². Wielu biologów długo uważało, że instynkty i inteligentne zachowanie stoją do siebie w opozycji, a sukces ich współdziałania wyłania się jako przypadkowy zbieg okoliczności²²³. Warto jednak podkreślić, że inteligencja nie jest nieuchwytną cechą, która znajduje się gdzieś w umyśle, lecz jest związana zdolnościami społecznymi, umiejętnością używania fizycznych znaków, a także możliwością gromadzenia i organizowania wiedzy. Peirce określił terminem semiotyczny „to, co musi być charakterystyczne dla wszystkich znaków używanych przez [...]”

220 J. Hoffmeyer, *From Thing to Relation. On Bateson's Bioanthropology*, [w:] *A Legacy for Living Systems*, t. 2, 2008, ss. 27–44.

221 J. Hoffmeyer, *Semiotic Scaffolding Of Living Systems...*, *op. cit.*, s. 95.

222 F. De Waal, *Are we smart enough to know how smart animals are?*, WW Norton & Company 2016; B. Heinrich, H. Kober, *Mind of the raven: investigations and adventures with wolf-birds*, Cliff Street Books New York 1999; R. Menzel, J. Fischer, *Animal thinking: contemporary issues in comparative cognition*, MIT press 2011, t. 8.

223 Patrz. C.P. Richter, *Animal Behavior and Internal Drives*, „The Quarterly Review of Biology”, t. 2, nr 3, 1927, ss. 307–343; E. Thorndike, *Animal intelligence: Experimental studies*, Routledge 2017, s. 150.

inteligencję zdolną do uczenia się przez doświadczenie”²²⁴. Rozszerzenie inteligencji na wszystkie żywe systemy jest próbą zniwelowania ontologicznej przepaści między człowiekiem a zwierzęciem. Pozwala również przyjrzeć się inteligencji maszynowej bez potrzeby odwoływania się do ludzkiego systemu kognitywnego.

Semioza odgrywa kluczową rolę w doświadczeniu i zdobywaniu wiedzy o świecie. Proces semiotyczny przenika początek, środek i koniec rozwoju każdego organizmu. W systemie semiotycznym formują się plany i wszystkie możliwe koncepcje, bez których nie zaszłoby inteligentne działania. Dlatego symbolizacja jest podstawą inteligencji. Biologia i semiotyka coraz bardziej się przenikają. W rzeczywistości biologia dostarcza silnej podstawy empirycznej dla zrozumienia, że zdolności semiotyczne są istotą inteligencji. Każdy kod biologiczny lub genetyczny jest *par excellence* systemem semiotycznym i może być interpretowany jako podstawa samej inteligencji: powiązania kodów genetycznych zawartych w żywych komórkach z samym charakterem tych komórek, różnicowanie się i integracja tych komórek w złożone organizmy, ich rozwój w różnych narządach, a ostatecznie ich organizacja w mózgu. Inteligencja wskazuje na proces doświadczenia świata fizycznego w taki sposób, aby miał sens zarówno dla podstawowych struktur jak i dla całości życia organizmu.

Zwierzęta wykazują się inteligentnym zachowaniem, takim jak planowanie przyszłości, polowanie na określony rodzaj zdobyczy, pamięć o znalezieniu lub zbudowaniu schronienia, branie udziału w zalotach a także używanie innych strategii powszechnie uważanych za występujące tylko u wyższych gatunków. Dla Uexküll'a jest to związane z naturalnym cyklem życia, opartego na zasadach regulujących czynności życiowe. W subiektywnym *Umwelcie* zwierzę produkuje warianty działań, dostrzega gradacje, przewiduje niepowodzenia i daje się zwieść iluzjom. Postrzega świat w sposób zintegrowany, nie dzieląc go na części, kolory lub dźwięki. Jego percepcje i działania są inteligentne, odnoszą się do kontekstu i posiadają głęboki sens. Doświadczenie konkretnej sytuacji determinują zmieniające się nastroje, obecne potrzeby i intencje, a także zdolności zmysłowo-motoryczne. Umiejętności te rozwijają się w trakcie doświadczenia świata,

224 C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 2: Elements of Logic*, Charlottesville, Va 1994, akap. 227.

poprzez praktykę i dlatego wydają się odgrywać zasadniczą rolę w ucieleśnieniu i usytuowaniu, tkwiących u podstaw wszystkich inteligentnych zachowań²²⁵.

Wcześniejsze próby poradzenia sobie z tą wiedzą przyniosły teorię zakładającą dualizm między umysłem a ciałem. W tym miejscu biosemiotyczne podejście może pomóc, ponieważ postrzega znaczenia (*sema*) jest właściwe dla ciała (*soma*). Ciało jest z natury zaangażowane w procesy komunikacyjne, które służą do koordynowania aktywności komórek, tkanek i narządów w organizmie, a także wymianę komunikatów integrujących różne poziomy hierarchii. Widziana w tym świetle inteligencja jest po prostu interfejsem, za pomocą którego organizm zarządza relacjami z otaczającym go środowiskiem i jego rzeczami. Żywe organizmy są systemami, których nie można interpretować jako niezależnych od środowiska, lecz jako systemy adaptacyjne.

Biosemiotyczna teoria intencjonalności jest ucieleśnieniem cech zorganizowanych, żywych, dynamicznych systemów żywych. Inteligentni agenci to podmioty ukierunkowane na cele odpowiednie do swojej sytuacji i wchodzące w interakcję ze środowiskiem w sposób, który pozwala te cele osiągać. Wszystkie żywe organizmy posiadają wspólne elementy inteligencji, np.: wrażliwe na kontekst zachowania adaptacyjne, umiejętność przetwarzania informacji, tworzenie społeczności. Inteligencja w perspektywie autonomiczności jest rozumiana jako zwiększenie samokontroli. Zwierzęta wartościują przepływające informacje i przewidują interakcje tak aby osiągnąć cele. Związany z tym proces uczenia się wynika z dążenia do poprawy antycypacji, w którym rozpoznawany jest kontekst oraz uwzględniona jest niejasność lub niepełność informacji²²⁶. Rozwój zachowań inteligentnych odbywa się poprzez poprawę błędów, rozpoznawanie kontekstu i budowę ulepszanego przewidywania.

Biosemiotyka zwraca również uwagę na funkcjonowanie rozproszonej inteligencji, która opiera się na zewnętrznych rusztowaniach semiotycznych w tej samej mierze, co na własnych zdolnościach i wiedzy podmiotu²²⁷. Działanie i wiedza nie pochodzą tylko z jednego wewnętrznego źródła. Rusztowanie zapewnia uzyskanie odpowiedniego poziomu

225 H.L. Dreyfus, *What computers can't do: The limits of artificial intelligence*, Harper & Row New York 1979; J.R. Searle, *The rediscovery of the mind...*, *op. cit.*

226 J. Deely, *Pars Pro Toto from Culture to Nature: An Overview of Semiotics as a Postmodern Development, with an Anticipation of Developments to Come*, „The American Journal of Semiotics”, t. 25, nr 1, 2009, ss. 167–192.

227 E. Thompson, *Empathy and Consciousness*, „Journal of Consciousness Studies”, t. 8, nr 5–7, 2001, ss. 1–32.

poznawczego, wspomagając poprawę funkcji poznawczych działającego podmiotu w celu prostszego wykonania zadania. Rozproszone poznawanie opisuje akty kognitywne jako rezultat działania systemu obejmującego zarówno podmiot działający, innych i przedmioty otoczeniu, oraz kontekst sytuacyjny. Rozproszona inteligencja uwzględnia wsparcie, umożliwiające wykonanie działania, które w innym razie byłyby podatne na błędy, trudne lub niemożliwe do osiągnięcia tylko w oparciu o zdolności agenta. Analiza działań opartych na zewnętrznych rusztowaniach może umożliwić zrozumienie inteligentnych zachowań, których podstawa znajduje się poza umysłem.

Biosemiotyka rozumie też współzależność między odniesieniem do siebie a odniesieniem do tego, co na zewnątrz, które wskazuje na przyszłe ukierunkowanie systemu żywego. Na najniższym poziomie, odnosi się do nabywania przez komórki wewnętrznych markerów strukturalnych, które odzwierciedlają przeszłe wydarzenia i wpływają na przyszłe działania. Autonomia znajduje swoje podłoże w umiejętności zachowania samokontroli, tożsamości, jedności systemu oraz kształtowania celów, zapewniających przetrwanie. Organizm podejmuje działania w celu zapewnienia dostępu do zasobów niezbędnych do dalszego życia. Autonomiczna aktywność agenta jest mediowana semiotycznie, przez znaki: jedzenia, niszy, tego, gdzie należy się znajdować, co jeść i jak uruchamiać właściwe wewnętrzne procesy produkcji składników organizmu we właściwym czasie²²⁸.

Życie mentalne opiera się na intencjonalności cielesnej, przejawiającej się w cyklach postrzegania i działania. Ta aktywność jest głęboko zintegrowana ze strategiami przetrwania organizmów. Przede wszystkim jest jednak działaniem zamierzonym, czyli prekursorem dla tego wymiaru życia, który u wyższych zwierząt jest postrzegany jako inteligencja, a ostatecznie świadomość. Biosemiotyka bada podstawowe procesy, z których powstało życie mentalne i twierdzi, że pojedyncze procesy komórkowe i działanie całości organizmu to procesy semiotyczne, które poprzedzają pojawienie się prawdziwej inteligencji i życia umysłowego. Semioza powinna być postrzegana jako pierwotna forma inteligencji. Zachowania inteligentne stają się natychmiast zrozumiałe w tym podejściu, ponieważ semioza łączy przyczynowość ze swoją triadyczną naturą, gdy zachodzi w kontekstowej sytuacji. Inteligencja jest procesem powstałym w danym kontekście, a jednym z głównych tematów ewolucji organicznej jest wnikanie kontekstu do organizmu,

228 C. Emmeche, *Organism and body: The semiotics of emergent levels of life*, [w:] *Towards a Semiotic Biology: Life is the action of signs*, World Scientific 2011, ss. 91–111.

co znaczy, że w trakcie ewolucji organizmy nauczyły się tworzyć coraz bardziej wyrafinowane reprezentacje wewnętrzne aspektów sytuacji zewnętrznych ²²⁹ . Taka aktywność jest źródłem biosemiozy, a ponieważ jest głęboko zintegrowana ze strategiami przetrwania organizmów, jest działaniem zamierzonym, czyli prekursorem tego wymiaru procesu życiowego, który uznajemy za inteligencję.

229 A. Sharov, T. Maran, M. Tønnessen, *Comprehending the Semiosis of Evolution*, „Biosemiotics”, t. 9, nr 1, 2016, ss. 1–6.

3 Modelowanie sztucznej inteligencji w biosemiotycznych ramach koncepcyjnych

Prace Uexküll'a, które okazały się przełomowe dla biologii, przyniosły wyjaśnienie, w jaki sposób w interakcji system-środowisko powstaje znaczenie. Jak wykazano w poprzednim rozdziale, ucieleśnienie, antycypacja i adaptacja środowiskowa odgrywają istotną rolę w procesie kształtowania się inteligencji. Opisane przez biosemiotykę procesy przetwarzania znaków przez żywe organizmy pokazują, w jaki sposób powstaje znaczenie wpływające na pojawianie się inteligentnych zachowań. Tradycyjne podejście do sztucznej inteligencji zostało mocno skrytykowane, głównie z powodu ograniczeń semantycznych sztucznych systemów. Budowa i sposoby rozwoju żywych systemów stały się inspiracją dla konstruowania ich sztucznych odpowiedników. Ważne jest zatem przyjrzenie się sygnałom odbieranych ze środowiska przez sztuczny system, które powinny wpływać na kształtowanie się inteligentnych działań. Biosemiotyczna perspektywa daje dużą szansę na opracowanie ram koncepcyjnych, które mogą zostać użyte w kreowaniu sztucznej inteligencji.

Dotychczasowe badania SI inspirowane biosemiotyką podejmowały kwestie związane ze znakami i powstawaniem znaczenia w sztucznych systemach. Prace obejmowały m.in. tworzenie architektury subsumpcyjnej²³⁰, komunikację synergiczną²³¹, użycie znaków i języka²³², mapowanie obiektów środowiska w strukturach wewnętrznych agenta (tzw. ugruntowanie symboli)²³³. Mimo tego, ostatnie badania w tej dziedzinie rzadko odnoszą się do fundamentalnych zagadnień semantycznych, interakcji system-środowisko lub innych problemów podejmowanych przez biologiczną teorię znaczenia. Większość badań w ogóle lub tylko w niewielkim stopniu nawiązuje do specyficznych

230 R.A. Brooks, *A robust layered control system for a mobile robot...*, *op. cit.*

231 J. Fredslund, M.J. Mataric, *A general algorithm for robot formations using local sensing and minimal communication*, „IEEE Transactions on Robotics and Automation”, t. 18, nr 5, 2002, ss. 837–846.

232 A. Billard, K. Dautenhahn, *Grounding communication in situated, social robots*, [w:] 1997, ; L. Steels, P. Vogt, *Grounding adaptive language games in robotic agents*, [w:] *Proceedings of the fourth european conference on artificial life*, t. 97, MIT Press 1997, ss. 474–482.

233 S. Coradeschi, A. Saffiotti, *An introduction to the anchoring problem*, „Robotics and Autonomous Systems”, t. 43, nr 2–3, 2003, ss. 85–96.

cech charakterystycznych dla życia, które są kluczowe dla rozwijania sztucznej inteligencji.

3.1 Możliwości powstania inspirowanych biosemiotycznie inteligentnych systemów

Od lat osiemdziesiątych XX wieku wiele czasu poświęcono badaniu ucieleśnionych autonomicznych agentów. Takie systemy zwykle uczą się, rozwijają i ewoluują w interakcji z ich otoczeniem. Można argumentować, że te samoorganizujące się właściwości rozwiązują problem ugruntowania symbolu lub reprezentacji w badaniach sztucznej inteligencji, a tym samym autonomiczni agenci mogą być rozważani w perspektywie semiotycznej. Należy więc przyjrzeć się możliwościom implikacji teorii znaczenia Jakoba von Uexküll'a w badaniu potencjału powstania sztucznych organizmów. Zastosowanie teorii niemieckiego biologa powinno umożliwić badanie sztucznych systemów w perspektywie biologicznych podstaw poznania poprzez sprawdzenie w jakim stopniu sztuczne organizmy są ucieleśnione, zdolne do semiozy i autonomii, oraz czy mogą stać się podmiotem poznawczym.

3.1.1 Status semiotyczny sztucznych systemów

Perspektywa biosemiotyczna zakłada, że inteligencja powstaje w trakcie subiektywnego procesu, który ma miejsce w pewnej określonej sferze. Za dynamikę rozwijania się zachowań inteligentnych odpowiada proces semiotyczny, który wiąże jednostkę poznawczą (agenta) ze szczególnym kontekstem fizycznym (środowiskiem) w wyjątkowym związku dialektycznym. Płaszczyznę działań organizmu i użycie kręgów funkcjonalnych można prześledzić skupiając się na rodzaju interakcji, które są możliwe i ograniczone przez fizyczną architekturę organizmu. W *Umwelcie* organizmu szczególne cechy środowiskowe zostają wyróżnione i zindywidualizowane. Przez agenta są postrzegane jako nacechowane wartością znaczenia, dlatego wybiera właśnie je spośród innych, zależnie od swoich potrzeb lub celów.

Złożony charakter systemu zmysłowego podkreślał nie tylko Uexküll, ale i chociażby Ernst Cassirer, mówiąc, że krąg działania przypomina punkt widzenia²³⁴. Żywe organizmy są systemami, które są podporządkowane swojej konstrukcji i oprogramowaniu. Każdy organizm posiada informację genetyczną, która kieruje systemem przez różne cykle rozwojowe i różne konteksty środowiskowe. Informacja ta determinuje budowę organizmu, która umożliwia użycie odziedziczonej pamięci i przedempirycznej wiedzy, która pozwala organizmowi rozpoznać i nadać znaczenie. Dzięki odziedziczonym informacjom rozpoznaje znaki środowiska i może odpowiadać na nie, wyzwalając odpowiednie zachowanie i zaspokajając wymagania własnego *Innenweltu*²³⁵. *Umwelt* stanowi tu percypowany świat organizmu a *Innewelt* to jego stan wewnętrzny, który wyłapuje pojawiające się znaczenie w informacjach płynących ze środowiska.

Działania inteligentne są wyrazem ciągłego procesu uczenia się, zbierania doświadczeń i rozwoju, dzięki któremu system definiuje i redefiniuje swoje postrzeganie i działanie w świecie, odpowiednio dostosowując swoje zachowanie²³⁶. Oznacza to, że inteligencja jest subiektywnie związana z otoczeniem i jej pojawienie się wynika z właściwości agenta, powstających w wyniku ustalonej wciąż na nowo perspektywy. Z ewolucyjnego punktu widzenia, forma systemu wydaje się wynikać z funkcji, ponieważ istnienie określonej struktury fizycznej jest zdeterminowane przez specyficzne potrzeby funkcjonalne, które kształtowały organizm i z którymi zmagał się w trakcie historii ewolucyjnej i rozwojowej²³⁷. Można więc przyjąć, że podobna sytuacja powinna mieć miejsce podczas konstruowania sztucznego agenta. Jego architektura kognitywna i zewnętrzna forma fizyczna powinny być zdeterminowane przez cechy jego *Umweltu* i sposób, w jaki ma z nim oddziaływać, przy zachowaniu autonomii inteligentnego zachowania wymuszanej przez stany wewnętrzne *Innenweltu*. W drugim rozdziale przedstawiono semiozę jako główne źródło powstawania inteligentnych zachowań. Wykorzystanie teorii biosemiotycznych w SI ma na celu przezwyciężenie istotnego ograniczenia sztucznych symboli związanych z brakiem własnego znaczenia (*symbol*

234 E. Cassirer, *The Philosophy of Symbolic Forms. Vol. 4...*, op. cit., t. 4, s. 282.

235 J. von Uexküll, *Umwelt und Innenwelt der Tiere*, Berlin 1921.

236 M.I.A. Ferreira, *On Meaning: A Biosemiotic Approach*, „Biosemiotics”, t. 3, nr 1, 2010, ss. 107–130.

237 J. Hoffmeyer, K. Kull, *Baldwin and biosemiotics: What intelligence is for*, [w:] *Evolution and learning: The Baldwin effect reconsidered*, B. Weber, D. Depew (red.), Cambridge 2003, ss. 253–272.

grounding problem). Umieszczenie procesu semiozy w centrum modelu podkreśla istotną rolę, jaką odgrywa proces przypisywania znaczenia w tworzeniu inteligentnego zachowania. Natomiast uwzględnienie roli *Umweltu* i *Innenweltu* daje podstawę do kształtowania zachowań sztucznego, autonomicznego systemu.

Sztuczne systemy biorą udział w procesach związanych ze sferą znaczenia, ponieważ operują znakami/symbolami. Przeciwnicy sztucznej inteligencji przeczą, by zachodziły tam procesy semiotyczne, twierdząc, że urządzenia cyfrowe mogą przetwarzać tylko sygnały na automatyczne reakcje²³⁸. Brakuje więc tu charakterystycznego dla semiozy triadycznego przetwarzania znaków. Newell przedstawił przetwarzanie symboli w komputerach jako proces odnoszący się do dwóch symboli fizycznych, X i Y, gdzie X koresponduje z Y²³⁹. Jest to tylko diadyczna relacja, ponieważ brakuje w niej znaków, które rozwinęłyby się jako denotacyjne stosunki znaczeniowe. Przetwarzanie znaków w komputerach stanowi podstawowy poziom, który redukuje znaczenie do sygnalizacji elektronicznej, a zatem quasi-semiozy, jakby określił to Peirce²⁴⁰. Z drugiej jednak strony, semioza maszynowa może być niedostatecznie opisana, podobnie jak semioza zachodząca w działaniu mózgu, gdzie procesy są opisywane jako sekwencje dyskretnych sygnałów, które występują jako wejścia i wyjścia miliardów komórek nerwowych²⁴¹.

Peirce, w swojej teorii maszyn logicznych wymieniał podobieństwa między ludźmi i maszynami twierdząc, że ludzki umysł w pewnych aspektach działa jak maszyna, np. rozwiązując zadanie, które może rozwiązać maszyna logiczna lub obliczeniowa, czyli gdy ściśle przestrzega reguł z góry określonego algorytmu²⁴². Peirce zauważył, że ludzki umysł, rozwiązując problem matematyczny lub logiczny, działa jak maszyna umysłowa. Zatem należałoby uznać, że maszyny obliczeniowe i logiczne są maszynami rozumującymi, ponieważ wykonują pracę umysłową podobną do pracy umysłowej człowieka. Wilfried Nöth twierdzi, że to podobieństwo między rozumowaniem rachunkowym a mechanicznym rozumowaniem można wyjaśnić wspólnym dziedzictwem

238 J.R. Searle, *Minds, brains, and programs*, „Behavioral and Brain Sciences”, t. 3, nr 03, 1980, ss. 417–424.

239 A. Newell, *Physical symbol systems*, „Cognitive Science”, t. 4, nr 2, 1980, ss. 135–183.

240 C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 5: Pragmatism and Pragmaticism*, Charlottesville, Va 1994, akap. 473.

241 U.T. Place, *Is consciousness a brain process?*, [w:] *The Mind-Brain Identity Theory*, C. V. Borst (red.), London 1970, ss. 42–51.

242 C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 2: Elements of Logic*, Charlottesville, Va 1994, akap. 59.

ewolucyjnym natury biologicznej i fizycznej: zarówno mózg ludzki, jak i działanie maszyn ewoluowały przy tych samych ograniczeniach fizycznych – tak pojawił się schemat przetwarzania znaków wspólny dla ludzi i maszyn²⁴³.

Olbrzymia większość sztucznych systemów to maszyny Turinga, nie posiadające samoodniesienia. To stanowi kolejny problem do rozwiązania w kontekście badania semiozy maszyn. Takie systemy mogą wchodzić w interakcję ze środowiskiem, ale nie mają możliwości odniesienia się do siebie samych. Samoidentyfikacja jest biologiczną własnością żywych organizmów, które by przetrwać w swoim środowisku, muszą mieć zdolność do rozróżniania między samymi sobą i zewnętrznymi obiektami środowiska. Zdolność do samodzielnego odniesienia się i autonomicznego odniesienia do swojego otoczenia, do samonaprawy i do samoodtworzenia są warunkami autopoezy²⁴⁴. Sztuczne systemy nie są systemami autopoietycznymi, ale allopoietycznymi, o ile są wytwarzane, programowane i serwisowane przez człowieka²⁴⁵.

Dystans między semiozą żywego organizmu i maszyny maleje, co zauważył już w 1998 roku Peter Cariani²⁴⁶. Część biologów uważa, że rozróżnienie między organizmami a maszynami ostatecznie zniknie. Nawet jeśli nie zostanie osiągnięta pełna synteza żywej komórki, w przyszłości powstaną sztuczne struktury organiczne, które będą pełniły funkcje komórek lub tkanek i narządów, gdzie granica między działaniem maszyny a istoty żywej zostanie zatarta²⁴⁷. W działaniu prawdziwej semiozy nie najważniejsze jest jednak, w jaki konkretny sposób wynik zostaje osiągnięty, lecz czy będzie miał pewien ogólny

243 W. Nöth, *Sign machines in the framework of Semiotics Unbounded*, „Semiotica”, t. 2008, nr 169, 2008, ss. 319–341.

244 F. Maturana, W. Shen, D.H. Norrie, *MetaMorph: an adaptive agent-based architecture for intelligent manufacturing*, „International Journal of Production Research”, t. 37, nr 10, 1999, ss. 2159–2173.

245 To rozróżnienie jest często kwestionowane, nie tylko na polu inżynierii, lecz także w biologii, psychologii i filozofii. Autonomia ludzkiej świadomości była poddawana w wątpliwość. Rozwój biologii, psychologii ewolucyjnej, współczesnej genetyki powoduje burzliwe dyskusje dotyczące autonomii ludzkiej jaźni i świadomości. Patrz. S. Harter, *The construction of the self: A developmental perspective*, Guilford Press 1999; R.M. Ryan, E.L. Deci, *Self-determination theory: Basic psychological needs in motivation, development, and wellness*, Guilford Publications 2017.

246 P. Cariani, *Towards an evolutionary semiotics: the emergence of new sign-functions in organisms and devices*, [w:] *Evolutionary systems*, Springer 1998, ss. 359–376.

247 Y. Kawade, *The two foci of biology: Matter and sign*, „Semiotica”, t. 127, nr 1–4, 1999, ss. 369–384.

charakter²⁴⁸. Możliwe jest przetwarzanie znaku znak przez normę lub zwyczaj semiotyczny, które nie musi być naśladowany w każdym momencie procesu, lecz pozwala na kreatywność w tworzeniu i przekształcaniu znaków.

Wraz z rozwojem systemów komputerowych, automatów i robotów, powstają maszyny, które nie wydają się już zwykłymi artefaktami alopoetycznymi, ale zaczynają nosić znamiona systemów autopoietycznych. Najlepszego przykładu istnienia takich agentów dostarczają badania nad sztucznym życiem, gdzie istnieje możliwość tworzenia sztucznych organizmów zdolnych do samodzielnego utrzymania się, a nawet do samoodtworzenia. Pewne rozwiązania znajdują też zastosowanie w dziedzinie robotyki. Szczególnie modelowanie procesów związanych z ucieleśnieniem może przyczynić się do rozwijania sztucznej inteligencji. Efekty ucieleśnienia były dostrzegane w wielu zadaniach poznawczych, w tym na przykład w obserwacjach działania, postrzegania obiektów, pamięci i przetwarzaniu języka. Ponadto, w porównaniu z wczesnymi badaniami SI, ucieleśnione modele obliczeniowe wykazują potencjał do uzyskania bardziej rozwiniętych zachowań, ponieważ przypominają realistyczne interakcje agenta i jego środowiska. Co więcej, ucieleśnione teorie poznania dostarczają możliwości kreowania wyrafinowanych zdolności poznawczych sztucznych systemów, których inne metody projektowania nie były w stanie zrealizować.

Organizm jest specyficznym rodzajem systemu fizycznego, ponieważ zawiera pewne funkcjonalne częściowe relacje, oparte na naturalnych systemach znaków, jak chociażby kod genetyczny. Pojęcie semiozy zachodzące w maszynach jest krytykowane za to, że jest zbyt wąskie lub nawet mentalistyczne. Proces semiozy opisany przez Peirce'a jest dużo bardziej wiarygodny i satysfakcjonujący. Jednak semioza oznacza, że coś odnosi się do czegoś w specyficzny triadyczny sposób, czyli znak odnosi się do przedmiotu i interpretatora. Jednak znak może oznaczać coś również dla innego systemu interpretacji niż umysł żywej istoty (np. dla komórki, dla powstania opinii, dla sztucznego agenta), gdzie odniesienie jest pośrednictwem w celu wywołania skutku (zatem interpretatorem) w jakimś systemie. Tak więc, semioza, czyli proces przetwarzania znaków, zawsze obejmuje nieredukowalnie triadyczny proces między znakiem, obiektem i interpretatorem.

248 C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 1: Principles of Philosophy*, Charlottesville, Va 1994, akap. 211.

Należy zatem przeprowadzić bardziej szczegółową analizę różnych systemów, aby zrozumieć odmienny sens semiozy. W tym celu biosemiotyka jest cenny, ponieważ stara się odejść od perspektywy antropocentrycznej. Ucieleśnione sztuczne systemy pozostają mechanizmami w sensie technicznym, lecz radykalnie różnią się od rodzaju mechanizmu, o którym dyskutowali Kartezjusz, Peirce czy Uexküll. Przyjęcie teorii *Umweltu* we współczesnych badaniach nad budową inteligentnych maszyn wymusza oddzielną epistemologię dla sztucznego systemu, który nie jest interpretowany przez kategorie antroposemiozy²⁴⁹. Porównania odmiennych systemów skutkują błędnym przenoszeniem cech i funkcji jednego przedmiotu badań na inny a także uniemożliwiają rozwój prac nad sztucznymi systemami. Jak więc widać status semiotyczny sztucznych systemów nie może być do końca sprecyzowany, niemniej jednak nie można całkowicie odrzucić możliwości, że w pewnym stopniu w sztucznych systemach może zachodzić semioza.

3.1.2 Koncepcja ucieleśnienia i usytuowania sztucznych systemów

Zadaniem inspirowanych biologicznie maszyn jest odczytywanie znaków i przetwarzanie je w zgodzie z ich systemem percepcyjnym. Sztuczne systemy, które mogą przetwarzać symbole i interpretować je w konkretnym kontekście, bazują na wiedzy i doświadczeniu zdobytych w trakcie interakcji ze środowiskiem. W *Umweltcie* każdy składnik posiada funkcjonalne znaczenie dla organizmu. Krąg funkcjonalny określa sposób interakcji organizmu z przedmiotu działania. W kręgu funkcjonalnym dane sensoryczne są przekształcane w konkretne znaczenie. Sztuczny system inspirowany działaniem kręgów funkcjonalnych powinien identyfikować znaki i nadawać sens sygnałom płynącym ze środowiska, odpowiadając na nie adekwatnym zachowaniem. Może to być podstawą do odtworzenia w sztucznym systemie dynamiki procesów poznania prowadzących do autonomii a także inteligentnego zachowania. Zarówno działanie w *Umweltcie* poprzez kręgi funkcjonalne jak i stany wewnętrzne *Innenweltu*, wpływają na funkcjonowanie każdej fizycznej jednostki. Dzięki sprzężeniu danych z *Umweltu* i reprezentacji z *Innenweltu* dane płynące ze środowiska zyskują znaczenie zapewniając poprawne odpowiedzi, które mogą być odczytane jako inteligentne zachowanie.

249 T. von Uexküll, *The Sign Theory of Jakob von Uexküll*, [w:] *Classics of Semiotics*, M. Krampen i in. (red.), Berlin, 1987, ss. 147–179.

Fizyczna forma maszyny i jej architektura poznawcza są konstruowane w oparciu o przewidywane sposoby oddziaływania ze środowiskiem. Inteligentny sztuczny agent poprzez analogię z organizmem musi posiadać zdolność do samodzielnego rozwoju oraz przystosowywania do środowiska mechanizmów działania. Ucieleśniona maszyna posiada układ sensomotoryczny (kamery, mikrofony, czujniki podczerwieni, sonary laserowe czujniki zasięgu, GPS, koła, chwytaki, manipulatory itd.), którego elementy wchodzi w fizyczną interakcję ze światem zewnętrznym, tworząc semantyczną interakcję. Możliwe jest zatem odtworzenie schematu połączeń między procesami percepcyjnymi i procesami aktywności. W ten sposób zbudowane systemy biorą udział w procesach znaczenia, przypisując je bodźcom środowiska, a następnie udzielając odpowiedzi, co pozwala wytworzyć związek między percepcją a działaniem²⁵⁰.

Problem stojący przed pracami nad sztuczną inteligencją jest związany z ucieleśnieniem i usytuowaniem sztucznych systemów. Mimo, iż większość badaczy zgadza się, że są to podstawowe zasady funkcjonowania w realnym świecie, istnieją różne interpretacje tej idei. Przede wszystkim nie istnieje kompromis co do roli i znaczenia ucieleśnienia. Niektórzy przedstawiciele dziedziny są mocno oddani koncepcji fizycznego ucieleśnienia, a tym samym fizycznego usytuowania w rzeczywistym świecie, innym natomiast wystarcza symulowanie sztucznych organizmów w fizycznym środowisku stworzonym przez programistę²⁵¹. Ci drudzy twierdzą, że bezcielesne systemy bez ciała w sensie fizycznym mogą być inteligentne, ponieważ są ucieleśnione w usytuowanym poczuciu bycia autonomicznymi agentami strukturalnie związanymi z ich środowiskiem²⁵². W interakcji ucieleśnionego i osadzonego agenta i środowiska powstają mechanizmy znaczeniowe, które mogą ukształtować fenomenalny *Umwelt*. Agent jako aktywny podmiot wykorzystuje środowisko jako zasób wiedzy niezbędnej do działania. Zatem taki sztuczny system jako osadzony w kontekście sytuacji istnieje w świecie rzeczywistym, w którym przypisuje znaczenie bodźcom. Dlatego można powiedzieć, że ucieleśnieni agenci, mimo iż są ugruntowani we własnej architekturze i oprogramowaniu, posiadają ciało, które pasuje nie tylko do ich zadań, lecz przede wszystkim do ich świata. Rozważania nad

250 R.A. Brooks, *Elephants don't play chess*, „Robotics and Autonomous Systems”, t. 6, nr 1–2, 1990, ss. 3–15.

251 G. Pezzulo i in., *The Mechanics of Embodiment: A Dialog on Embodiment and Computational Modeling*, „Frontiers in Psychology”, t. 2, 2011.

252 S. Franklin i in., *Artificial minds*, „Computers in Physics”, t. 11, nr 3, 1997, ss. 258–259.

ucieleśnieniem sztucznych systemów z perspektywy biosemiotycznej wymusza zrozumienie różnicy między pojęciem organizmu a pojęciem ciała. Żywy organizm funkcjonuje w fenomenalnym *Umwelcie* poprzez aktywność poznawczą, postrzeganie, przyjmowanie konkretnej perspektywy oraz tworzenie specyficznej historii, łączącej doświadczenia z działaniem. Ciało zaś może być traktowane jako materialna rzecz lub rodzaj cielesnej maszyny. Posiadanie ciała jest warunkiem wstępnym bycia organizmem. Nic ani nikt nie może być ucieleśniony, jeśli nie jest materialnym bytem współistniejącym ze światem egzystencjalno-fenomenalnym²⁵³.

Naturalną zdolnością każdego żywego organizmu jest podtrzymanie homeostazy systemu. Zdolność utrzymywania stałych parametrów wewnętrznych oraz skoordynowanie funkcjonowania ciała ze środowiskiem, jest warunkiem utrzymania samoorganizacji oraz wynika z aktywnego i stałego rekompensowania uszkodzeń²⁵⁴. Fizycznie ucieleśnione maszyny pozostają heteronomiczne; ich części są wytwarzane podczas procesów zewnętrznych, które są od nich niezależne. W przeciwieństwie do sztucznych systemów, organizmy kształtują się odśrodkowo w trakcie embriogenezy²⁵⁵. Maszyny nie mogą, gdy ulegną uszkodzeniu, same naprawiać się ani się regenerować. Żywe organizmy mogą, ponieważ posiadają materiał protoplazmatyczny, który może zostać wykorzystany do autonomicznego naprawienia uszkodzeń. Ten aspekt cielesności żywych organizmów znacząco odróżnia je od fizycznie ucieleśnionych sztucznych systemów i sprawia, że różnią się one radykalnie od żywych odpowiedników.

Można to podsumować mówiąc, że maszyny działają zgodnie z planami ludzkich projektantów, podczas gdy żywe organizmy są planami działania²⁵⁶. Uexküll opisał jako zasadniczą różnicę między konstrukcją mechanizmu a żywym organizmem: „organy żywych istot mają wrodzoną jakość znaczenia, w przeciwieństwie do części maszyny; dlatego mogą się rozwijać tylko odśrodkowo”²⁵⁷. Mimo ich zdolności do pewnej

253 T. Ziemke, N.E. Sharkey, *A stroll through the worlds of robots and animals: Applying Jakob von Uexküll's theory of meaning to adaptive robots and artificial life*, „Semiotica”, t. 134, nr 1/4, 2001, ss. 701–746.

254 C. Emmeche, *Organicism and qualitative aspects of self-organization*, „Revue Internationale de Philosophie”, nr 2, 2004, ss. 205–217.

255 R. Magnus, *Time-plans of the organisms: Jakob von Uexküll's explorations into the temporal constitution of living beings*, „Σημειωτική-Sign Systems Studies”, t. 39, nr 2–4, 2011, ss. 37–57.

256 J. von Uexküll, *Umwelt und Innenwelt der Tiere...*, op. cit., s. 301.

257 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten*, Rohwohlt Verlag, Hamburg 1934, s. 119.

samoorganizacji, sztuczni agenci są w rzeczywistości dalecy od posiadania cielesności żywych organizmów. Składają się głównie z części mechanicznych i programów sterujących. Dostarczanie maszynie elementów i sterowników można uznać za próbę zapewnienia im sztucznej ontogenezy. Jednak próba zapewnienia im w ten sposób autonomii jest skazana na niepowodzenie. Samosterowność, samokontrola i samodzielne działanie mogą zostać zaaplikowane w sztucznym systemie, jednak utrzymanie tożsamości, kształtowanie własnych celów i przetrwanie są dla badań SI dużym wyzwaniem.

3.1.3 Problem autonomii sztucznych systemów

Agent, który może korzystać z różnych danych sensorycznych i wykonywać odpowiadające im działania posiada wiele możliwości wyboru aktywności. Wykorzystanie modularnej architektury Brooksa, która rozkłada mechanizmy percepcji i działania na wiele czynników oraz integruje je i pośredniczy między nimi, pozwala wyłonić się samoadaptacji, która jest zależna od kontekstu w jakim system się znajduje. Sztuczny system może sam zmieniać swoje własne zachowanie w przyszłości w zależności od nabytego doświadczenia. Wykorzystanie sieci neuronowych, w których wagi połączeń odwzorowują sensomotoryczne dane mogą być dostosowane w każdym momencie. Takie mechanizmy sterowania neuronowego pozwalają wytwarzać adaptacyjne, zależne od kontekstu odpowiedzi na bodźce. Daje to szansę na wyłonienie się inteligentnych zachowań autonomicznych.

Istnieje spór dotyczący tego, co właściwie oznacza autonomia²⁵⁸. Jako że rozważania o autonomii maszyn rozpatrywano w cybernetyce²⁵⁹, największy nacisk położony jest na samosterowność i przeciwdziałanie jej utracie. Samosterowność oznacza niezależność od projektanta systemu i od zmian czynników środowiskowych. Istnieje wiele propozycji stworzenia obliczeniowych autonomicznych sztucznych systemów, np.

258 Propozycję zautonomizowania sztucznych systemów przedstawiłam również w: A. Sarosiek, *Kręgi funkcjonalne jako wzór dla modelowania procesów autonomicznych sztucznych systemów*, [w:] *Studia z filozofii informatyki*, K. Sołoducha, P. Stacewicz (red.), Warszawa 2018, ss. 171–190.

259 Zob. np. M. Mazur, *Cybernetyka i charakter*, Warszawa 1999.

samonaprawialne systemy Bickharda²⁶⁰, modelowanie celu i motywacji d’Inverno i Lucka²⁶¹, biologiczna teoria znaczenia Ziemke²⁶² lub taksonomiczna teoria agentów programowalnych Franklina i Graessera²⁶³. Autonomia sztucznych systemów była też rozważana przez Maturanę i Varełę w sensie wspomnianej już autopoezy²⁶⁴, co w świetle teorii biosemiotycznej wydaje się być najlepszym rozwiązaniem, ponieważ autopoetyczna maszyna powinna mieć zdolność do interakcji ze środowiskiem, wchodzenia w relacje z jego przedmiotami, transformacji i regeneracji²⁶⁵. Wszystkie teorie zgadzają się jednak, że autonomia opiera się na zachowaniu tożsamości, samokontroli, jedności systemu oraz zapewnia własne ustalanie celów, które mają zapewnić przetrwanie. Dlatego tak ważnym celem badań nad sztuczną inteligencją jest zapewnienie sztucznym systemom autonomii. Istnienie autonomii jest niezbędne do powstania inteligentnego systemu, ponieważ odnosi się do samorządności i samostanowienia, niezbędnych dla intencjonalnego podmiotu.

Nie ma wątpliwości, co do tego, że biologicznie inspirowane maszyny są atrakcyjnymi modelami autonomicznych organizmów. Wielu badaczy krytykowało jednak możliwość autonomii sztucznych systemów²⁶⁶. Częstym zarzutem było twierdzenie, że nie są „naprawdę” autonomiczne, ponieważ charakter autonomii przypisuje się im tylko jako funkcję modelu. Ich autonomiczne zachowania są zbyt proste w porównaniu z zachowaniami żywych organizmów, ponieważ nie wykazują żadnych innych umiejętności

260 M.H. Bickhard, *Autonomy, function, and representation*, „Communication and Cognition-Artificial Intelligence”, t. 17, nr 3–4, 2000, ss. 111–131.

261 M. d’Inverno, M. Luck, *Creativity through autonomy and interaction*, „Cognitive Computation”, t. 4, nr 3, 2012, ss. 332–346.

262 T. Ziemke, N.E. Sharkey, *A stroll through the worlds of robots and animals: Applying Jakob von Uexküll’s theory of meaning to adaptive robots and artificial life...*, op. cit.

263 S. Franklin, A. Graesser, *Is It an agent, or just a program?: A taxonomy for autonomous agents*, [w:] *Intelligent Agents III Agent Theories, Architectures, and Languages*, Springer, Berlin, Heidelberg 1996, ss. 21–35.

264 W języku polskim istnieje również określenie autopoeiza, autopojeza. Patrz. *System autopojetyczny*, [w:] *Wikipedia, wolna encyklopedia*, 2018.

265 H.R. Maturana, *Autopoiesis and Cognition: The Realization of the Living*, Springer Science & Business Media 1980.

266 Patrz W. Christensen, C. Hooker, *Anticipation in autonomous systems: foundations for a theory of embodied agents*, „Int J Comput Anticip Syst”, t. 5, 2000, ss. 135–154; H.R. Maturana, *Autopoiesis and Cognition...*, op. cit.; C.T.A. Schmidt, F. Kraemer, *Robots, Dennett and the Autonomous: A Terminological Investigation*, „Minds and Machines”, t. 16, nr 1, 2006, ss. 73–80; T. Ziemke, N.E. Sharkey, *A stroll through the worlds of robots and animals: Applying Jakob von Uexküll’s theory of meaning to adaptive robots and artificial life...*, op. cit.

poza zaprogramowanymi. Jest to jednak sposób definiowania autonomii sztucznego systemu w opozycji do żywego organizmu. Tymczasem biosemiotyka przedstawia działanie podmiotu jako unikatowe, które nie może być porównywane z działaniem żadnego innego podmiotu należącego do innego gatunku. Dlatego radykalne porównywanie sztucznego i żywego systemu jest zbędne, ponieważ nie zachodzi tam jakościowe zróżnicowanie poziomu autonomii. Podstawą autonomicznego działania jest tworzenie własnych niezależnych relacji ze światem. Maszyna skonstruowana w zgodzie z założeniami biosemiotycznymi, może brać udział w procesach semiozy, które stanowiłyby podstawę jej aktywności w *Umwelt*cie i gwarantowały jej niezależność działań. Autonomia i subiektywność żywych systemów wyłania się z interakcji ich składników (komórek, tkanek i narządów). Zarówno Uexküll oraz Maturana i Varela twierdzili, że interakcja między nimi a ich otoczeniem jest wynikiem tej strukturalnej zgodności. Zatem autonomia wyłania się jako własność organizacji żywego organizmu, już na samym początku, jako zgodność komórkowa, czyli harmonijne działanie w środowisku wynika z udanej reprodukcji.

Autonomia w perspektywie technicznej i biologicznej posiada odmienne znaczenia. Biolodzy podkreślają znaczenie mechanizmów zapewniających przetrwanie, inżynierowie szukają sposobów uzyskania niezależności maszyny od projektanta. Obydwa punkty widzenia może połączyć perspektywa biosemiotyczna. Tom Ziemke proponuje, by uznać poziomy istnienia autonomii: w silnym sensie biologicznym, który uwzględnia te rozmyte pojęcia oraz w słabszym sensie, który może zostać zastosowany do sztucznych systemów²⁶⁷. Daje to możliwość pogodzenia biologicznego i inżynierskiego stanowiska, bez zgłębiania trudnych filozoficznych pojęć związanych ze świadomością. Pewne rozwiązania daje teoria biosemiotyczna, która proponuje sposób interpretacji i przetwarzania znaków przez żywy organizm. Szczególnego rozwiązania dostarcza teoria działania kręgów funkcjonalnych²⁶⁸, która być może pozwoli lepiej zrozumieć, jak pojęcie autonomii może być postrzegane poprzez działanie aktywnego podmiotu, nie tylko żywego, lecz również sztucznego.

267 T. Ziemke, *On the role of emotion in biological and robotic autonomy*, „Biosystems”, t. 91, nr 2, 2008, ss. 401–408.

268 J. von Uexküll, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten*, op. cit.

Sztuczne organizmy mogą przypominać pod wieloma względami ich żywe odpowiedniki. Zastosowanie sieci koneksjonistycznych w miejsce układu nerwowego, oraz umiejętność uczenia się przez sieć sprawia, że system taki ma możliwość bycia autonomicznym. Może wykorzystać swoją pamięć i dostosować zachowanie do zmieniającego się środowiska. Samoorganizacja pozwala sztucznemu układowi reagować na bodźce płynące ze środowiska w sposób, który nie jest zaprogramowaną czy mechaniczną reakcją. Zamiast użyć ogólnych procedur, system autonomiczny reaguje w zgodzie ze swoim doświadczeniem. Ponadto można zauważyć, że udział maszyny w semiozie jest jak najbardziej rzeczywisty. Maszyna samodzielnie i prywatnie przetwarza znaki. Proces ten zachodzi nie tylko w sposób niezależny od projektanta, lecz również często niespodziewany i odmienny niż założony. Procesy przetwarzania znaków stają się niezależne od inżynierów i od interpretacji obserwatorów. Sztuczne organizmy zaczynają brać udział w interakcjach z otoczeniem w sposób samodzielny. Zastosowanie technik ewolucyjnych pozwala na samodzielny rozwój programu. Umiejętność dekompozycji zadań jest wynikiem adaptacji i zmian powstałych w modułowej architekturze. Aktywność oparta na modułach pod wieloma względami przypomina kręgi funkcjonalne, działające zarówno jako oddzielne struktury jak i całościowo.

Autonomia istot żywych zawiera się w takich właściwościach jak umiejętność przetrwania, regeneracji uszkodzeń i reprodukcji. Umiejętności te są wynikiem semiotycznych interakcji z otoczeniem kształtowanym poprzez lata ewolucji. Badania nad inteligentnymi systemami autonomicznymi wzorują się na biologicznie inspirowanych zasadach projektowania, ponieważ umożliwia to znalezienie sposobu na rozwiązanie problemu w sposób bardziej naturalny. Zachowanie żywych organizmów da się często wyjaśnić prostymi potrzebami, a złożoność ich zachowania zależy od złożoności sygnałów odbieranych ze środowiska. Dla żywego organizmu zdolność do adaptacji i elastyczność zachowania wynikają z odmiennych procesów niż odczyt odpowiedzi lub reguł z bazy wiedzy przy użyciu gotowego skryptu.

Inteligentne sztuczne systemy powinny być zdolne do wykonywania zadań samodzielnie. Musi to wynikać z autonomicznego procesu semiozy, która dostarcza im znaczenia, niezbędnego do wyboru działań. Istotną cechą sztucznego semiotycznego systemu powinna być umiejętność uczenia się w trakcie działania w środowisku. Dopiero agent, który uczy się poprzez doświadczenie w interakcji ze środowiskiem i reaguje zmianą własnych programów w celu poprawy swojej przyszłej wydajności, nie jest już

maszyną deterministyczną, ale prawdziwie semiotyczną²⁶⁹. Wiele z aspektów semiozy jest obecne w sztucznych systemach. Dodatkowo w trakcie rozwoju badań z zakresu sztucznej inteligencji powstało tak wiele rodzajów systemów przejawiających inteligentne zachowanie, że każde pojedyncze kryterium semiotyczne jest obecne w różnych poszczególnych tworcach sztucznej inteligencji. Pozostałe różnice między semiozą żywego a sztucznego systemu są kwestią stopnia. Ta różnica stopnia jest szczególnie widoczna, jeśli weźmiemy pod uwagę kreatywność semiotyczną²⁷⁰.

Autonomia sztucznych agentów postrzegana jako cecha wyższego poziomu, wynika z samorozwoju i adaptacji ucieleśnionego systemu. Dzięki przyjęciu perspektywy zakładającej istnienie poziomów autonomii, można badać strukturę systemu, który może zrealizować autonomiczne zachowanie, nie rozważając sposobu jego powstania i rozwoju. Modelowanie autonomicznych zachowań w sztucznych systemach przy użyciu teorii kręgów funkcjonalnych Uexküll'a, pokazuje jak powstaje dynamiczny związek systemu i jego środowiska. Dodatkowo uwzględnia się fakt, że percypowanie świata przez system znacząco różni się od percepcji ludzkiego konstruktora. Autonomia sztucznych systemów nie powinna pokrywać się z pojęciem autonomii żywego organizmu. Tym bardziej, że jeśli w sztucznym systemie miałyby istnieć autonomia w silnym biologicznym sensie, jego programista/konstruktor/użytkownik nie miałby do niego dostępu a próba opisu jego działania pozostałaby eksternalistyczna i niewiele wyjaśniająca.

3.1.4 Znaczenie zgodności strukturalnej percepcji i działania

Sprzężenie percepcji z działaniem ma dla systemów autonomicznych kluczowe znaczenie, ponieważ informacje zwrotne pozwalają kontrolować i modyfikować funkcjonowanie. Ucieleśniony i usytuowany agent otrzymuje wszystkie dane ze swoich czujników. Każda informacja przetwarzana przez system wpływa na całościowe działanie systemu. Struktury działające oddolnie bezpośrednio wchodzą w relacje ze środowiskiem. Wyższe poziomy interpretują, lecz nie determinują ich działania. Interpretacja następuje w oparciu o fizycznie ugruntowane hipotezy, powstałe w wyniku zgromadzonej wiedzy o

269 W. Noth, *Semiotic machines...*, op. cit.

270 *Ibid.*

świecie²⁷¹. Implementacja symbolicznej reprezentacji świata w systemie zostaje zastąpiona danymi o świecie, który sam w sobie jest modelem²⁷². Znaczenie zawarte w informacjach płynących ze środowiska może być poddane analizie. Poszczególne elementy systemu mogą odpowiadać na poszczególne bodźce płynące ze środowiska. Pojedyncze moduły systemu działają samodzielnie wykrywając sygnały, obliczając, modelując procesy i uruchamiając układy wykonawcze. Mogą się kontaktować z innymi, lecz informacje, które są przesyłane, są ubogie w dane. Skomplikowane zadanie może być podzielone między kilka modułów, mimo że nie pracują ani hierarchicznie ani liniowo²⁷³.

Dwukierunkowe dopasowanie między organizmem a *Umweltem* jest tym, co Maturana i Varela określają jako strukturalną zgodność między nimi, co jest wynikiem ich strukturalnego sprzężenia. Ontogeneza w tej perspektywie oznacza historię zmian strukturalnych przy pozostawianiu jednością²⁷⁴. Zmiany strukturalne wywołane są przez interakcje ze środowiskiem albo jako wynik wewnętrznej dynamiki systemu. W interakcjach struktura środowiska wyzwała zmiany strukturalne w jednostkach autopoetycznych, ale nie określa ich ani nie ukierunkowuje. Rezultatem jest historia zgodnych zmian strukturalnych, dopóki organizm bądź jego otoczenie nie ulegną dezintegracji²⁷⁵.

Powyższa charakterystyka organizmu jest dokładnie tym, co Uexküll opisał jako podmiot naznaczający swoją jakość ego na obiektach, w zależności od aktualnego zachowania. Dzieje się to również w powtarzających się neuronowych kontrolerach robotów: w każdym czasie bodźce przychodzące są interpretowane na podstawie aktualnych odchyłeń i odwzorowywania sensoryczno-motorycznego. Agent znajduje się w konkretnej sytuacji i kontekście, które weryfikuje doświadczając strukturalnego sprzężenia ze środowiskiem. Tak jest w przypadku zarówno kleszcza Uexkülla jak i autonomicznego agenta. W obu przypadkach głównym mechanizmem odnajdywania się w konkretnej

271 S. Harnad, *Computation is just interpretable symbol manipulation; cognition isn't*, „Minds and Machines”, t. 4, nr 4, 1994, ss. 379–390.

272 K.J.W. Craik, *The nature of explanation*, Cambridge 1967, s. 61; D. Favareau, *Essential Readings in Biosemiotics: Anthology and Commentary*, Springer Science & Business Media 2010, s. 265.

273 R.A. Brooks, *A robust layered control system for a mobile robot...*, *op. cit.*

274 H.R. Maturana, F.J. Varela, *The tree of knowledge...*, *op. cit.*, s. 74.

275 *Ibid.*

sytuacji jest sprzężenie strukturalne, które za pomocą układu sensomotorycznego, dopasowuje się do różnych sytuacji i wymagań²⁷⁶.

3.1.5 Biosemiotyczna propozycja rozwiązania problemu ugruntowania symboli

Przyjęcie perspektywy biosemiotycznej pozwala na analizę semiozy we wszelkich układach system-środowisko. Relacje semantyczne odnoszą się do warunków lub zasad, zgodnie z którymi możliwe jest powstawanie znaczenia. Jak twierdził Charles Morris, jeśli coś kontroluje zachowanie w określonym celu, który jest podobny, choć niekoniecznie identyczny ze sposobem, w jaki coś innego będzie kontrolować zachowanie w odniesieniu do tego celu, wtedy mówimy o działaniu znaku²⁷⁷. W badaniach nad sztuczną inteligencją należy się skupić na stworzeniu i utrzymaniu relacji danych otrzymanych ze środowiska agenta z jego symbolicznymi reprezentacjami. Dzięki temu sygnały sensoryczne pozwolą na sklasyfikowanie otrzymanych danych. Znaki w biosemiotyce stanowią obiekty o różnych właściwościach referencyjnych, które są wykorzystywane przez wszystkie organizmy do wielu celów. Co jednak najważniejsze, ich użycie nie ogranicza się do celów opisowych, lecz do każdej możliwej aktywności, którą chcielibyśmy aplikować w sztucznych systemach np. do nawiązywania relacji, nawigacji, utrzymywania uwagi i autonomicznego działania.

W badaniach nad SI, techniki kreowania autonomicznych agentów zaczynają się oddzielenia znaków od podstawowej reprezentacji. Podstawową sposobem jest tworzenie etykiet, które początkowo są reprezentowane w umyśle architekta systemu, a następnie aplikowane sztucznemu systemowi²⁷⁸. Jest to próba wykorzystania funkcji klasycznego trójkąta semiotycznego, w którym znak (symbol) oznacza referent (obiekt) poprzez swój związek z pojęciem (znaczeniem). Problem techniczny w polega na utworzeniu odpowiedniej reprezentacji zapewniającej działanie tego połączenia. W symbolicznym

276 T. Ziemke, *On 'parts' and 'wholes' of adaptive behavior: Functional modularity and diachronic structure in recurrent neural robot controllers*, [w:] *From animals to animats*, t. 6, 2000, ss. 171–180.

277 C. Morris, *Signs, language and behavior*, Oxford 1946, s. 84.

278 S. Harnad, *The symbol grounding problem*, „Physica D: Nonlinear Phenomena”, t. 42, nr 1, 1990, ss. 335–346.

podejściu znak służy jednak tylko do stworzenia idei referenta tylko w umyśle obserwatora.

Biosemiotyka proponuje inne podejście do semantyki aktów znakowych. Kalevi Kull sugeruje, że reprezentacje mogą być błędne lub poddawane w wątpliwość, ponieważ są ściśle powiązane z dynamicznie zmieniającymi się interakcjami agenta i środowiska²⁷⁹. Tymczasem w działaniu sztucznych agentów, reprezentacja związana jest z pożądanym i przewidzianym wynikiem interakcji. Uzyskanie semantycznej autonomii zakłada użycie przez agenta znaków, które posiadają znaczenie w odniesieniu do jego własnych celów lub działań. Znak jest związany z interakcją i jej wynikiem, a nie, jak w przypadku żywych organizmów, z działaniem znaku na układ sensoryczny. Symbole i znaki muszą mieć znaczenie ze względu na właściwości, które posiadają dla agenta w interakcjach ze światem.

Ugruntowanie symboli w sztucznych systemach powinno uwzględniać relatywność pojęć agenta. Budowanie sztucznego systemu powinno respektować pytanie, dlaczego reprezentowany obiekt jest ważny dla agenta. Znaki użyte w konkretnym znaczeniu mają charakter narzędziowy, ponieważ służą do osiągnięcia pożądanego wyniku interakcji²⁸⁰. Można tu użyć przykładu Heideggera opisującego użycie znaków przez kierowców: znaku zmiany kierunku i zatrzymania się. Użycie znaku w tym przykładzie jest oparte na znajomym wyniku interakcji²⁸¹. Sygnały te są konsekwencją nauki i doświadczenia kierowcy i sprawiają, że jazda samochodem jest łatwiejsza i bezpieczniejsza. Użycie znaków przez sztuczne systemy autonomiczne powinno wyglądać podobnie. Erich Prem rozważał przykład systemu, który uczy się wykrywać znaki w celu podjęcia decyzji i wskazywał kilka przykładów przedstawiających możliwe znaki, które mogą korzystnie wpłynąć na działanie sztucznych autonomicznych systemów²⁸².

279 K. Kull, *Semiosis stems from logical incompatibility in organic nature: Why biophysics does not see meaning, while biosemiotics does*, „Progress in Biophysics and Molecular Biology”, t. 119, nr 3, 2015, ss. 616–621.

280 T. Susi, T. Ziemke, *On the subject of objects: Four views on object perception and tool use*, „tripleC”, t. 3, nr 2, 2005, ss. 6–19.

281 M. Heidegger, *Bycie i czas*, Warszawa 1994, s. 111.

282 E. Prem, *A Biosemiotic Framework for Artificial Autonomous Sign Users*, [w:] <https://www.aaai.org/Papers/Workshops/2004/WS-04-03/WS04-03-007.pdf>.

znak	zachowanie
powitanie	agent reaguje na powitanie lub na innego agenta określonym zachowaniem
oznaczanie	agent zaznacza lokalizację lub obiekt, aby łatwiej je odnaleźć
ostrzeżenie	agent wysyła sygnał alarmowy, aby poinformować o niebezpieczeństwie lub zainicjować ucieczkę
ucieczka	agent reaguje ucieczką na sygnał ostrzegawczy
podążanie	agent porusza się w oznaczonym kierunku, znak może wyznaczać również umiejscowienie celu
cel	agent przechodzi do wyznaczonego miejsca; tutaj, tam, dom, stacja itp.

Rys.2. Powiązanie znaków z działaniem wg Ericha Prema

Interesującym aspektem tych procesów jest wyraźne zaznaczenie odniesienia dla sztucznego systemu. Znak wskazuje nie tylko na kierunek, ale reprezentuje go w symboliczny sposób. Dodatkowo wykorzystuje też znaki dostarczane przez innych. Procesy semiozy tutaj mają jednak nieco inny charakter niż w przyrodzie ożywionej. Znaki nie są zwykłymi bytami w stosunku do innych bytów, do których odnoszą się w konkretnym związku, lecz są traktowane przede wszystkim jako narzędzia do nadawania znaczenia zarówno samemu agentowi jak i innym. Udana interakcja polega na wskazaniu celu tej interakcji i skutkuje pożądanym rezultatem dla użytkowników znaku. Dlatego też w sztucznych systemach, sam referencyjny charakter znaku nie potrzebuje aplikacji pełnej semiotycznej denotacji występującej u żywych organizmów (a w szczególności człowieka) lecz ma doprowadzić do udanej interakcji bądź działania.

Procesy semiozy pozostają tu na poziomie zwykłych znaków, a nie języka. Należy zbadać w jaki sposób autonomiczne systemy tworzą znaki i wykorzystują je dla

osiągnięcia celu²⁸³. Dopiero wtedy można próbować opisać złożony system umiejętności komunikacyjnych, który może służyć wielu różnym celom, od wzywania pomocy aż do prowadzenia filozoficznych dysput. Obecnie rozważanie powstania języka sztucznego systemu, bardzo przypomina refleksję Wittgensteina, który stwierdził, że „gdyby lew umiał mówić, nie potrafilibyśmy go zrozumieć”²⁸⁴, ponieważ język jako system znaków i reguł w przypadku sztucznych systemów nie dostarcza nam na razie żadnego kryterium, które wskazywałoby na to, że możemy zrozumieć ten język.

Znaki muszą być utworzone przez agentów, dla innych agentów i używane między agentami. To oznacza, że różne konteksty interakcji muszą pojawić się w użyciu znaku. Dodatkowo wewnętrzne znaczenie, o ile powstanie, pozostaje niezależne od programisty. Znaki nie mogą pozostać prostymi przedmiotami, ale muszą wskazywać agentowi na coś. Dzięki temu można obserwować i sprawdzać zachowanie maszyny. Wewnętrzne przetwarzanie znaczenia pozostaje ukryte. Nawet w przypadku konstruowanych przez człowieka systemów śledzenie tworzonych w nich reprezentacji jest trudne ze względu na dynamiczną interakcję ze światem. Ramy konceptualne w systemie będą nieprzejrzyste dla obserwatora, ponieważ są wynikiem złożonego procesu adaptacji, którego celów lub kryteriów wyboru nie musimy nawet znać. Agent używający znaku będzie dążył do uzyskania odpowiedniego efektu. Głównym celem jest wygenerowane pożądanego działania, a nie wyjaśnianie całego procesu dla zewnętrznego obserwatora.

3.1.6 Rola rusztowań poznawczych i rozproszonego poznania

Inteligentny sztuczny system powinien radzić sobie ze zmieniającymi się kontekstami środowiska, nawet w sytuacjach braku dostępnej wiedzy bądź w wypadku deficytu umiejętności. Rusztowania poznawcze zapewniają działania poprzez interakcje semiotyczne z elementami środowiska, bodźcami lub sygnałami, które pojawiają w dynamicznych sytuacjach. Służą jako zewnętrzne repozytorium pamięci lub model pozwalający wykonywać złożone czynności i umożliwiają unikanie błędów. Ich rola polega na kierowaniu i kształtowaniu zachowania jako narzędzie do sterowania i

283 L. Steels, P. Vogt, *Grounding adaptive language games in robotic agents...*, *op. cit.*

284 L. Wittgenstein, *Dociekania filozoficzne*, B. Wolniewicz (tłum.), Państwowe Wydawnictwo Naukowe 1972, s. 313.

kontrolowania działania, jak również na przekazywaniu informacji między agentami²⁸⁵. Koncepcja rusztowania została podjęta w badaniach nad autonomicznymi agentami, w celu zapewnienia im zdolności do funkcjonowania w złożonym i nieprzewidywalnym świecie rzeczywistym. Podejście to pozwoliło na odrzucenie fundamentalnej części klasycznego podejścia, zakładającego istnienie głównej bazy danych, która zawiera wszystkie dostępne informacje²⁸⁶. Tym samym można było porzucić trud przekształcania nadchodzących informacji sensorycznych w pojedynczy kod symboliczny, aby mogły zostać wprowadzone do bazy.

Wzorce aktywności, których struktura wysokiego poziomu nie może być zredukowana do konkretnych sekwencji ruchów, mogą wynikać z interakcji między odruchami prostymi a środowiskiem, do którego są przystosowane. Wyłaniające się zachowanie systemu jako całości jest wynikiem różnych autonomicznych działań współdziałających ze sobą i z otoczeniem, a nie scentralizowanego systemu podejmującego decyzje w oparciu o reprezentowane wewnętrznie kierunki działań lub cele. Horst Hendriks-Jansen nazywa to zasadą „interaktywnego pojawiania się”²⁸⁷. Z punktu widzenia biosemiotyki, interakcyjnego zachowania nie można wyjaśnić w kategoriach prawa dedukcyjnego lub generatywnego. Wymaga to wyjaśnienia historycznego, ponieważ nie może być żadnych reguł przewidujących rodzaje zachowań, które mogą się pojawić²⁸⁸. Andy Clark natomiast zauważył, że ewoluujące istoty nie przechowują ani nie przetwarzają informacji w kosztowny sposób, kiedy mogą korzystać ze struktury środowiska²⁸⁹.

Sztuczni agenci będący aktywnymi użytkownikami środowiska i jego obiektów mogą otrzymać dostęp do zasobów, które pozwala im kształtować swoje zachowanie z zewnątrz. Prace nad poznaniem ucieleśnionym, usytuowanym i rozproszonym w badaniach nad inteligentnymi zachowaniami wyjaśnia, w jaki sposób agenci aktywnie wykorzystują swoje własne zachowanie i ruch w środowisku w celu uporządkowania sensorycznych danych wejściowe, dzięki którym mogą rozwiązywać problemy,

285 J. Hoffmeyer, *Semiotic Scaffolding of Living Systems*, [w:] *Introduction to Biosemiotics*, M. Barbieri (red.), Springer Netherlands 2008, ss. 149–166.

286 A. Clark, *Being there: Putting brain, body, and world together again...*, *op. cit.*, s. 21.

287 H. Hendriks-Jansen, *Catching ourselves in the act: Situated activity, interactive emergence, evolution, and human thought...*, *op. cit.*, ss. 8–9.

288 *Ibid.*, s. 9.

289 A. Clark, *Being there: Putting brain, body, and world together again...*, *op. cit.*, s. 46.

nierozwiązywalnych w prosty sposób rozwiązać za pomocą wewnętrznych procesów obliczeniowych. W przykładzie Ziemke agent wyposażony w czujniki zbliżeniowe i fotokomórkę zaczyna podróż na początku labiryntu w kształcie litery T. Napotyka tam źródła światła i po okresie opóźnienia musi skrócić w jego kierunku na skrzyżowaniu, aby osiągnąć cel²⁹⁰. Zadanie z opóźnioną odpowiedzią wymaga pewnej formy pamięci krótkotrwałej, ponieważ agent w tym czasie musi zatrzymać w pamięci informację, po której stronie było źródło światła. W eksperymencie okazało się jednak, że agenci ewoluowali, aby niezawodnie rozwiązać zadanie. Maszyny przesuwają się w kierunku źródła światła, gdy mogły je zobaczyć, a następnie podążały za odpowiednią ścianą aż do celu²⁹¹. Można zatem uznać, że agenci zamiast używać pamięci roboczej lub wewnętrznej reprezentacji używali ściany lub własnego odniesienia do niej.

Większość przetwarzania poznawczego odbywa się w interakcji ze środowiskiem. W systemach ewoluujących może powstać bardziej złożona aktywność, wzmacniana przez doświadczenie i naukę. Dzięki temu mogą powstać nowe procedury, do których stosuje się agent. Po pojawieniu się nowej informacji, zachodzi filtrowanie danych i ponowne wiązanie ich z danymi uzyskanymi w wyniku wcześniejszego doświadczenia. Działanie podlega testowi poprzez jego zastosowanie w środowisku²⁹². Autonomiczny agent wybiera spośród możliwych rozwiązań opierając swoje działania na wcześniejszych doświadczeniach zdobytych w poprzednim cyklu.

3.2 Badania nad inteligentnymi agentami

Podstawowe pojęcia typowe dla biosemiotyki to semioza, *Umwelt*, *Innenwelt*, *Wirkwelt*. Umieszczenie podstawowego procesu interpretacyjnego w centrum każdego z prezentowanych modeli podkreśla istotną rolę, jaką odgrywa proces przypisywania znaczenia w zjawisku poznania i w konsekwencji w produkcji inteligentnego zachowania. Poniżej przedstawione zostały zarówno historyczne już eksperymenty dotyczące

290 T. Ziemke, *Cybernetics and embodied cognition: on the construction of realities in organisms and robots*, „Kybernetes”, t. 34, nr 1/2, 2005, ss. 118–128.

291 *Ibid.*

292 W.P. Hall, *Biological nature of knowledge in the learning organisation*, „The Learning Organization”, t. 12, nr 2, 2005, ss. 169–188.

modelowania sztucznej inteligencji, jak i późniejsze próby modelowania inteligentnych zachowań w perspektywie biosemiotycznej.

3.2.1 Architektura subsumpcyjna Rodneya Brooksa

W latach dziewięćdziesiątych dwudziestego wieku zastosowano nowe koncepcje i metody, aby ulepszyć cybernetyczne podejście do projektowania sztucznych systemów. Wtedy też pojawiły się idee, mające prowadzić do powstania autonomicznych systemów, zwanych też agentami lub animatami²⁹³. W połowie lat osiemdziesiątych Rodney Brooks i jego grupa z laboratorium sztucznej inteligencji *Massachusetts Institute of Technology* przedstawili krytykę „paradygmatu myślenia deliberatywnego”, zakładającego, że inteligentne zadania muszą być realizowane przez procesy działające na symbolicznym modelu wewnętrznym (GOFAI)²⁹⁴. Dzięki tej krytyce, w wyniku której powstały nowe pojęcia, rozwinęły się prace nad nowymi zestawami technik modelowania i zasadami budowy, które pozwoliły kształtować procesy postrzegane jako ucieleśnione. Jedną z głównych inspiracji dla nowych badań, była teoria *Umweltu* Uexküll’a, która dowodziła, że istnieje ścisły dynamiczny związek między ciałem zwierzęcia a jego doświadczanym światem oraz że świat postrzegany przez zwierzę jest całkowicie odmienny od świata postrzeganego przez inne gatunki. Brooks zauważył wtedy, że również sztuczny system będzie „postrzegał” świat w zupełnie inny sposób niż jego architekt²⁹⁵. Pozwoliło mu to przyjąć nową perspektywę, której celem było odrzucenie antropomorfizmu w dziedzinie sztucznej inteligencji.

Brooks rozpoczął pracę nad inteligentnymi systemami robotycznymi, budując fizyczne i osadzone w środowisku maszyny, które działały dzięki pracującym równolegle modułom, z których każdy przypominał krąg funkcjonalny Uexküll’a²⁹⁶. Każdy z nich był połączony na wejściu z określonymi czujnikami sensorycznymi (kamerą, mikrofonem, sonarem itp.). Zadaniem każdego z modułów było przetwarzanie otrzymanych danych i kontrola zachowań określonych efektorów (kół, manipulatorów, głośników itp.), dlatego

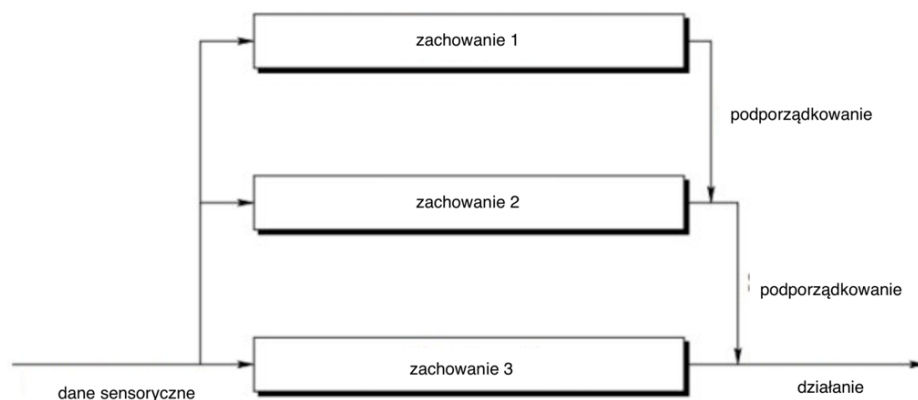
293 J.A. Meyer, *From natural to artificial life: Biomimetic mechanisms in animat designs*, „Robotics and Autonomous Systems”, t. 22, nr 1, 1997, ss. 3–21.

294 R.A. Brooks, *Intelligence without reason...*, op. cit.

295 R.A. Brooks, *Achieving Artificial Intelligence through Building Robots*, 1986.

296 R.A. Brooks, *A robust layered control system for a mobile robot...*, op. cit.

też zostały nazwane modułami behawioralnymi. Moduły zostały połączone ze sobą hierarchicznie. W tego typu architekturze jedne moduły były podporządkowane aktywności innych (architektura subsumpcyjna)²⁹⁷. Celem takiej budowy było skonstruowanie agenta, który jest kierowany przez połączenie działania modułów zgodnie ze sobą i w opozycji do siebie.



Rys. 3. Model architektury subsumpcyjnej wg Rodneya Brooksa²⁹⁸

Dzięki zastosowaniu takiej budowy możliwe było na przykład zbudowanie poruszającego się robota, którego moduł sonaru filtrował dane w poszukiwaniu nieprawidłowych odczytów i tworzył mapę przeszkód, w tym samym czasie moduł kolizji monitorował tę mapę i wysyłał sygnał zatrzymania się do modułu poruszania się, który podsumowywał wyniki w oparciu o dane pochodzące z pozostałych modułów i wywoływał pracę silnika, a tym samym ruch robota²⁹⁹. Oczywiście im bardziej skomplikowany system, tym współpracujących modułów było więcej.

Architektura subsumpcyjna łączyła informacje sensoryczne z wyborem działania w sposób prywatny i oddolny. Całościowe zachowanie było rozłożone na czynniki. Częściowe zachowania były zorganizowane w hierarchię warstwową. Każda warstwa implementowała określony poziom kompetencji behawioralnych, a wyższe poziomy były w stanie zintegrować działanie niższych poziomów i wytworzyć kompleksowe zachowanie. Działanie systemów niższego poziomu było wykorzystane przez systemy

²⁹⁷ Ibid.

²⁹⁸ Ibid.

²⁹⁹ Ibid.

wyższego poziomu. Warstwy, które otrzymały informacje z czujników, pracowały równolegle i generowały dane wyjściowe. Wyjścia te mogły być poleceniami dla silników lub sygnałami, które tłumiły lub wzmacniały działanie innych modułów. Uszkodzenie konkretnego modułu pozwalało nadal działać robotowi, w przeciwieństwie do klasycznego systemu, który mógł zawieść, nawet wtedy, jeśli zmienił się pojedynczy bajt informacji.

Roboty oddziaływały ze swoim środowiskiem, tak jak robią to organizmy, co wskazywało, że posiadają jakiś *Umwelt* w tym sensie, że były zależne od własnych czujników (zmysłów) i mogły rozpoznać swoje fizyczne środowisko, selektywnie przetwarzając bodźce środowiskowe oraz tworzyły własną wewnętrzną reprezentację swojego *Umweltu*. Propozycja Brooksa zakładała, że można stworzyć roboty oparte na innym pojęciu inteligencji, uwzględniającym działanie procesów sensorycznych. Zamiast modelować aspekty inteligencji żywych organizmów poprzez manipulację symbolami, podejście to miało na celu wywołanie interakcji w czasie rzeczywistym i wytworzenie realnych odpowiedzi w dynamicznym środowisku³⁰⁰. Wpływ badań biosemiotycznych był również widoczny, gdy Brooks stwierdził, że system nie powinien reprezentować świata za pomocą zaimplementowanego wewnętrznego zestawu symboli, a następnie działać na tym modelu³⁰¹. Zgodził się z Kennethem Craikiem, że „świat jest swoim własnym najlepszym modelem”³⁰², co oznacza, że połączenie percepcji z działaniem doprowadzi do bezpośredniej interakcji ze światem w miejsce tworzenia symbolicznego modelu tego świata. Mimo, że każdy moduł modeluje tylko pewne aspekty świata, robi to na niskim poziomie, korzystając w tym celu sygnałów sensomotorycznych. Oczywiście, moduły używają założeń dotyczących świata zakodowanych na stałe w samych algorytmach, ale unikają użycia pamięci do przewidywania zachowania świata. Zamiast tego polegają na bezpośrednim sprzężeniu sensorycznym w możliwie największym stopniu.

Brooks twierdził też, że ucieleśnienie systemu pozwala na tworzenie i testowanie zintegrowanego systemu kontroli fizycznej, a nie teoretycznych modeli lub symulowanych robotów, które mogą niewystarczająco efektywnie działać w świecie fizycznym³⁰³. Kolejnym argumentem za ucieleśnieniem była możliwość rozwiązania problemu

300 R.C. Arkin, R.C. Arkin, *Behavior-based robotics*, MIT press 1998, s. 130.

301 R.A. Brooks, *Intelligence without representation...*, *op. cit.*

302 K.J.W. Craik, *The nature of explanation...*, *op. cit.*, ss. 51–52.

303 R.A. Brooks, *Artificial life and real robots*, [w:] *Toward a practice of autonomous systems: Proc. of the 1st Europ. Conf. on Artificial Life*, 1992, s. 3.

ugruntowania symboli, poprzez bezpośrednie połączenie danych zmysłowych ze znaczącymi działaniami. Rozwijanie umiejętności percepcyjnych i ruchowych uznano za niezbędną podstawę inteligencji³⁰⁴. Wcześniej reprezentacje konstruktora/programisty stanowiły punkt wyjścia dla badań nad sztuczną inteligencją, tymczasem architektura Brooksa wskazywała, że inteligencja powstaje w interakcji ze światem. Poszczególne moduły nie były inteligentne, lecz efekt ich współdziałania można było za taki uznać³⁰⁵. Brooks budował roboty z sieci połączonych automatów, prostych modułów, które działały równolegle i wdrażały zachowania. Tego typu mechanizmy miały małe zapotrzebowania na moc obliczeniową, w przeciwieństwie do klasycznych maszyn, co sprawiało, że bardziej przypominały ekonomiczne zużycie energii przez żywe istoty.

Zainteresowanie Brooksa pracami Uexküll'a przyniosło idee ucieleśnienia i usytuowania, które okazały się wystarczające dla modelowania prostych poruszających się robotów, lecz całkowicie niewystarczające dla kształtowania wyższych zachowań inteligentnych. Powód leżał prawdopodobnie w braku głębszego zainteresowania mechanizmami powstawania coraz bardziej rozbudowanych znaczeń. Brooks początkowo twierdził, że badania powinny posuwać się zgodnie z układem łańcucha ewolucyjnego, lecz zamiast krok po kroku budować coraz bardziej wyrafinowane zwierzęce roboty, dodając nowe aspekty inteligentnych zachowań, postanowił zająć się humanoidami. Powstałe w pracowni Brooksa roboty Cog i Kismet modelowały co prawda pewne aspekty ludzkiego zachowania, jak chociażby interakcje społeczne, ale nie były rozszerzeniami wcześniejszych maszyn³⁰⁶. Brooks porzucił te badania, nie mogąc skonstruować kompletnego systemu. Należy jednak zaznaczyć, że mimo iż rozwój badań nad architekturą subsumpcyjną został przerwany, to architektura ta wciąż jest wykorzystywana przez architektów systemów poruszających się, które osiągają znacznie lepsze wyniki, niż maszyny DARPA oparte na symbolicznych strukturach³⁰⁷.

304 R. Pfeifer, J. Bongard, *How the Body Shapes the Way We Think: A New View of Intelligence*, Cambridge 2006.

305 R.A. Brooks, *Elephants don't play chess...*, *op. cit.*

306 S. Turkle i in., *Encounters with kismet and cog: Children respond to relational artifacts*, „Digital Media: Transformations in Human Communication”, t. 120, 2006, ss. 313–330.

307 R.A. Brooks, *The Bitter Lesson*, [w:] *Rodney Brooks Robots, AI, and other stuff*.

3.2.2 Badania nad komunikacją synergiczną

Biosemiotyka wskazuje, że inteligentne działanie jest nierozzerwalnie związane z przetwarzaniem informacji a także z komunikacją. Istnienie każdego gatunku zależy od przekazywania wiadomości genetycznych z pokolenia na pokolenie. Początkowo badanie komunikacji i kontroli w maszynach i organizmach żywych było rozpatrywane w świetle teorii informacji³⁰⁸. Cybernetyce brakowało jednak środków na opisanie komunikacji żywych organizmów, które są zbyt złożone. Sztuczni agenci poprzez swoją złożoność coraz bardziej przypominali żywe organizmy. Powstała więc konieczność sformułowania nowych zasad projektowania i analizy, które wykraczają poza tradycyjne podejścia, dlatego zaczęto się przyglądać możliwości integracji tej dziedziny z biosemiotyką³⁰⁹. Sztuczne systemy, by osiągnąć swoje cele, potrzebowały skutecznej komunikacji wewnętrznej (między autonomicznymi subagentami i ich przyszłymi stanami), jak również zewnętrznej komunikacji z innymi agentami³¹⁰. Komunikacja obejmuje produkcję i interpretację znaków, dlatego analiza złożonych czynników wymaga podejścia semiotycznego. Połączoną metodologię cybernetyczno-biosemiotyczną można zastosować również do zarządzania wiedzą, która obecnie rozwija się w kierunku wielości modeli bazowych i integracji danych.

Szczególnych wyników dostarczyła analiza przetwarzania informacji przez żywe organizmy na wielu poziomach. Funkcje na poziomie komórkowym obejmują wychwytywanie zasobów, wzrost, metabolizm, modyfikację cytoszkieletu i właściwości powierzchni, wykrywanie warunków zewnętrznych, cykl komórkowy i kontrolę głównych procesów wewnętrznych. Każda z tych funkcji wymaga tysięcy i milionów interakcji molekularnych. Organizmy wielokomórkowe mają jeszcze bardziej złożone zestawy funkcji związanych z różnicowaniem komórek. Każdy rodzaj komórek specjalizuje się w

308 C.E. Shannon, *A mathematical theory of communication*, „Bell System Technical Journal”, t. 27, nr 3, 1948, ss. 379–423; N. Wiener, *Cybernetics or Control and Communication in the Animal and the Machine*, MIT press 1965, t. 25.

309 Chodzi tu o cybernetykę drugiego rzędu, opisującą autoreferencyjny system, który ewoluje i różnicuje się w procesie komunikacji z samym sobą. Patrz. M.M. Kurowski, *Niklas Luhmanna radykalizacja projektu fenomenologii*, „Fenomenologia”, nr 1, 2003, ss. 63–74.

310 S. Brier, *Cybersemiotics: Why information is not enough!*, University of Toronto Press 2008, ss. 244–246.

wykonywaniu unikalnych funkcji, które są niezbędne dla całego organizmu. Funkcje całego organizmu obejmują jedzenie, trawienie, wydalanie, wykrywanie, ruch, kojarzenie i rozmnażanie. Wreszcie istnieją ekosystemy, populacje, rodziny, gatunki, których funkcje obejmują komunikację, budowę siedlisk i obronę³¹¹. Istnieje wiele różnych form synergii, które ułatwiają komplementarność funkcjonalną, zwiększenie efektywności lub udoskonalenie działania, umożliwiają wspólne uwarunkowania środowiskowe w tym podział ryzyka i kosztów, wymianę informacji, podział pracy, lecz przede wszystkim, co tutaj istotne dla tej pracy, w wyniku najróżniejszych interakcji kształtują inteligencję.

Komunikacja synergiczna zwiększa prawdopodobieństwo osiągnięcia sukcesu w działaniu i znacznie zmniejsza ryzyko porażki. Skoordynowany ruch, w tym wykorzystanie formacji w celu zwiększenia wydajności aerodynamicznej lub hydrodynamicznej, może zmniejszyć indywidualne wydatki energetyczne i ułatwić nawigację. Roboty, które wykorzystują ten typ komunikacji używają zazwyczaj tylko lokalnego wykrywania i minimalizują komunikację, aby utrzymać globalny cel³¹². W badaniach nad komunikacją Jakoba Fredslunda i Mai Matarić każdy robot podążał za wyznaczonym „przyjacielem”, poruszając się pod odpowiednim kątem i w pewnej odległości, używając czujnika³¹³. Główną ideą było, by algorytm utrzymywał przyjaciela w polu widzenia. Prosty schemat komunikacji minimalnej, który opierał się na monitorowaniu pulsowania innych robotów, obserwowanie zmian w ich ruchu i wysyłaniu jednobajtowych informacji w celu formowania kolumny i unikania przeszkód, okazał się być wyjątkowo skuteczny. Zademonstrowano, że nie trzeba tworzyć skomplikowanego algorytmu pozycjonowania. Globalna mapa i globalna wiedza o pozycjach wszystkich robotów umożliwiły niezawodne działanie, podczas którego grupa robotów wykazywała globalne, skoordynowane zachowanie.

Ustalenia biosemiotyki przeniesione z żywych organizmów na sztucznych agentów dają w dziedzinie komunikacji świetne efekty. Mimo, że badania nad inteligencją roju wciąż trwają, należałoby dodatkowo uwzględnić w nich kontekst hierarchii agentów,

311 Por. A. Sharov, T. Maran, M. Tønnessen, *Comprehending the Semiosis of Evolution*, „Biosemiotics”, t. 9, nr 1, 2016, ss. 1–6; J.S. Turner, *Semiotics of a superorganism*, „Biosemiotics”, t. 9, nr 1, 2016, ss. 85–102.

312 M.J. Mataric, R.A. Brooks, *Learning a distributed map representation based on navigation behaviors...*, *op. cit.*

313 J. Fredslund, M.J. Mataric, *A general algorithm for robot formations using local sensing and minimal communication...*, *op. cit.*

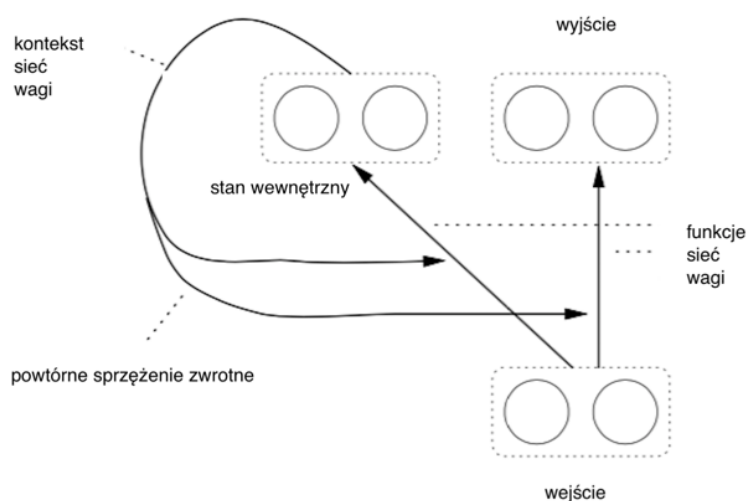
ponieważ różnice w poziomach przetwarzania informacji mogą być dziedziczone i wpływać na wzrost inteligencji stada bądź różnicowania się zachowań. Ponadto agenci o różnej hierarchii mogliby komunikować się i wchodzić w interakcję na różne sposoby, lecz odmienne zachowania budowałyby podmiot zbiorowy o wielu cechach i umiejętnościach. Agenci, którzy byliby ze sobą połączeni hierarchicznie lub genealogicznie, mogliby produkować zachowanie o dużo większej złożoności. Funkcje agentów, które byłyby kodowane przez zestaw znaków (jak znaki pamięciowe, informacje zewnętrzne i naturalne znaki) w pewnym stopniu indukowałyby semiozę. W przeciwieństwie do klasycznego programowania systemów z naciskiem na kontrolę, pętle zwrotne i atraktory, takie podejście pozwala przyjrzeć się ewolucji, funkcjonalności i komunikacji agentów.

3.2.3 Działanie sieci neuronowych w perspektywie biosemiotycznej

Zadaniem sztucznej sieci neuronowej (SSN) jest zapewnić samodzielne tzn. autonomiczne oddziaływanie na otoczenie. Większość sztucznych systemów używających SSN wciąż pozostaje heteronomiczna, ponieważ ich mechanizm sterowania pozostaje zależny od projektanta. Niezbędnym elementem inteligencji jest zdolność samodzielnego określenia i dostosowania mechanizmów leżących u podstaw zachowania sztucznego systemu. Stąd też wynika potrzeba samoorganizacji. Zazwyczaj to podejście opiera się na wykorzystaniu technik uczenia, aby umożliwić agentom dostosowanie wewnętrznych parametrów ich mechanizmów kontrolnych. Z punktu widzenia kognitywistyki potrafią to robić roboty adaptacyjne kontrolowane przez SNN dzięki równoległemu i rozproszonemu charakterowi reprezentacji wag i jednostek w SNN. Reprezentacje mogą być tworzone w interakcji ze środowiskiem agenta, który nie tworzy z góry określonych modeli zewnętrznej rzeczywistości. Mimo tego zazwyczaj środowisko jest wciąż ograniczone do wartości wejściowych i wyjściowych. Oznacza to, że SSN, w przeciwieństwie do prawdziwego układu nerwowego, nie są osadzone w kontekście agenta i jego świata zewnętrznego.

Należy więc wykorzystać sztuczną sieć neuronową jako sterownik robota, dzięki czemu będzie mogła kontrolować jego sztuczny układ przetwarzania informacji od wejść sygnałów sensorycznych aż do wyjść silników. Wtedy za pomocą sensomotorycznej architektury agenta sieć może faktycznie wchodzić w interakcję z obiektami fizycznymi w

jego otoczeniu niezależnie od mediatora lub interpretacji obserwatora. Już w 1997 roku Georg Dorffner zaproponował podejście, które nazwał radykalnym koneksjonizmem, które miało na celu wykorzystanie SSN do kontroli i uczenia się agenta w naturalnym środowisku ³¹⁴. Najważniejszą rolę w perspektywie podejścia do badania funkcji poznawczych odgrywają rekurencyjne sieci neuronowe (RSN). Ich działanie opiera się na funkcji mapowania przepływu sygnału od wejścia do wyjścia, a także jej stanu wewnętrznego oraz na działaniu sieci kontekstowej odwzorowującej stan wewnętrzny na wagi sieci dla funkcji następnego kroku. Zatem ten typ sieci wykorzystuje sprzężenie zwrotne drugiego rzędu, co powoduje, że mapowanie wejścia-wyjścia sieci może zmieniać się z kroku krok po kroku w sposób zależny od kontekstu lub stanu wewnętrznego agenta.



Rys.4. Działanie rekurencyjnej sieci neuronowej

Sieci te stanowią podstawę dla długoterminowych reprezentacji doświadczeń w zakresie wag połączeń oraz krótkoterminowych reprezentacji bieżącego kontekstu lub bezpośredniej przeszłości agenta dzięki istniejącej wewnętrznej informacji zwrotnej³¹⁵. Podobnie twierdził Maturana uznając je za zdeterminowaną strukturę podobną do prawdziwego układu nerwowego³¹⁶. Oznacza to, że ich reakcje na bodźce środowiska zawsze zależą od aktualnego stanu lub struktury systemu, nie są zatem nigdy określone

314 G. Dorffner, *Radical connectionism-a neural bottom-up approach to AI*, „Neural Networks and a New Artificial Intelligence”, 1997, ss. 93–132.

315 *Ibid.*

316 H.R. Maturana, *Autopoiesis and Cognition...*, *op. cit.*

przez same dane wejściowe. Powstaje tu więc pewnego rodzaju autonomia systemu reprezentacyjnego. W nawiązaniu do teorii biosemiotycznej należy zatem uznać, że *Umwelt* takiego agenta jest ontologicznie odmienny od sposobu istnienia jego fizycznego środowiska. RSN funkcjonuje podobnie jak dynamiczna sieć biofizyczna aktywnego podmiotu poznającego, którego nie można opisać z czysto zewnętrznego punktu widzenia³¹⁷.

W szkoleniu sieci neuronowych stosuje się wiele technik i algorytmów. Różnią się one przede wszystkim stopniem sprzężenia zwrotnego oraz poziomem samoorganizacji. Uczenie nadzorowane polega na wymuszeniu oczekiwanych odpowiedzi na konkretne dane wyjściowych i kontrolowaniu każdego kroku. Sieć jest instruowana o tym, których wejść użyć i jakie sygnały wyjściowe wyprodukować. Zazwyczaj jednak techniki nadzorowane nie są wykorzystywane w systemach robotycznych, aby umożliwić maksymalną autonomię robota i ograniczyć interwencję projektanta do minimum. Możemy to prześledzić na przykładzie nauki jazdy na rowerze: instruktor mówi jak w danym momencie poruszać nogami, rękami, w jakiej pozycji utrzymywać ciało. Takie instrukcje niewiele wnoszą, ponieważ zarówno uczący się jazdy jak i agent musi sam skonstruować sobie rozwiązanie tego problemu przez dostosowywanie aktywności sensomotorycznej w interakcji z obiektem działania i swoim otoczeniem. Dlatego sztuczne systemy powinny być przede wszystkim szkolone za pomocą technik wzmacniających, ewolucyjnych lub adaptacyjnych. Sztuczny agent powinien sam zorganizować swoje wewnętrzne struktury i parametry, aby dopasować się do pewnych ograniczeń środowiskowych. Można to porównać do konstruktywistycznego poglądu na organizmy dążące do koncepcyjnego zrównoważenia poprzez dostosowanie ich struktur wewnętrznych i do nabytego uprzednio doświadczenia³¹⁸. Uczenie nienadzorowane, którego przykładem sieć Donalda Hebba, połączenie jest wzmacniane niezależnie od błędu, ale jest wyłącznie funkcją zbieżności między potencjałami działania neuronów. Metoda Hebba powstała na podstawie badania funkcji poznawczych, takich jak rozpoznawanie wzorców i uczenie się przez doświadczenie. Sieć taka grupuje dane, które nie zostały oznakowane, sklasyfikowane lub skategoryzowane, analizując podobieństwa i reaguje na ich podstawie lub ich braku na

317 Por. T. von Uexküll, *Introduction: Meaning and science in Jakob von Uexküll's concept of biology*, „Semiotica”, t. 42, nr 1, 1982, ss. 1–24.

318 S. Harter, *The construction of the self: A developmental perspective.*, op. cit.

nowe dane³¹⁹. Podczas nienadzorowanego uczenia się przez wzmacnianie agent otrzymuje tylko sporadyczne informacje zwrotne, zazwyczaj w kategoriach pozytywnych lub negatywnych. Dzięki temu może dostosować swoje zachowanie do środowiska w taki sposób, aby zmaksymalizować pozytywne wzmocnienie i zminimalizować negatywne wzmocnienie.

Koordinacja przepływu sygnału sieci między wejściem a wyjściem zależy od samoorganizacji sieci. Stąd wewnętrzne reprezentacje (wagi i aktywacje ukrytych jednostek) można uznać za znaki, które są prywatne dla sieci i często nieprzejrzyste dla zewnętrznych obserwatorów. Zatem, w przeciwieństwie do tradycyjnej sztucznej inteligencji, koneksjonizm nie promuje symbolicznych reprezentacji, które odzwierciedlają wcześniej określoną zewnętrzną rzeczywistość. Podkreślają one samoorganizację adaptacyjnego przepływu sygnałów między jednostkami, co jest zgodne z interaktywnym poglądem na reprezentację³²⁰. Koneksjonizm oferuje zatem alternatywne podejście do badania reprezentacji poznawczej i używania znaków. W szczególności równoległy i rozproszony charakter reprezentacji wagi i jednostek w SSN oraz możliwość tworzenia reprezentacji z doświadczenia sprawiają, że koneksjonizm jest w dużej mierze zgodny z konstruktywistycznym poglądem na wymuszoną adaptację napędzaną przez potrzebę dopasowania się do ograniczeń środowiskowych³²¹.

RSN pozwalają na użycie klasycznych rozwiązań sztucznej inteligencji, pozbawionych podstaw i interakcji ze światem, który mają reprezentować, oraz na niezależność reaktywnych systemów, których działania od tego świata zależą. Wewnętrzne struktury i reprezentacje utworzone w RSN są wynikiem strukturalnego sprzężenia między nimi, co zapewnia ich strukturalną zgodność. Zachodzi więc tu konstruktywny proces, odzwierciedlający subiektywne osadzenie agenta w świecie, który pozwalając mu przypisywać różne znaczenia bodźcom zgodnie z jego własnym aktualnym zachowaniem³²².

319 G.E. Hinton, T.J. Sejnowski (red.), *Unsupervised learning: foundations of neural computation*, Cambridge, Mass 1999.

320 M.H. Bickhard, *The interactivist model*, „Synthese”, t. 166, nr 3, 2009, ss. 547–591.

321 G. Dorffner, *Radical connectionism-a neural bottom-up approach to AI...*, *op. cit.*

322 Por. F. Maturana, W. Shen, D.H. Norrie, *MetaMorph: an adaptive agent-based architecture for intelligent manufacturing...*, *op. cit.*; T. Ziemke, N.E. Sharkey, *A stroll through the worlds of robots and animals: Applying Jakob von Uexküll's theory of meaning to adaptive robots and artificial life...*, *op. cit.*

Radykalny koneksjonizm ma wiele wspólnego z perspektywą Brooksa przyjętą w tworzeniu struktury subsumpcyjnej³²³. Architektura ta tworzy sieć adaptacyjną i oferuje wiele możliwości reprezentacyjnych. Połączenie jej z działaniem RSN pozwala sztucznemu agentowi skonstruować własne odniesienie do świata. Adaptacyjna neurorobotyka proponuje interesujące podejście do badania interaktywnej inteligencji. Sieć koneksjonistyczna może przy pomocy ciała robota (sensorów i efektorów) oddziaływać z obiektami fizycznymi i stać się integralną częścią agenta, który w perspektywie Brooksa staje się ucieleśniony i usytuowany. Reprezentacje wewnętrzne takiego systemu są tworzone w wyniku interakcji fizycznej ze światem, który reprezentują lub odzwierciedlają. Sieć sterownika robota pełni rolę pełnego kręgu funkcjonalnego i może do pewnego stopnia skonstruować własny *Umwelt*. Jako przykład może posłużyć samojezdny robot w pomieszczeniu wypełnionym różnymi fizycznymi obiektami. Robot wyposażony w czujniki/receptory jest wrażliwy na percepcyjne wskazówki jego otoczeniu. Unika przeszkód, które są częścią jego świata percepcyjnego, natomiast obrazy na ścianie lub widok z okna nie są częścią jego świata percepcyjnego i nie mają dla niego znaczenia. W ten sposób robota można uznać za osadzonego w swoim własnym *Umwelcie*, składającym się z jego świata percepcyjnego (*Merkweltu*) zawierającego obiekty lub ich brak i świata operacyjnego (*Wirkweltu*) powstającego w wyniku poruszania się. W *Merkwelcie* odbywa się transformacja znaku i powstawanie reprezentacji. Wewnętrzny świat jest tu samoorganizującym się działaniem prywatnych znaków w interakcji agent-środowisko. Tak więc adaptacje w robocie mogą być postrzegane jako konstrukcja interaktywnych reprezentacji poprzez tworzenie, adaptację i optymalizację kręgów funkcjonalnych w interakcji ze środowiskiem.

3.2.4 Metody ewolucyjne i robotyka adaptacyjna

Metody ewolucyjne tworzone są przez analogię do mechanizmów ewolucji biologicznej Darwina. Wykorzystanie technik ewolucyjnych jest podejściem, które ma na celu umożliwienie robotom uczenia się w interakcji z ich otoczeniem przy minimalnej

323 G. Dorffner, *Radical connectionism-a neural bottom-up approach to AI...*, op. cit.

interwencji człowieka³²⁴. Informacja zwrotna nie wyznacza wprost działania, lecz je wartościuje. W algorytmie genetycznym istnieje pewna ilość rozwiązań problemu (populacja tzw. osobników lub fenotypów), która ewoluuje w kierunku lepszych rozwiązań. Każde kandydujące rozwiązanie ma zestaw właściwości (chromosomy lub genotyp), które mogą być rekombinowane i mutowane³²⁵. Ewolucja rozpoczyna się od procesu iteracji losowo generowanych osobników (pokolenia). W każdym pokoleniu ocenia się przystosowanie (*fitness*) każdego osobnika w środowisku. Przystosowanie jest zazwyczaj wartością kosztu w rozwiązywanym problemie optymalizacji³²⁶. Lepiej dopasowane osobniki są wybierane z populacji i tworzą nowe pokolenie, po czym proces zaczyna się od początku. Algorytm kończy się, gdy zostanie wyprodukowana maksymalna liczba pokoleń lub osiągnięty zostanie zadowalający poziom przystosowania populacji, czyli znajdzie się optymalne rozwiązanie problemu. Ze względu na presję selekcyjną przeciętne przystosowanie populacji prawdopodobnie wzrośnie.

Konkretnym przykładem badań nad robotami ewolucyjnymi są prace Dario Floreano, które pokazały, jak wizualne umiejętności sieci neuronowej robotów zostały wzmocnione przez algorytmy genetyczne. Receptory wizyjne rozwinęły wrażliwość na obrazy napotkane w środowisku dużo bardziej niż receptory biernie eksponowane na ten sam zestaw obrazów. Roboty z ewoluującym systemem wizyjnym rozwinęły wrażliwość na drobniejsze aspekty środowiska i aktywnie rozwijały cechy utrzymujące stabilne zachowanie w przeciwieństwie do innych robotów, które nie były w stanie prawidłowo się poruszać. Późniejsza transplantacja neuronów ewoluujących robotów do grupy kontrolnej pokazała, że nawet wtedy nie mogły poprawnie ocenić bodźców wizualnych³²⁷. Wskazuje to na istotę współdziałania rozwoju ewolucyjnego i środowiska. Tego typu eksperymenty pokazują, że sieci neuronowe z dynamiką sprzężenia zwrotnego ewoluowały do rozróżnienia odpowiednich bodźców środowiskowych. Zatem w tych eksperymentach zarówno wewnętrzny przepływ sygnałów, wykorzystanie sprzężenia zwrotnego, a także

324 S. Nolfi i in., *Evolutionary robotics*, [w:] *Springer Handbook of Robotics*, Springer 2016, ss. 2035–2068.

325 D. Whitley, *Next Generation Genetic Algorithms: A User's Guide and Tutorial*, [w:] *Handbook of Metaheuristics*, M. Gendreau, J.-Y. Potvin (red.), Cham 2019, ss. 245–274.

326 Oznacza to, że odwzorowuje wartość zmiennych na rzeczywistą liczbę, reprezentując koszt/stratę związane ze zdarzeniem.

327 D. Floreano, P. Husbands, S. Nolfi, *Evolutionary Robotics*, [w:] *Springer Handbook of Robotics*, B. Siciliano, O. Khatib (red.), Berlin, Heidelberg 2008, ss. 1423–1451.

użycie zewnętrznego znaku, czyli bodźców zewnętrznych można interpretować jako znaki konstruowane w sztucznym procesie ewolucyjnym.

Procesy ewoluujących znaków są często trudne do szczegółowej analizy i zrozumienia, ponieważ są one prywatne dla robota i w wielu przypadkach radykalnie różnią się od rozwiązań zaprojektowanych przez ludzi w oparciu o ich własną dalszą perspektywę. Wydaje się zatem, że procesy ewolucyjne często znajdują sposoby na spełnienie kryteriów przystosowania, które są sprzeczne z intuicjami twórcy systemu, co do tego, w jaki sposób problem powinien być rozwiązany. Neuroadaptacyjna robotyka umożliwia sztucznym agentom tworzenie własnych mechanizmów sensomotorycznych i reprezentacji poprzez samoorganizację w interakcji z otoczeniem. Budowa fizyczna i początkowa konfiguracja tego typu robotów pozostają produktem człowieka, lecz nie zostaje tam wbudowana wiedza na temat transformacji sensomotorycznej. Wiedza ta jest tworzona przez doświadczenie a adaptacja jest wynikiem dążenia do dopasowania.

Użycie sztucznych systemów nerwowych w połączeniu z technikami uczenia się obliczeniowego pozwala autonomicznym agentom dostosować się do ich środowiska. Co więcej, przy użyciu wewnętrznego sprzężenia zwrotnego, behawioralne znaczenie różnych czynników może się zmieniać w czasie. Roboty takie dostosowują się do swojego środowiska, zyskując to, co Uexküll określił jako historyczna podstawa reakcji (*Reaktionsbasis*)³²⁸. Samoorganizujące się roboty posiadają więc pewne subiektywne znaczenie w tym sensie, że ich reakcje nie są określone przez odgórne reguły, lecz są specyficzne dla nich, ich doświadczenia i samoorganizacji.

Sztuczni agenci są także zaangażowani w procesy tworzenia znaczenia i wykorzystują znaki. W przeciwieństwie do typowych programów komputerowych, procesy znaczenia wśród sztucznych agentów nie są całkowicie określone przez ich ludzkich projektantów i niezależne od interpretacji przez zewnętrznych obserwatorów. Dodatkowo, w wielu przypadkach, interpretacja człowieka przy bliższym przyjrzeniu się procesom wewnętrznym jest niemożliwa. Znaczna część znaków używanych przez te systemy jest zatem, z powodu ich samoorganizacji, rzeczywiście prywatna i specyficzna dla nich. Można zatem stwierdzić, że tego typu sztuczni agenci mają pewien stopień

328 J. von Uexküll, *Theoretische biologie...*, op. cit., s. 180.

autonomii epistemicznej³²⁹, czyli tak jak żywe organizmy, same wchodzą w interakcje z otoczeniem.

Konkretnym przykładem badań nad robotami ewolucyjnymi są prace Husbandsa, który opracował powtarzalne kontrolery do zadania rozróżniania celu. Zadaniem robota wyposażonego w kamerę było zbliżenie się do trójkąta z białego papieru zamontowanego na ścianie i unikanie prostokątów. W tych eksperymentach zarówno topologia sieci i lub pole odbiorcze (piksele obrazu kamery sieci) podlegały procesowi ewolucyjnemu. Analiza przebiegów eksperymentalnych wykazała, że strukturalnie proste sieci sterujące ze złożoną wewnętrzną dynamiką sprzężenia zwrotnego ewoluowały, wykorzystując rozróżnienie między bodźcami środowiskowymi. Zarówno wewnętrzny przepływ sygnałów, wykorzystanie sprzężenia zwrotnego a także wybór wśród bodźców pojawiły się w sztucznym procesie ewolucyjnym.

Wpływ ludzkiego projektanta można jeszcze bardziej zredukować, stosując metody koewolucyjne. Nolfi i Floreano, ukształtowali ewolucyjnie dwa roboty kontrolowane przez sieci rekurencyjne, aby uzyskać zachowanie drapieżników i ofiar³³⁰. Drapieżnik był wyposażony w prosty aparat, który pozwalał mu obserwować zdobycz z dystansu, a jego zadaniem było dotknięcie ofiary. Ofiara była tylko wyposażona w czujniki podczerwieni krótkiego zasięgu, ale także umiejętność szybszego poruszania się niż drapieżnik. Obydwa roboty rozwinęły złożone strategie pościgu i ucieczki. Drapieżnik rozwinął sprawność, która pozwoliła mu obserwować i interpretować zachowanie ofiary jako objaw jej przyszłego zachowania w najbliższej przyszłości, co pozwalało mu atakować wtedy, gdy miał realną szansę dotknąć ofiarę. Ofiara natomiast pozostawała nieruchoma do momentu wykrycia drapieżnika, po czym dopiero uciekała.

Roboty wykorzystywane w tego typu eksperymentach są produktem człowieka, jeśli chodzi o budowę fizyczną, konfiguracje eksperymentalne i inne z góry określone aspekty, lecz żadna odgórna wiedza na temat zachowania nie została w nie wbudowana. Zdobywanie wiedzy w dużym stopniu pozostawia się samemu robotowi. Podejście takie jest w dużej mierze zgodne z konstruktywistycznym poglądem na wiedzę jako aktywnie

329 M.H. Bickhard, *Robots and representations*, [w:] *From animals to animats 5-Proceedings of the Fifth International Conference on Simulation of Adaptive Behavior*, Cambridge 1998, ss. 58–63.

330 S. Nolfi, D. Floreano, *Coevolving Predator and Prey Robots: Do "Arms Races" Arise in Artificial Evolution?*, „Artificial Life”, t. 4, nr 4, 1998, ss. 311–335.

budowaną przez poznawanie obiektów i użycie systemu sensomotorycznego w interakcji z otoczeniem. Adaptacja jako dążenie do dopasowania lub przeżycia wykorzystuje informację zwrotną we wzmacnianiu i uczeniu się.

3.2.4.1 Koncepcja adaptacji środowiska przez roboty

Robotykę ewolucyjną znacząco wpłynęła na zmiany w dziedzinie SI, łącząc ich oprogramowanie, sztuczny system nerwowy, fizyczne ciało ze środowiskiem³³¹. Wielu kognitywistów i robotyków podkreślało, że ścisła interakcja mózgu, ciała i środowiska jest kluczowa dla pojawienia się inteligentnych, adaptacyjnych zachowań i procesów poznawczych³³². W konsekwencji pojęcia kognitywistyki, takie jak ucieleśnienie i usytuowanie, spotkały się z dużym zainteresowaniem i były badane eksperymentalnie. Pomimo wiedzy na temat znaczenia interakcji agent-środowisko, robotycy poświęcili mało uwagi pracom nad poznaniem rozproszonym. Zamiast tego robotyka ewolucyjna czerpała wiele inspiracji z SI opartej na projektowaniu oddolnym i modułowym. Mimo, że to podejście przyniosło wiele sukcesów, zignorowano rolę środowiska jako zasobów obliczeniowych i reprezentacyjnych.

Zdolność do modyfikowania środowiska dla własnych celów jest bardziej charakterystyczna dla ludzi niż dla innych gatunków zwierząt. Jednak wszystkie gatunki zmieniają środowisko w sposób adaptacyjny: ptaki budują gniazda, bobry stawiają tamy zaporowe, szympansy zostawiają kamienie w miejscach, gdzie łupią orzechy, itd. Jak zauważył David Kirsh, każdy gatunek wykazuje trzy sposoby na zmaksymalizowanie swoich działań w otoczeniu: przystosowanie się do środowiska, przeniesienie się w inne miejsce lub dostosowanie otoczenia do własnych potrzeb³³³. W badaniach nad robotami ewolucyjnymi tylko pierwsza z tych opcji została zbadana. Mimo, że istnieje wiele badań

331 A. Clark, *Being there: Putting brain, body, and world together again...*, *op. cit.*

332 A. Clark, *Reasons, robots and the extended mind*, „Mind & Language”, t. 16, nr 2, 2001, ss. 121–145; D. Floreano, P. Husbands, S. Nolfi, *Evolutionary Robotics...*, *op. cit.*; J. Fredslund, M.J. Mataric, *A general algorithm for robot formations using local sensing and minimal communication...*, *op. cit.*; G. Pezzulo i in., *The Mechanics of Embodiment...*, *op. cit.*; T. Ziemke, *On the role of emotion in biological and robotic autonomy...*, *op. cit.*

333 D. Kirsh, *Adapting the Environment Instead of Oneself*, „Adaptive Behavior”, t. 4, nr 3–4, 1996, ss. 415–452.

na temat agentów / robotów dostosowujących się do danego środowiska, aktywna adaptacja lub modyfikacja środowiska została zbadana stosunkowo słabo.

Ewolucyjne badania robotyki agentów lub gatunków, które aktywnie modyfikują swoje środowisko materialne dla własnych celów, a nie tylko inteligentnie wykorzystują dane środowisko, są rzadkie. Istnieją oczywiście eksperymenty z robotami, które ewoluują, aby zmienić ich otoczenie, np.: robot Khepera zbierający śmieci³³⁴. Jednak charakter zadania robota się nie zmienia: podnosi przedmioty i pozbywa się ich, ale nigdy nie wykorzystuje ich w celu ułatwienia sobie późniejszej pracy. Dlatego interesujące jest zbadanie, w jaki sposób gatunki w trakcie swojego życia mogą modyfikować swoje środowisko dla korzyści poznawczych własnych lub innych agentów, aby zmniejszyć wymagania dotyczące na przykład pamięci lub czasu niezbędnego dla wykonania zadania.

Adaptacja środowiska jest ściśle związana z koncepcją niszy w biosemiotyce, co oznacza zdolność organizmów do konstruowania, modyfikowania i wybierania ważnych składników ich lokalnych środowisk. Doświadczalny *Umwelt* jest zakorzeniony w materialnym i semiotycznym ciele organizmu, które znajduje się w określonej części siedliska – niszy. Każdy gatunek ma ograniczony dostęp do całości świata, ponieważ zdolność każdego gatunku do wykrywania i interpretowania potencjalnych sygnałów w jego otoczeniu, ewoluowała tak, aby pasowała do określonej niszy ekologicznej. Forma postaci pojedynczego organizmu w swojej niszy ekologicznej, zakłada działanie ciała fizycznego jako systemu termodynamicznego w stanie nierównowagi. Interakcja podmiotu ze środowiskiem jest mediowana przez właściwości niszy: gdzie być, co jeść i jak działać, aby osiągnąć sukces³³⁵. W tym celu nie tylko przystosowanie się do środowiska, ale też ulepszanie go dla własnych celów jest niezbędne dla powstania niezbędnych warunków zapewniających przetrwanie.

Badania te nadal znajdują się na najwcześniejszych etapach. Tymczasem pytania takie jak: jaka jest minimalna pojemność pamięci roboczej potrzebnej do wykonania zadania lub jaka jest minimalna ilość obliczeń wymaganych do podjęcia decyzji, wciąż domagają się odpowiedzi. Wiele z tych informacji może być zakodowanych w środowisku. Perspektywa robotyki ewolucyjnej pokazuje, że podczas gdy szereg badań dotyczy roli koordynacji czuciowo-ruchowej, współdziałania morfologicznej budowy i wewnętrznych

334 D. Floreano, P. Husbands, S. Nolfi, *Evolutionary Robotics...*, *op. cit.*

335 J.N. Hoffmeyer, *Biosemiotics: Towards a new synthesis in biology*, „European Journal for Semiotic Studies”, t. 9, nr 2, 1997, ss. 355–376.

obliczeń, sposób w jaki agenci (lub też gatunki i populacje) modyfikują swoje środowisko materialne, aby pasowało do własnych celów i zdolności poznawczych pozostaje ignorowane³³⁶. Robotyka ewolucyjna, ze względu na minimalistyczne, przejrzyste i stosunkowo bezstronne modelowanie poprzez samoorganizację powinna wnieść znaczący wkład w adaptację środowiska w systemach naturalnych i sztucznych, ale też w zrozumienie działania rusztowań poznawczych.

3.2.5 Koncepcja działania filtra semiotycznego

Wszystkie cechy środowiska postrzegane przez agenta będą miały inne znaczenie w różnych momentach i w różnych kontekstach, dlatego należy wymodelować postrzeganie przez agenta jego otoczenia (*Umweltu*), tak aby pogląd agenta na świat wpłynął zarówno na jego działania, jak i na zmianę stanu wewnętrznego. Naturalne żywe systemy są wyposażone w informacje genetyczne, które prowadzą system przez różne cykle rozwojowe i różne konteksty środowiskowe. Ta ucieleśniona informacja pozwala dojrzałemu systemowi zidentyfikować i przypisać znaczenie poszczególnym wskazówkom środowiskowym, wykazując odpowiednie zachowanie jako odpowiedź i zaspokajając w ten sposób wymagania własnego *Innenweltu*.

W zakresie badań sztucznej inteligencji należy się skoncentrować na procesach przebiegających od dołu do góry. Daje to nadzieję na powstanie autonomicznych zachowań i poprawnych interakcji systemu z otoczeniem. Znaczącą rolę odgrywają podstawowe działania, które przebiegają równolegle tworząc kompleksowe zachowania. Podejście oddolne tzw. *bottom up*, w przeciwieństwie do opisywanego wcześniej klasycznego podejścia koncentruje się na rozwoju podstawowych elementów systemu. Przetwarzanie informacji zachodzi na skutek otrzymywania danych z otoczenia, przetwarzanych przez poszczególne części systemu, po czym elementy są łączone w całość, by ostatecznie wytworzyć zachowanie.

336 S. Chandrasekharan, N.J. Nersessian, *Building Cognition: The Construction of Computational Representations for Scientific Discovery*, „Cognitive Science”, t. 39, nr 8, 2015, ss. 1727–1763.

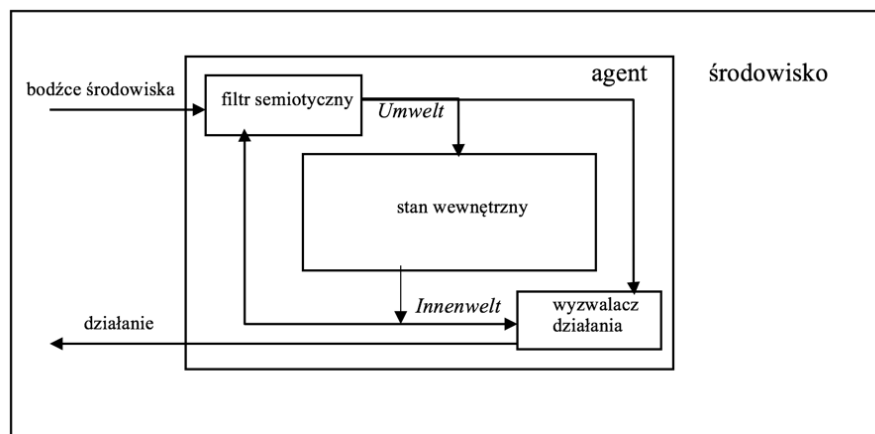
Maria Ferreira i Miguel Caldas proponują model tworzenia systemu poznawczego uwzględniający działanie tzw. filtra semiotycznego³³⁷. Zadaniem filtra jest wychwytywanie charakterystycznych cech środowiska, które mogą zmienić stan wewnętrzny agenta (*Innenwelt*), ale w równej mierze zależne są od aktualnego stanu. Percepcja jest tu całkowicie sprzęgnięta z działaniem. Podczas implementacji tego modelu znaczna część przetwarzania wstępnego zostanie przeprowadzona w przekształcaniu surowych danych z czujników sensorycznych robota w dane, które są znaczące dla systemu.

Informacje otrzymywane ze środowiska i wyszczególniane przez filtr semiotyczny stają się instrukcjami aktywności. Podejmowana przez agenta aktywność zapewnia duże prawdopodobieństwa skutecznego wykonania. Działania te mają wymierny wpływ na środowisko, które agent może zaobserwować w następnym kroku, umożliwiając powstanie sprzężenia zwrotnego i naukę bez nadzoru. Agent dzięki filtrowi sam interpretuje znaki i produkuje zachowanie. Dzięki temu znaczenie jest wytwarzane przez system, a sam agent uwalnia się tym samym od semantyki architekta.

Filtr semiotyczny jest zależny zarówno od stanu wewnętrznego (*Innenweltu*) agenta i aktualnego kontekstu środowiska. System interpretuje dane zależnie od aktualnego stanu i potrzeb. Przejście z jednego stanu wewnętrznego do innego jest modelowane w zależności od poprzedniego stanu i przy uwzględnieniu aktualnych cech *Umweltu* oraz wewnętrznych informacji o samym agencie. W przypadku mobilnego robota może to być np. przenoszenie informacji o poziomie energii w akumulatorze. W sytuacji niskiego stanu energii priorytetem staje się podłączenie do źródła zasilania i naładowanie akumulatora. Działa to w ten sam sposób w przypadku żywej istoty, która odczuwa głód i w wyniku tego przyjmuje pożywienie. Jeśli poziom energii jest wysoki, robot wykonuje zadania, nie zważając na znajdującą się stację ładowania. Powyższe działania, które zawierają w sobie cechy *Umweltu* i *Innenweltu* będą miały wpływ na proces podejmowania decyzji, który wywoła odpowiednie działania³³⁸.

337 M.I.A. Ferreira, M.G. Caldas, *Modelling Artificial Cognition in Biosemiotic Terms*, „Biosemiotics”, t. 6, nr 2, 2012, ss. 245–252.

338 M.I.A. Ferreira, *On Meaning...*, *op. cit.*



Rys.5. Działanie modelu poznawczego zawierającego filtr semiotyczny³³⁹

Prawdopodobnie nie wszystkie cechy środowiska mogą zostać zostaną dostrzeżone przez agenta a także w różnym czasie i w różnych kontekstach inne aspekty będą miały różne znaczenie. Autorzy proponują by modelować postrzeganie *Umwelt* przez agenta poprzez wybór cech, które w wyniku zastosowania filtra semiotycznego, którego działania będą zależały od reprezentowanego stanu wewnętrznego agenta (*Innenweltu*). Dzięki temu, agent uzyska specyficzną perspektywę swojego *Umweltu*, które to będzie wpływać na jego następne działania, jak i na zmianę stanu wewnętrznego. Nowy stan wewnętrzny z kolei wpłynie zarówno na działania agenta, jak i na jego filtr semiotyczny, a przez to na jego postrzeganie środowiska. *Umwelt* i *Innenwelt* pozostaną zatem w związku dialektycznym i oba będą miały znaczenie w określaniu działań agenta.

Tak zdefiniowany model ten pozwala na sformułowanie problemu, unikając w ten sposób konieczności rozwiązywania równoczesnych zadań, których system nie potrafi hierarchizować. Proces filtrowania sygnałów pozwala agentowi zidentyfikować i przypisać znaczenie nie tylko poszczególnym cechom środowiska, ale także zestawom cech, np. dla jednej danej funkcji środowiska można przypisać różne znaczenia w zależności od kontekstu, obecności innych specyficznych cech lub stanu wewnętrznego.

Proces poznania jest ciągłym procesem uczenia się i dojrzewania oraz zdobywania doświadczenia, w trakcie którego podmiot nieustannie określa i redefiniuje swój *Umwelt* świata odpowiednio dostosowując swoje zachowanie. Podstawowe pojęcia semiozy i *Umwelt*, *Innenwelt*, stanowią rdzeń modelowania matematycznego zaproponowanego

339 M.I.A. Ferreira, M.G. Caldas, *Modelling Artificial Cognition in Biosemiotic Terms...*, op. cit.

przez Ferreirę i Caldas. Model ma na celu uchwycenie dynamiki związanej z semiozą nieodłączną dla wszystkich form poznania. Umieszczając fundamentalny proces interpretacji w centrum modelu, podkreśla istotną rolę, jaką odgrywa proces przypisywania znaczeń w procesie poznawczym, a tym samym w wytwarzaniu inteligentnego zachowania.

3.2.6 Propozycja rozwiązania problemu intencjonalności agentów

Ostatnie osiągnięcia w inżynierii oprogramowania i sztucznej inteligencji, kładą coraz większy nacisk na koncepcje autonomicznych i intencjonalnych agentów jako cechę charakterystyczną dla inteligentnych sztucznych systemów. Sposoby programowania zmieniły się z modeli proceduralnych i funkcjonalnych do modeli zorientowanych obiektowo, co jest przyczyną rozwoju podejścia opartego na działaniu agentów. Potrzeba zbadania intencjonalności sztucznych systemów wynika z potrzeby rozwoju takich systemów informatycznych, jak chociażby: układy robotyczne, interfejsy użytkownika (np. boty pomocnicze), dynamiczne układy sztucznego życia, symulacje systemów naturalnych, takich jak żywe organizmy, systemy ekologiczne, gospodarcze lub społeczeństwa.

Użycie terminu agent spotyka się w dużej liczbie podejść i aplikacji. Prawdopodobnie to roboty stały się paradygmatycznymi przykładami agentów, które wykazywały zachowania autonomiczne w miejsce zaprogramowanej aktywności³⁴⁰. Wykorzystanie terminu agent przeniosło się też na inne systemy informatyczne, gdzie używane jest na określenie niezależnie działających systemów, często mających właściwości antropomorficzne. Szczególną rolę odgrywają jednak agenci w badaniach sztucznej inteligencji i sztucznego życia. Dla większości agentów, charakterystyczne są podstawowe cechy jak: asynchroniczne działanie, interaktywność, możliwość poruszania się, niedeterministyczne zachowanie. Termin agent łączy się nierozdzielnie z pojęciem *agency*, którą tu będziemy nazywać intencjonalnością. Jednak cechy charakterystyczne sztucznych agentów nie pokrywają się z tym, co filozofia łączy z pojęciem intencjonalności, mianowicie niezależność, samokontrolę i poczucie odrębności.

340 M. Bunge, *System boundary*, „International Journal of General System”, t. 20, nr 3, 1992, ss. 215–219.

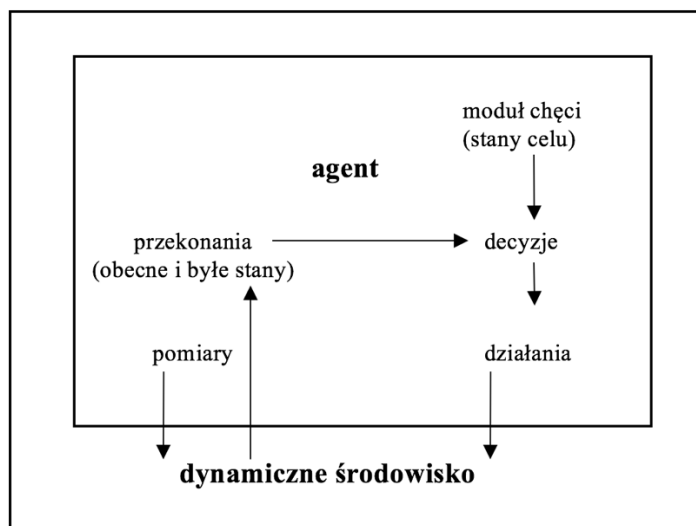
Pojęcie intencjonalności jest wyjątkowo ważne dla sztucznej inteligencji, której celem jest badanie inteligentnej aktywności. Biosemiotyka wskazuje, jak badać to zjawisko z perspektywy semiotycznej. Intencjonalność jest nierozdzielnie związana z organizacją, znaczeniem i koncepcją autonomicznego agenta. Można wyszczególnić kilka właściwości lub warunków, które muszą być spełnione w tworzeniu struktury systemu, przedstawiającego model tworzenia relacji semiotycznej³⁴¹. Badanie znaczenia intencjonalności może pomóc rozwijać inteligencję sztucznych systemów, przy założeniu istnienia samodzielnego działania i ukierunkowania działania na cel. Te atrybuty są kluczowe dla funkcjonowania każdego inteligentnego systemu. Z perspektywy biosemiotyki taki pragmatyczny opis jest minimalną definicją intencjonalności, lecz może zostać wykorzystany w dziedzinie SI, ponieważ jest odpowiednia dla badania funkcji rekurencyjnych w sztucznych systemach, które potem mogą zostać rozwinięte o kryteria dotyczące natury relacji, przewidywania, nawyków, postrzegania siebie i innych, cykli akcji i percepcji itp.³⁴².

Intencjonalność agenta zakłada autonomię, dla której można wyróżnić konkretną domenę, gdzie agent posiada kontrolę, a tym samym – granicę między tym, co agent kontroluje agent a czego nie. Takie granice mogą istnieć w odniesieniu do różnych wymiarów: przestrzennych, czasowych lub funkcjonalnych. Oczywiście zakładamy, że autonomia jest kwestią stopnia, czyli agent może być mniej lub bardziej autonomiczny. Problemem związanym z intencjonalnością pozostaje poczucie tożsamości sztucznych systemów. Istnienie tego problemu można rozwiązać, implikując pewną zdolność do rozróżniania tego, co znajduje się wewnątrz i na zewnątrz agenta. Wtedy to, co znajduje się wewnątrz systemu, łącznie z jego stanami tworzy tożsamość agenta. Jest to forma pewnego zamknięcia systemu, gdzie te aspekty świata, które są uwięzione w granicach tożsamości agenta, są zamknięte przed innymi interakcjami. Zamknięcie to może przyjąć różne formy, od fizycznych (granice strukturalne), przyczynowych (jakiś rodzaj hermetyzacji), lub funkcjonalnych (zamknięcie wejść / wyjść)³⁴³.

341 M. Tønnessen, *The biosemiotic glossary project: Agent, agency*, „Biosemiotics”, t. 8, nr 1, 2015, ss. 125–143.

342 *Ibid.*

343 J.H. Holland, *Hidden order: how adaptation builds complexity*, Cambridge, Mass. 2003.



Rys.6. Model systemu działającego intencjonalnie³⁴⁴

System wykonuje pomiary swojego środowiska i konstruuje uogólnione przekonania, reprezentacje bieżącego stanu i wspomnień z dawnych stanów. Agent musi mieć również nadaną wartość „chęci”, która mobilizuje go do wyłonienia działań zorientowanych na cel. Moduł decyzyjny ma wpływ na potencjalne działania, które są następnie testowane w środowisku. Działania te oddziałują z dynamicznym środowiskiem, które następnie agent odbiera w postaci przyszłych przekonań. W ten sposób konsekwencje decyzji agenta mają wyraźny wpływ na jego przyszły rozwój. System posiada architekturę sterowania, określaną jako negatywną kontrolą sprzężenia zwrotnego³⁴⁵, co oznacza, że istniejący związek ze środowiskiem pozwala mu na podejmowanie własnych decyzji, które minimalizują możliwość błędu. Ścisłe powiązanie agenta z jego otoczeniem jest absolutnie niezbędne. Konsekwencje decyzji agenta są zawsze odzwierciedlane w dynamice interakcji w środowisku i pamięci agenta. System taki posiada przekonania i cele, do których może się samodzielnie odnosić. W ten sposób można skonstruować agenta, który będzie zarazem intencjonalny i semiotyczny.

System intencjonalny bierze udział w procesach percepcji, interpretacji, decyzji i działania ze środowiskiem. Przez wzgląd na możliwość zmiany celu lub działania jest

344 I. Stewart, *Self-organization in evolution: a mathematical perspective*, „Philosophical Transactions of the Royal Society of London. Series A: Mathematical, Physical and Engineering Sciences”, t. 361, nr 1807, 2003, ss. 1101–1123.

345 R.J. Bennett, R.J. Chorley, *Environmental Systems: Philosophy, Analysis and Control*, Princeton University Press 2015, s. 41.

bardzo dynamiczny i plastyczny, podobnie jak w przypadku żywych organizmów³⁴⁶. Samodzielnie działający agenci mogą być postrzegani jako tworzący symboliczne reprezentacje zmiennych konfiguracji, które są przedstawiane jako symbole, których dynamika jest nieistotna. Reakcje systemu nie są wywoływane przez dynamikę, lecz wartość informacji, którą są w stanie wyodrębnić w procesie semiozy. Procesy przetwarzania symboli nie mogą być jednak dynamicznie niespójne. Zmieniająca się dynamika działań musi być rozumiana w ramach stabilnego istnienia całej maszyny, która rozpoznaje na bieżąco reprezentacje warunków początkowych³⁴⁷. Możemy traktować je jako reprezentacje symboliczne, jeśli dynamiczne komponenty są używane do określania warunków początkowych dla dynamicznych konfiguracji samoorganizujących się agentów.

Aby rozwiązać problem spójnego poczucia tożsamości agenta, który będzie umiał odróżnić się od innych systemów i obiektów, należy skupić się na umożliwieniu autonomii w odniesieniu do działania. Agent musi być systemem, który ma wrodzoną swobodę dokonywania wyborów lub podejmowania decyzji dotyczących możliwych działań. Celem jest więc skonstruowanie agentów semiotycznych, które są wystarczająco, ale minimalnie skomplikowane, aby mieć autonomię działania. Aby umożliwić agentom odróżnianie się od innych i ich środowisk Vladimir Lefebvre proponował, by w modelowanym środowisku agentów uwzględnieni zostali inni agenci³⁴⁸. W ten sposób można stworzyć małą grupę agentów o określonym stanie, których działanie może zostać połączone, by osiągnąć bardziej skomplikowaną aktywność lub duży zbiór prostych agentów, który pozwoli na pojawienie się dynamicznych zjawisk lub też mały zbiór złożonych agentów, którzy współdziałają albo konkurują ze sobą³⁴⁹.

Intencjonalni agenci mogą koordynować działania, poznawać środowisko fizyczne i komunikować się. Decyzje agentów funkcjonujących w grupie mogą być ograniczone przez wspólny zestaw wiedzy lub przekonań, wynikłe ze wspólnej ewolucji, rodzaju przekazu lub ograniczenia możliwości uczenia się. Wiedza w tych systemach jest podstawowym atrybutem, ponieważ powoduje powstanie pewnego rodzaju

346 L.M. Rocha, *Evolution with material symbol systems*, „Biosystems”, t. 60, nr 1, 2001, ss. 95–121.

347 H. Patee, *Evolving Self-Reference: Matter, Symbols and Semantic Closure*, „Communication and Cognition - Artificial Intelligence”, t. 12, 1995, ss. 9–27.

348 V. Lefebvre, *Conflicting Structures*, Los Angeles 2015, ss. 88–89.

349 I. Stewart, *Self-organization in evolution: a mathematical perspective...*, *op. cit.*

subiektywizmu, relatywizmu, lub konstruktywizmu. Wynika to z faktu, że biosemiotycznie inspirowani agenci interpretują znaki tak, by były dla nich znaczące. Znaki i symbole nie mają znaczenia same w sobie, lecz znaczenie powstaje w wyniku interpretacji. Wiedza jest lokalna, a agenci posiadają dostęp tylko do świata postrzeganego. Znaczenie, które powstaje w takich systemach jest indywidualne, a zatem działanie nim naznaczone zyskuje intencję. Mimo, że w przypadku sztucznych systemów trudno połączyć intencjonalność ze świadomością typu ludzkiego, widać, że agenci inspirowani biosemiotycznie mają w pewnym sensie świadomość, w sensie umiejętności intencjonalnego ujmowania obiektów i sytuacji, interpretowania dostępnych mu danych, przetwarzania symboli oraz nadawaniu im znaczenia.

Wolność wyboru interpretacji i odniesienia się do obiektów lub sytuacji, którą muszą mieć inspirowani biosemiotycznie agenci, jest formą indeterminizmu. Można argumentować, że wiele prostych automatów stochastycznych może posiadać cechy takiego agenta i na odwrót, że intencjonalny agent działa z perspektywy zewnętrznej jako prosty automat stochastyczny. Jednak sztuczny system dokonuje zamierzonego, rozmyślnego i aktywnego wyboru spośród zestawu możliwych działań. Możliwe, że system w rzeczywistości jest deterministyczny, ale jego złożoność nie pozwala zidentyfikować poprawnie jego intencji. Perspektywa epistemiczna pozwala przyglądać się modelom, które są skonstruowane tak, by możliwy był ich opis. Z perspektywy zewnętrznej nie możemy jednak określić, czy system działa intencjonalnie, czy jest systemem niedeterministycznym. Jednak powyżej opisani agenci mają formę dynamicznej nieprzezroczystości, ponieważ nie można ich modelować jako dynamicznych systemy lub proste systemy określane przez stan, nawet jeśli faktycznie implementują jakąś formę formalizmu. Opisane systemy są niewątpliwie systemami deliberatywnymi, ponieważ są kształtowane jako systemy decyzyjne.

Od czasu powstania kognitywizmu wielu badaczy sztucznej inteligencji zgodziło się, że układ nerwowy żywego organizmu może być modelem obliczeniowym, którego sztuczny odpowiednik SSN odwzorowuje wejścia sensoryczne na wyjścia silników w sztucznych systemach. Mimo tego, oddolne biosemiotyczne podejście mocno dystansuje się od symbolicznego reprezentacyjnego podejścia i skupia się na interaktywności, widzianej jako struktura kształtująca zachowania inteligentne. Zgodnie z teorią Uexküll'a i kilku innych badaczy, każdy inteligentny agent powinien mieć swój własny subiektywny pogląd na świat. Sztuczny układ nerwowy w połączeniu z technikami uczenia się i

metodami ewolucyjnymi jest wykorzystywany w próbach tworzenia sztucznych organizmów, umożliwiając im samoorganizację procesów znaczenia. Kilka przykładów pokazało, jak takie techniki pozwalają robotom znaleźć własny sposób zorganizowania ich kręgów funkcjonalnych, czyli wewnętrznego sposobu wykorzystania znaków i odpowiedzi na bodźce pochodzące ze środowiska. Wykorzystanie takiej samoorganizacji czyni sztuczne organizmy wyjątkowym rodzajem mechanizmu, który powinien stać się przedmiotem badań, jako system semiotyczny.

Można zauważyć, że podejście do autonomicznych agentów i koncepcja ucieleśnionego poznania, są w dużej mierze kompatybilne, lecz istnieje zbyt mało wspólnych opracowań, które mogłyby wnieść znaczący wkład w rozwój SI. Sztuczne systemy, mimo implementacji biologicznych inspiracji i możliwości samoorganizacji, są wciąż dalecy od posiadania inteligencji żywych systemów. Wynika to głównie z faktu, że sztuczne systemy składają się głównie z części mechanicznych i programów sterujących, które nie powstały w wyniku samorozwoju ani ewolucji w celu naturalnego przystosowania się do środowiska. Z drugiej jednak strony inteligencja, autonomia i subiektywność żywych organizmów wynika z integracji ich składników i ich strukturalnej zgodności z otoczeniem.

W prezentowanych badaniach, starano się odsunąć zarzut, że sztucznym systemom z konieczności brakuje semantyki, wewnętrznego znaczenia i własnych celów. Dlatego należałoby coraz bardziej minimalizować rolę ludzkiego programisty i badać alternatywne podejścia. Prace Uexküll'a i jego następców wydają się wyjątkowo istotne, ponieważ mogą pomóc uniknąć nadmiernie mechanistycznych teorii. Mogą też pomóc wyjaśnić jak olbrzymią rolę odgrywa w powstawaniu inteligencji ucieleśnione poznanie i osadzenie sztucznych systemów w ich światach, co pomoże zrozumieć możliwości i ograniczenia sztucznej inteligencji.

4 Radykalny konstruktywizm jako paradygmat rozwoju badań nad sztuczną inteligencją

Konstruktywizm jest tezą, że wiedza nie jest biernym odzwierciedleniem rzeczywistości, ale musi być związana z konstrukcją aktywnego podmiotu. Pewne aspekty konstruktywizmu są szczególnie istotne dla sztucznej inteligencji. Reprezentacje tworzone przez agenta niekoniecznie odzwierciedlają strukturę środowiska oglądanego przez zewnętrznego obserwatora. Są to raczej modele kompatybilne ze środowiskiem agenta i interakcjami, które agent z nim podejmuje. Reprezentacje powstają w odniesieniu do potrzeb i celów agenta, w szczególności w odniesieniu do jego prób kontrolowania własnych działań i otoczenia. Nie powstają w wyniku biernej obserwacji i wnioskowania, lecz jako wynik aktywnej interakcji ze środowiskiem. W rzeczywistości reprezentacje te wymagają interakcji ze środowiskiem, aby w ogóle móc funkcjonować jako podstawy mechanizmów wyboru akcji.

Konstruktywizm został przyjęty przez niektórych badaczy sztucznej inteligencji i sztucznego życia (np. [10, 29, 31]) jako podejście do budowania i odkrywania podstaw kształtowania się inteligentnych zachowań. Zamiast szczegółowo opracowywać architekturę sztucznego systemu powstającego jako wynik rozważań *a priori*, mechanizmy i funkcje poznawcze agentów są opracowywane przy użyciu mechanizmów samoorganizacji i ewolucji. Aby takie podejście było wykonalne, agenci muszą być usytuowani w swoim środowisku, ponieważ wykorzystanie cech środowiska i silna praktyczna interakcja wywołuje skuteczny rozwój inteligentnych zachowań agenta, w przeciwieństwie do inżynierskiego podejścia, w którym agenci są najpierw projektowani i konfigurowani, a następnie dopiero osadzani w środowisku. Konstruktywizm w sztucznej inteligencji można postrzegać jako podejście, które obejmuje zarówno ucieleśnione i usytuowane podejście w SI a także założenia biosemiotyczne. Zamiast rozwoju sztucznych systemów poprzez analizę, projektowanie i testy wykonywane przez ludzkich projektantów w oparciu o ich wiedzę, rozwój inteligentnych działań agentów jest osiągany poprzez samoorganizujące się i ewolucyjne procesy działające na agenta znajdującego się w swoim *Umwelcie*.

Dla rozwoju SI podejście konstruktywistyczne jest ważne z kilku powodów. Przede wszystkim badania inteligentnych działań sztucznych systemów opierają się na

właściwościach badanych agentów, zamiast zakładać istnienie zewnętrznej niezależnej rzeczywistości, w której istnieje generalna inteligencja. Ponadto sztuczni agenci są budowani w konstruktywistycznym stylu SI, w którym treść i rozwój inteligencji agenta są określone w możliwie najmniejszym stopniu, a reprezentacje agenta są wewnętrzne i ugruntowane w jego działaniu. Znaczenie powstające w sztucznym systemie nie jest wcześniej ustalone przez programistę, więc semantyka agenta jest ugruntowana [20] poprzez budowanie doświadczenia i wiedzy o środowisku, w którym się znajduje oraz jednocześnie o jego własnych stanach wewnętrznych. Wreszcie konstruktywizm można uznać za rozsądne wyjaśnienie obserwowanego zachowania agentów w opisanym modelu, a zatem też jako prawdopodobne wyjaśnienie jego inteligentnych zachowań.

Źródło powstania konstruktywizmu należy szukać w ideach Kanta, lecz termin ten został po raz pierwszy użyty przez Jeana Piageta w celu określenia procesu, w którym jednostka konstruuje swoje spojrzenie na świat. W książce *La construction du réel chez l'enfant*, Piaget przedstawił tezę twierdzącą, że człowiek rozwija swoje umiejętności poznawcze takie jak koncepcje przestrzeni, czasu, przedmiotów i przyczyn jako podstawowe ramy koncepcyjne podczas interakcji sensomotorycznej ze środowiskiem³⁵⁰. W ten sposób konstruuje w umyśle rzeczywistość, która współgra ze środowiskiem, oraz pozwala rozpoznać własne możliwości i ograniczenia. Interakcja ze środowiskiem jest podstawową formą poznania świata, która prowadzi do działań, które przynoszą oczekiwany cel, a tym samym przyczyniają się do formowania się zachowań, które określić można jako inteligentne. Piaget nie miał prawdopodobnie wglądu w prace Jakoba von Uexküll'a, lecz wiele pomysłów tych dwóch myślicieli wykazuje spore podobieństwo. Opisany przez Uexküll'a *Merkwelt* i *Wirkwelt*, jest przedstawiony przez Piageta jako poziom sensomotoryczny w rozwoju umiejętności poznawczych. Należy zauważyć, że obaj badacze byli pod dużym wpływem filozofii Kanta, przyjmując, że to co nazywamy wiedzą, musi być determinowane w dużym stopniu przez sposoby postrzegania i pojmowania przez podmiot poznający. Zostało podkreślone już w pierwszym zdaniu *Krytyki czystego rozumu*, które mówi, że doświadczenie jest pierwszym produktem, który rodzi nasza inteligencja, działając na materiale wrażeń zmysłowych³⁵¹.

Zdobywanie wiedzy jako umiejętność adaptacyjna było opisane już wcześniej, lecz to Piaget zauważył, że adaptacja w dziedzinie poznawczej to umiejętność stworzenia

350 J. Piaget, *La construction du réel chez l'enfant*, Neuchâtel; Paris 1937.

351 I. Kant, *Krytyka czystego rozumu...*, op. cit.

pewnej równowagi pojęciowej³⁵², która pozwoli wydajniej funkcjonować w danej rzeczywistości. Mechanizmy poznawcze postrzegamy jako mentalne, ponieważ poznanie nie posiada żadnej ikonicznej korespondencji z ontologiczną rzeczywistością, ponieważ nie powstaje jako jej obraz lub kopia. Poznanie było postrzegane przez Piageta jako „instrument adaptacji, jako narzędzie dopasowania do świata naszych doświadczeń”³⁵³. Inteligentne zachowanie powstaje zatem przez nabycie nowych doświadczeń oraz ich asymilację w istniejących już strukturach poprzez modyfikację lub stworzenie zupełnie nowych mechanizmów. Dla Piageta oznaczało to powstanie równowagi pojęciowej, która oznaczała dopasowanie przekonań i oczekiwań podmiotu z jego nowymi doświadczeniami.

Ernst von Glasersfeld sformułował teorię radykalnego konstruktywizmu, którego podstawowe zasady są następujące:

- proces poznania nie może być przekazany ani nie powstaje podczas biernego odbioru danych zmysłowych;
- poznanie jest procesem aktywnie konstruowanym przez podmiot poznający;
- poznanie ma charakter adaptacyjny w sensie biologicznym jako mające na celu zwiększenie dopasowania do środowiska i poprawę własnej wydajności;
- poznanie nie prowadzi do odkrycia obiektywnej rzeczywistości ontologicznej świata, lecz jest podstawą organizacji empirycznej podmiotu³⁵⁴.

Konstruktywizm nie jest jednak rozwinięciem sensualizmu w teorii percepcji, ponieważ sensualiści uważali poznanie za wynik działania na organizm wrażeń i danych sensorycznych. Empirystyczni sensualiści twierdzili, że rozwój poznawczy zachodzi spontanicznie, bez aktywności podmiotu – samo doświadczenie tworzy skojarzenia między doznaniem³⁵⁵. W konstruktywizmie przedmiot odgrywa aktywną rolę w percepcji. Sam je buduje i konstruuje wiedzę, przetwarzając informacje sensoryczne przy pomocy określonych zasad, przy użyciu doświadczeń i wykorzystaniu własnej historii. Empirystyczny sensualizm i konstruktywizm w interpretacji percepcji są pozycjami przeciwnymi i wzajemnie się wykluczającymi. W teorii konstruktywistycznej podmiot jest

352 J. Piaget, *The construction of reality in the child*, M. Cook (tłum.), New York 1954, s. xii.

353 E. von Glasersfeld, *The Radical Constructivist View of Science*, „Foundations of Science”, t. 6, nr 1, 2001, s. 14.

354 *Ibid.*, s. 51.

355 V.A. Lektorski, *Realism, Antirealism, Constructivism, and Constructive Realism in Contemporary Epistemology and Science*, „Journal of Russian & East European Psychology”, t. 48, nr 6, 2010, ss. 5–44.

interpretowany jako aktywny konstruktor a świat zewnętrzny jest rodzajem wyzwacza procesów mentalnych. Rzeczywistość przyczynowo działa na narządy zmysłów, to, co jest mu bezpośrednio dane, to nie świat zewnętrzny, ale jego własne subiektywne stany, które mogą nie przypominać tego, co istnieje poza świadomością. Rzeczywistość zewnętrzna jest pewnego rodzaju bytem idealnym, z którego w wyniku przetwarzania, wyłania się inny byt – percepcja.

Glaserfeld nalegał, by ten rodzaj konstruktywizmu nazwać „radykałnym”, ponieważ uważał, że konstrukcja rzeczywistości jest wszechobecna i wobrazie rzeczywistości nie istnieje nic, co nie jest wynikiem procesu konstrukcji. Niewątpliwie jest to duży kontrast w porównaniu z innymi teoriami wiedzy, a szczególnie stwierdzenie, że konstrukcje koncepcyjne reprezentują w pewien sposób niezależną od umysłu rzeczywistość zewnętrzną³⁵⁶.

Radykalny konstruktywizm został poszerzony w celu uwzględnienia powiązanych podejść. Te pokrewne podejścia to cybernetyka drugiego rzędu Heinza von Foerster i teoria systemów autopoetycznych Maturany i Vareli. Mimo, że te podejścia przedstawiają odmienne przesłanki, to w wielu punktach są ze sobą zbieżne, a w szczególności, jeśli chodzi o nacisk kładziony na funkcje poznawczego zamknięcia. W radykalnym konstruktywistycznym sensie zamknięcie odnosi się do jakości systemu poznawczego, w którym zmiany jego stanów propagowane są wyłącznie w obrębie tworzących je sieci procesów. Zamknięcie nie oznacza jednak, że systemy poznawcze są izolowane od środowiska, a tym samym nie reagują na zmiany w otoczeniu, wręcz przeciwnie – podmiot aktywnie percypuje świat.

Teoria radykalnego konstruktywizmu jest w dużej mierze zgodna z biologicznie inspirowanymi badaniami w dziedzinie sztucznej inteligencji i sztucznego życia, które badają adaptację agentów i powstawanie ich wewnętrznej struktury poznawczej w trakcie interakcji sztucznego systemu i jego środowiska. Sztuczne systemy, które zostały opisane do tej pory, aktywnie budują własną wiedzę, w miejsce odgórnego programowania, opierając się na transformacji sensomotorycznej zamiast na wewnętrznym modelu zewnętrznej rzeczywistości zaimplementowanym przez programistę. Ten kierunek badań nad sztuczną inteligencją można określić jako próbę wyposażenia sztucznych systemów w zdolności umysłowe i behawioralne żywych organizmów. Zwrot badań w kierunku

356 E. Von Glasersfeld, *Radical Constructivism: A Way of Knowing and Learning*. London, Wasington, DC, London 1995, s. 15.

radikalnego konstruktywizmu jest adekwatny z przyjętą perspektywą ucieleśnionego i osadzonego w środowisku podmiotu, umożliwia również zbudowanie sztucznych systemów, których inteligencja mogłaby być cechą powstałą w wyniku samodzielnego rozwoju umiejętności poznawczych. Można zatem przyjąć, że ten kierunek rozwoju w dziedzinie sztucznej inteligencji jest oparty na radykalnym konstruktywizmie, ponieważ badania wykorzystują dane sensomotoryczne kształtujące stany wewnętrzne sztucznego systemu i bezpośrednio wpływają na wyłanianie się zachowań, które można określić jako inteligentne.

Radykalny konstruktywizm ujmuje szereg specyficznych cech współczesnych nauk kognitywnych, w tym cech charakterystycznych dla sztucznej inteligencji. Podejście to zawiera szereg ważnych cech charakteryzujących aktywność poznawczą. Zgodnie z klasycznym rozumieniem, które pochodzi od Platona, wiedza, w przeciwieństwie do opinii i przekonań, z konieczności zakłada związek z rzeczywistością. Ogólnie rzecz biorąc, w życiu codziennym używa się terminu „wiedzieć” nie tylko w odniesieniu do „tego, co jest” w rzeczywistości, ale także w odniesieniu do zdolności, zdolności do robienia czegoś, do wykonania, a zatem wiąże się z inteligentnym działaniem. Organizm lub sztuczny agent wie, jak poruszać się by unikać przeszkód, reagować na specyficzne konteksty środowiska, komunikować i tak dalej. Główną ideą radykalnego konstruktywizmu epistemologicznego jest więc to, że „wiedzę o” można sprowadzić do „wiedzy, w jaki sposób”.

Z punktu widzenia radykalnego konstruktywizmu, badanie inteligencji nie polega na zajmowaniu się realnymi przedmiotami, a jedynie konstrukcjami podwójnego rodzaju. Po pierwsze, są to różnego rodzaju produkty interakcji. Konstrukty te będą się różnić w zależności od rodzaju badanego przedmiotu i jego historii. Dlatego jeśli by było pojawienie się osobowości, jaźni, inteligencji i świata subiektywnego, miałyby one różny wygląd w różnych systemach. Po drugie, badający inteligencję człowiek sam konstruuje przedmiot badań, który nie istnieje poza tym procesem. To, co uważa się za inteligencję, w rzeczywistości nie jest tym, co istnieje realnie poza jego umysłem. W tym przypadku badacz sztucznej inteligencji nie jest obiektywnym uczestnikiem badania, ale uczestnikiem tworzącym obraz pewnych zachowań inteligentnych, o których można mówić tylko w sensie hipotetycznym, ponieważ istnieje ona tylko w ramach konstruktywnej działalności. Dlatego też ten paradygmat doskonale oddaje sposób myślenia o sztucznej inteligencji,

łącząc w sobie założenie biosemiotycznego poznania i perspektywę patrzenia na odmienną inteligencję maszynową.

4.1 Sztuczne Umwelty – perspektywa powstania sztucznych fenomenalnych światów

Zastosowanie sztucznego odpowiednika systemu nerwowego pozwala faktycznie łączyć ciało fizyczne sztucznego agenta z jego czujnikami i silnikami. W ten sposób oddziałuje z fizycznymi obiektami swojego otoczenia, niezależnie od mediacji lub interpretacji obserwatora. Sieć neuronowa staje się integralną częścią agenta, który jest nie tylko ucieleśniony na sposób fizyczny, lecz również fizycznie osadzony w swoim środowisku. Ponadto jego reprezentacje wewnętrzne są formowane w fizycznej interakcji ze światem, który odzwierciedlają, co sprawia, że jego działania są skonstruowane poprzez procesy adaptacyjne. Sieć kontrolera robota może być częścią kręgu lub wielu kręgów funkcjonalnych, co sugeruje, że sztuczny agent może budować własny *Umwelt*.

Odejście od teorii mechanistycznej w biologii sprawiło, że zwierzęta przestały być postrzegane jako maszyny, lecz podmioty, których życie opiera się na działaniu i postrzeganiu³⁵⁷. Zauważono, że wszystko, co podmiot postrzega, staje się jego światem percepcyjnym, a wszystko, co robi, światem efektorów. Radykalny konstruktywizm stawia tezę, że świat jest jedynie konstruktem powstającym w wyniku otrzymywania danych z zewnątrz. Analiza takich systemów jest możliwa tylko wtedy, gdy systemy te są interpretowane jako naprawdę istniejące w prawdziwym środowisku i w interakcji z tym ostatnim. Niemożliwe jest zrozumienie inteligencji w takich systemach bez uwzględnienia ich potrzeby uzyskania informacji z zewnątrz. Podobnie w teorii biosemiotycznej Uexküll'a tworzą się *Umwelty*. Możliwość przypisania *Umweltu* sztuczemu systemowi jest wciąż dyskutowanym problemem³⁵⁸. Z jednej strony wydaje się całkiem możliwe, by sztuczne

357 D. Hammond, *The Science of Synthesis: Exploring the Social Implications of General Systems Theory*, University Press of Colorado 2010, ss. 37–40.

358 C. Emmeche, *Does a robot have an Umwelt? Reflections on the qualitative biosemiotics of Jakob von Uexküll*, „Semiotica”, t. 134, nr 1/4, 2001, ss. 653–694; M.I.A. Ferreira, M.G. Caldas, *Modelling Artificial Cognition in Biosemiotic Terms*, „Biosemiotics”, t. 6, nr 2, 2012, ss. 245–252; K. Kaveh, M.D. Bui, P. Rutschmann, *Development of an Artificial-Neural-Network-based*

systemy posiadały *Umwelt*, skoro mogą go posiadać nawet najprostsze organizmy. Z drugiej strony maszyny pozostają wciąż programowalnymi systemami; bez względu na to, jak skomplikowane są obwody sztucznej sieci neuronowej nie są w stanie nic odczuwać ani rozumieć – pozostają zamknięte w obrębie Chińskiego Pokoju³⁵⁹.

W jednej z niewielu publikacji poświęconej temu zagadnieniu Claus Emmeche zastanawiał się, czy roboty mogą posiadać swój własny *Umwelt*³⁶⁰. Jego odpowiedź przechyliła się ku stwierdzeniu, że jest to niemożliwe. Przekonanie, że sztuczne systemy nie posiadają *Umweltu*, przyjmowane jest przez tych, którzy uważają, że co prawda system rzeczywiście używa kręgów funkcjonalnych w sensie pętli sprzężeń zwrotnych, nie postrzega się jednak tego kręgu za prawdziwy przykład semiotycznego formowania sensu poprzez aktywne, uczestniczące przetwarzanie znaków. W tej perspektywie charakter *Umweltu* stanowi funkcjonalno-cybernetyczny aspekt przetwarzania sygnału w systemie, w miejsce fenomenalnego świata żywego organizmu. Dodatkowo uznaje się, że tylko żywy organizm jest konstruowany jako aktywny, intencjonalny podmiot. Zatem można powiedzieć, że tylko autentyczne żywe organizmy (a zwłaszcza zwierzęta) żyją w *Umweltach*³⁶¹.

Umwelt i środowisko, są istotne zarówno dla biologii, nauk poznawczych i sztucznej inteligencji. Organizmy i roboty wchodzą w interakcje ze środowiskiem w procesach semiozy. Jednak termin środowisko musi być rozróżnione od terminu *Umwelt*. *Umwelt* w sensie Uexküll'a może być tylko doświadczany przez podmiot, a jego doświadczenia żaden badacz nie jest w stanie rozpoznać, środowisko zaś należy postrzegać jako zbiór wszelkich możliwych przedmiotów rzeczywistości zewnętrznej. Charakterystyczną cechą dla fenomenalnych światów jest to, że żadna z doświadczanych właściwości materii nie pozostałaby stała, gdyby możliwe było jej zbadanie we wszystkich *Umweltach*. W różnych percypowanych światach każdy obiekt zmienia znaczenie wraz ze

concept for hydro-morphodynamic modelling in rivers, [w:] *New Challenges in hydraulic research and engineering. Proc. of the 5th IAHR-Europe Congress*, 2018, ss. 721–722; D.C. Lahti, *The limits of artificial stimuli in behavioral research: the umwelt gamble*, „Ethology”, t. 121, nr 6, 2015, ss. 529–537; L. Rudolph, *A Unified Topological Approach to Umwelts and Life Spaces Part II: Constructing Life Spaces from an Umwelt*, [w:] *The Observation of Human Systems*, Routledge 2017, ss. 117–140.

359 J. Searle, *Umysły, mózgi i programy...*, op. cit.

360 C. Emmeche, *Does a robot have an Umwelt? Reflections on the qualitative biosemiotics of Jakob von Uexküll...*, op. cit.

361 *Ibid.*

strukturą swoich właściwości. W konstruktywistycznej perspektywie nie istnieje żadna podstawa do kompleksowego, obiektywnego spojrzenia na rzeczywistość³⁶². Ta niemożność bardzo ogranicza dyskusję na temat powstawania znaczenia i *Umweltów* sztucznych systemów, ponieważ zakłada, że nie ma możliwości wglądu w to czy sztuczny system czegoś doświadcza i czy to doświadczenie niesie ze sobą znaczenie.

Istnieje jednak możliwość złagodzenia sprzeczności między konstruktywistycznym poglądem na świat zbudowany z percepcji, których nie możemy zbadać i obiektywistycznym przekonaniem, że gdzieś istnieje prawdziwy świat zbudowany z obiektywnych przedmiotów, a co za tym idzie, możliwość zbadania czy w świecie sztucznego systemu powstaje znaczeniowy *Umwelt*. W tym celu należy się przyjrzeć procesowi przetwarzania znaków w sztucznym systemie. Według Peirce'a semioza jest procesem mediacji i może zostać wykorzystana w celu zbadania relacji podmiotu z przedmiotem *Umweltu*³⁶³. Na początku należy dokonać rozróżnienia na dynamiczne i bezpośrednie obiekty świata. Obiekt dynamiczny istnieje poza rzeczywistym znakiem i pozwala ustalić odniesienie znaku do jego przedstawienia. Obiekt bezpośredni pokazuje, w jaki sposób obiekt dynamiczny jest reprezentowany symbolicznie poprzez indeks lub ikonę³⁶⁴. Dlatego też można ustalić, że obiekt bezpośredni znajduje się w *Umwelcie* w takiej postaci, w jakiej został skonstruowany przez *Innenwelt* żywej istoty lub też w przypadku sztucznego agenta – przez jego stany wewnętrzne. Dynamiczny obiekt zaś pozostaje w *Umwelcie*, niezależnie od jakiegokolwiek konkretnego aspektu³⁶⁵. Dzięki badaniu procesu semiozy możemy obserwować bezpośredni obiekt w *Umwelcie* sztucznego systemu i dynamiczny obiekt poza nim, co daje dostęp do środowiska znacznie szerszego niż *Umwelt* badacza.

Przykładem może być tu mobilny robot zawierający program pozwalający mu uczenie się w trakcie działania w środowisku. Dzięki temu modułowi percepcyjnemu, maszyna tworzy reprezentację zewnętrznego świata i jego przedmiotów. Mimo, że znajomość percepcyjna środowiska nie jest wynikiem działania umysłu robota, lecz symulacją programisty, to agent w każdym działaniu wchodzi w kontakt ze światem

362 J. von Uexküll, *Bedeutungslehre*, Leipzig 1940, ss. 230–231.

363 W. Noth, *Semiotic machines*, „Cybernetics & Human Knowing”, t. 9, nr 1, 2002, ss. 5–21.

364 C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 4: The Simplest Mathematics*, Charlottesville, Va 1994, akap. 536.

365 C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 8: Reviews, Correspondence, and Bibliography*, Charlottesville, Va 1994, akap. 133.

dynamicznych obiektów. Wynik pracy programisty jest tu niezwykle istotny, ponieważ przedstawia symbolizację obiektu jego *Umweltu*, który dla sztucznego systemu jest obiektem dynamicznym. W procesie nauki, podczas prób i błędów, obiekty te zyskują reprezentację w stanach wewnętrznych maszyny. Wynik interpretacji rzeczywistości powstały podczas obliczeń w dynamicznych interakcjach ze środowiskiem jest interpretantem w procesie semiozy. Zatem rzeczywisty proces semiozy, w którym powstaje interpretacja (zawierająca znaczenie) wpływa na przyszłe postrzeganie *Umweltu*, a zdobyta wiedza o subiektywnym środowisku jest dostępna jako bezpośredni obiekt *Umweltu*.

Istnieją też inne argumenty przemawiające za tym, że niektóre sztuczne systemy (przeważnie roboty) współdziałają ze swoim środowiskiem podobnie jak żywe organizmy. Sztuczni agenci tak jak organizmy mają zdolność rozróżnienia między sobą a lokalnym środowiskiem, które jest przez nie percypowane w subiektywny sposób. Poruszający się robot, którego czujniki wykrywają tylko konkretne obiekty środowiska, jak na przykład przeszkody, które omija, tworzy swoją reprezentację środowiska z tych właśnie przedmiotów. Jest to jego *Umwelt* zbudowany z przeszkód, w świecie pełnym okien, drzwi, plakatów na ścianach, zawartości maszyn wendingowych itd. Te i inne aspekty świata pozostają dla niego niedostępne, tak jak ultradźwięki wydawane przez nietoperze, pozostają niesłyszalne dla ludzi. Taki sztuczny system przetwarza bodźce środowiskowe selektywnie, podobnie jak żywe organizmy. Reaguje na sygnały percepcyjne wszystkich przeszkód pojawiających się w zasięgu jego czujników, zatem postrzega świat, jak chociażby człowiek, który słyszy tylko dźwięki o ograniczonej częstotliwości. Roboty wyposażone w różne czujniki postrzegają swoje środowisko jako subiektywne otoczenie, dzięki czemu wchodzi w interakcję ze światem w sposób opisany przez Uexküllą. Wskazuje to na fakt, że konstruują własne reprezentacje świata. Te symboliczne reprezentacje są związane z jego stanem wewnętrznym – *Innenweltem*, podobnie jak dzieje się to w przypadku żywych istot. Również porównanie systemu sensomotorycznego biologicznego organizmu z sztucznym systemem wskazuje wiele podobieństw. Sztuczny system ma czujniki percepcyjne i moduły efektorowe, które pozwalają mu na interakcję ze środowiskiem, podobnie jak organizm jest wyposażony w narządy i receptory percepcyjne oraz organy efektorowe, które pozwalają im poznać naturę *Umweltu*.

Wydaje się zatem, że sztucznego agenta można uznać za osadzonego we własnym *Umwelcie*, składającym się ze świata percepcyjnego (*Merkwelt*), zawierającego

przedmioty stałe (lub ich brak) i niosące znaczenia, np. przeszkoda i wolna przestrzeń, oraz świata operacyjnego (*Wirkwelt*) powstającego w wyniku pracy silników. Wewnętrzny świat agenta jest jego wewnętrznym przepływem znaków w sztucznej sieci neuronowej i interaktywnych reprezentacji. Jest to samorganizowany przepływ znaków prywatnych w interakcji agent-środowisko. Adaptacja kontrolerów sieci neuronowej w robocie może być postrzegana jako konstrukcja interaktywnych reprezentacji działania poprzez tworzenie, dostosowanie i optymalizację kręgów funkcjonalnych w interakcji z otoczeniem.

Mimo powyższych argumentów istnieją znaczące różnice semiotyczne między żywym organizmem a sztucznym systemem w sposobie interakcji ze środowiskiem. Życie biologicznych istot jest do pewnego stopnia autoreferencyjne, zaś sztuczny system działający w *Umwelcie* pozostaje alloreferencyjny. Idea samoodniesienia jest mocno związana z interakcją organizmu z *Umweltem*. Bycie biologiczną istotą według Uexküll'a było głównym punktem odniesienia dla koncepcji znaczenia, ponieważ poczucie podmiotowości leży w swoim własnym istnieniu, w swoim własnym ja³⁶⁶. Semioza zachodzi w centrum *Umweltu*, a znaki są interpretowane w zgodzie z kodem gatunkowym³⁶⁷. Jednym z istotnych problemów sztucznego *Umweltu* jest fakt, że świat fenomenalny żywego organizmu zawiera w sobie znaczenie, które jest związane z celami podmiotu. Maszyny nie posiadają własnych celów, ponieważ są konstruowane tak, by wypełniały zadania narzucone im przez człowieka – nie można stwierdzić, że posiadają swoje ostateczne cele same w sobie. Ostatecznym układem dla sztucznego systemu jest przecież przyczynowość narzucona przez programistę bądź użytkownika, a zatem jego proces semiotyczny pozostaje alloreferencyjny.

Wydaje się jednak, że można rozważyć istnienie pewnego rodzaju celowości w sztucznym systemie, starając się połączyć jego cele z celami użytkownika. Można argumentować, że taki rodzaj celowości nie ma nic wspólnego z wolnością wyboru, agent nie próbuje utrzymać swojej homeostazy, nie wykorzystuje zasobów, aby przeżyć lub uniknąć zagrożeń. Te względy powinny mieć wpływ na kształtowanie *Innenweltu* agenta. Zasoby obliczeniowe maszyny mogą być na tyle duże, aby część z nich została przeznaczona do rozpoznania sytuacji (zagrożenia) i uruchomiłaby mechanizmy odpowiedzialne za dodatkowe działanie (przetrwanie). Zaletą tego rozwiązania jest to, że

366 J. von Uexküll, *Umwelt und Innenwelt der Tiere...*, op. cit.

367 J. von Uexküll, *Theoretische biologie...*, op. cit.; T. von Uexküll, *Introduction: Meaning and science in Jakob von Uexküll's concept of biology...*, op. cit.

można by je stosować niezależnie od zadania. Konkretny cel dla agenta, taki jak zadanie przeżycia, można by osiągnąć poprzez promowanie tego mechanizmu w celu nadania większego znaczenia sytuacjom związanym z celem, a tym samym przydzielenie większej ilości zasobów obliczeniowych do ich zidentyfikowania³⁶⁸. *Umwelt* agenta musi być wyłaniającą się relacyjną właściwością sprzężenia między podmiotem a jego otoczeniem. Musi być jego własnym konstruktem. Jeśli system może wykorzystać stan wewnętrzny w oparciu o historię spostrzeżeń i działań, znalezienie mechanizmu, który określi ten stan wewnętrzny może doprowadzić do pojawienia się znaczącego celu w *Umwelt*ie.

4.2 Konsekwencje przyjęcia radykalnego konstruktywizmu – odmienność inteligencji maszynowej

Badanie samoregulacji, autonomii i hierarchicznych układów doprowadziło do krystalizacji pojęć, takich jak przyczynowość, sprzężenie zwrotne, równowaga, adaptacja, kontrola oraz do pojęcia funkcji, systemu i modelu. Większość z tych terminów jest tak popularna, że stały się modnymi słowami i pojawiają się w wielu kontekstach. Sztuczna inteligencja stara się systematycznie powiązać te pojęcia, które zostały ukształtowane i powiązane z tymi terminami w interdyscyplinarnej analizie. Już podczas powstawania cybernetyki istniały dwie główne orientacje, które zostały przeniesione do dziedziny sztucznej inteligencji. Pierwsza z nich dotyczy koncepcji i projektu rozwoju technologicznego opartego na mechanizmach samoregulacji za pomocą sprzężenia zwrotnego i przyczynowości. Rezultatem tych badań było powstanie robotów przemysłowych, automatycznych pilotów i innych automatów oraz komputerów. Komputery z kolei doprowadziły do opracowania funkcjonalnych modeli mniej lub bardziej inteligentnych procesów. To pole sztucznej inteligencji dziś obejmuje nie tylko systematyczne badania w zakresie rozwiązywania problemów, dowodzenia twierdzeń, teorii liczb i innych obszarów logiki i matematyki, ale także wyrafinowane modele procesów wnioskowania, sieci semantycznych i umiejętności takich jak np. gra w szachy lub interpretacja języka naturalnego. Druga orientacja skoncentrowała się na szerokim

368 C. Salge, C. Guckelsberger, *Does empowerment maximisation allow for enactive artificial agents?*, [w:] MIT Press 2016, ss. 704–711.

pytaniu dotyczącym poznania i inteligencji, starając je umieścić w ramach koncepcyjnych samoorganizacji³⁶⁹. Tutaj pojawił się konflikt z tradycyjnym założeniem naukowym, który zakłada, że opisy naukowe i wyjaśnienia powinny, a nawet mogą, przybliżać strukturę obiektywnej rzeczywistości, która istnieje niezależnie od obserwatora. Radykalny konstruktywizm w biologii, psychologii i w sztucznej inteligencji, który bierze pod uwagę podstawowe pojęcia samoregulacji, autonomii i informacyjnie zamknięty charakter organizmów poznawczych, wskazuje na odmienną perspektywę. Zgodnie z tym poglądem, postrzeganie rzeczywistości jest sprawą subiektywną. Zatem również działania postrzegane jako inteligentne przez człowieka, z perspektywy innego obserwatora mogą być interpretowane tylko jako seria „jakichś” działań³⁷⁰.

Najsilniejszy argument dotyczący przymusu oderwania się od przekonania o obiektywności obserwatora pochodzi z fizyki. Teoria względności i mechanika kwantowa zmuszają do namysłu, czy dotyczą obiektywnej rzeczywistości, czy raczej świata opisanego w wyniku obserwacji. Albert Einstein miał nadzieję, że interpretacja realistyczna ostatecznie doprowadzi do jednorodnego spojrzenia na wszechświat, istniejący poza ludzkimi zmysłami³⁷¹, lecz Werner Heisenberg i Niels Bohr wykazali podczas eksperymentów, że szanse na obiektywizm są nikłe, ponieważ badane przez nich cząstki, są tym, czym są dopóki ktoś ich nie zauważy³⁷². Można zatem tworzyć modele, lecz problem istnienia obiektywnej, niezależnej od obserwatora rzeczywistości prędzej czy później pojawi się ponownie. Dlatego też tak skomplikowanych pojęć jak wiedza czy inteligencja nie można traktować jako wyniku działania w obiektywnej rzeczywistości, ale raczej jako szczególny sposób organizowania doświadczenia. Obiektywność w tradycyjnym sensie, jak zauważył Heinz von Foerster, jest poznawczą wersją fizjologicznej

369 H.R. Maturana, *Autopoiesis and Cognition...*, op. cit.

370 Problem obserwatora pojawia się w również w naukach społecznych i jest związany z pojęciem zrozumienia. Szczególnie w antropologii uświadomiono sobie, jak trudnym przedsięwzięciem jest analiza struktury obcej kultury, mimo podjęcia wysiłku, aby zrozumieć tę kulturę w kategoriach stworzonych przez nią struktur pojęciowych. Podobnie jest w badaniach literatury zagranicznej lub historycznej: podejście hermeneutyczne zyskuje na popularności, ponieważ jego celem jest rekonstrukcja znaczenia w kategoriach pojęć i ram koncepcyjnych w czasie i miejscu życia autora.

371 J. Shelton, *The role of observation and simplicity in Einstein's epistemology*, „Studies in History and Philosophy of Science Part A”, t. 19, nr 1, 1988, ss. 103–118.

372 D.R. Hawkins, *I: Reality and Subjectivity*, Hay House, Inc 2013, ss. 573-580.

ślepej plamy: nie widzimy tego, czego nie widzimy³⁷³. Zatem przekonanie o możliwości powstania sztucznej inteligencji, która przypominałaby człowiekowi jego własną inteligencję może okazać się subiektywnym złudzeniem.

Pomimo inspiracji biologicznej współczesne sztuczne systemy wciąż różnią się radykalnie od żywych organizmów. Posiadają pewien stopień samoorganizacji i autonomii, lecz ich zdolności są nie tylko odmienne, lecz także dalekie od zdolności żywych organizmów. Wynika to z faktu, że systemy te zazwyczaj składają się z części mechanicznych i programów sterujących (czyli sprzętu i oprogramowania) stworzonych przez człowieka w celu zapewnienia działań, które człowiekowi mają służyć. Autonomia i podmiotowość żywych systemów wyłania się z interakcji ich składników, czyli autonomicznej jedności komórkowej (*Zellautonomie*)³⁷⁴. Znacząca interakcja między tymi jednościami oraz jednością z otoczeniem jest wynikiem ich strukturalnej zgodności, co wskazali Uexküll, a także Maturana i Varela³⁷⁵. Tak więc inteligencja jest własnością organizmu jako autonomicznej jedności komórkowej, a początkowa strukturalna zgodność z otoczeniem wynika ze specyficznego rozwoju ewolucyjnego, który umożliwił przetrwanie i rozmnażanie. Zapewnienie artefaktom możliwości rozwoju samoorganizacji i umożliwienie im konstruowania autonomicznych procesów może być postrzegane jako próba zapewnienia im sztucznej ontogenezy. Jednak próba zapewnienia im autonomii w ten sposób wydaje się skazana na niepowodzenie, ponieważ autonomii nie można zaprogramować. Próba wprowadzenia sztucznego systemu w jakąś formę strukturalnej zgodności z otoczeniem może odnieść sukces, ale tylko wtedy, gdy kryterium zgodności nie jest wynikiem samorozwoju sztucznego systemu, lecz leży w koncepcji konstruktora systemu i w ocenie obserwatora. Tak właśnie dzieje się, gdy robot jest szkolony w zakresie dostosowywania swojej struktury do rozwiązywania zadania określonego przez projektanta.

Projektanci inteligentnych systemów kładą nacisk na równowagę między wewnętrznymi strukturami behawioralnymi i koncepcyjnymi maszyny a doświadczeniem środowiska. Podstawy tych procesów u żywych organizmów, które podkreślali Uexküll, Piaget, Merleau-Ponty, Searle, Maturana i Varela, są pomijane w pracach nad SI. Pogląd na znaczenie ciała pozostaje zgodny z teoriami mechanistycznymi, które mechanizmy kontrolne rozumieją jako formę obliczeń. Interfejs maszyny jest postrzegany jako rodzaj

373 E. von Glasersfeld, *The Radical Constructivist View of Science...*, *op. cit.*, s. 149.

374 J. von Uexküll, *Theoretische biologie...*, *op. cit.*, s. 200.

375 Patrz. poprzedni rozdział.

urządzenia wejściowego i wyjściowego, które ma zapewnić fizyczne ugruntowanie wewnętrznych mechanizmów obliczeniowych. Tak więc biologicznie inspirowana SI łączy mechanistyczny pogląd na rolę ciała z konstruktywistycznym poglądem przemawiającym za powstawaniem wiedzy w wyniku transformacji sensomotorycznej.

Poprzedni rozdział rozważał możliwości powstania zachowań inteligentnych w sztucznych systemach, a przedstawione przykłady wykazały, że różne aspekty inteligencji są w dużej mierze obserwowane u sztucznych agentów. Można zatem stwierdzić, że współcześnie działające sztuczne systemy są w pewnym stopniu inteligentne. Aspekty inteligencji takie jak autonomia lub poczucie jaźni lub podmiotowości, możliwość samoodtwarzania nie są możliwe do zrealizowania współcześnie przez żaden sztuczny system. Jak wykazano, cechy te mogą być postrzegane jako pewien poziom w rozwoju systemu, lecz szansa na zaistnienie ich w całej okazałości (jak u zwierząt) jest wysoce wątpliwa. Jeszcze większy problem powoduje brak emocji lub odczuć w takim systemie, a które to wpływają na rozwój zachowań inteligentnych żywych organizmów, uwzględniających osobiste odniesienie do sytuacji, zadania, bądź samego działania.

Problem subiektywizmu dotyczy postrzegania sztucznych systemów jako inteligentne lub nie. Inteligencja człowieka umożliwia mu interpretację świata w elastyczny i kontekstowy sposób, który nie może być zinterpretowany przez sztuczny system, z powodu jego sensomotorycznych ograniczeń. Sztuczne systemy posiadają jednak odmienne właściwości, mimo że naśladują w pewnych aspektach działanie mózgu człowieka. Inteligencja człowieka i sztucznego systemu nie może więc być porównywana, ponieważ ich sposób działania i wyznaczone cele są odmienne. Przyjęcie perspektywy radykalnego konstruktywizmu zmusza do wyciągnięcia konkretnych wniosków dotyczących natury sztucznej inteligencji. Sztuczna inteligencja jest radykalnie odmienna i zawsze będzie się różnić od inteligencji każdego żywego organizmu. Powstania sztucznej inteligencji typu ludzkiego jest niemożliwe, tym bardziej w świetle radykalnego konstruktywizmu. Zatem sztuczna inteligencja w silnym sensie pozostaje jedynie ideą, która może znaleźć zastosowanie w fantastyce naukowej. Z drugiej jednak strony, przyjęcie subiektywistycznego poglądu na istnienie innej inteligencji niż inteligencja żywych organizmów, zmusza do uznania, że sztuczne systemy mogą być, a nawet są inteligentne.

Głównym problemem związanym z badaniami sztucznej inteligencji jest to, że dołożono stosunkowo niewiele wysiłku, aby nawiązać połączenie z innymi teoriami i

dyscyplinami, w których kwestie takie jak autonomia, reprezentacja, położenie i ucieleśnienie były dyskutowane przez długi czas, choć niekoniecznie pod tymi nazwami. Przykładem może być tu biosemiotyka lub teorie autopoezy. Nowa SI odróżnia się od swojego tradycyjnego odpowiednika pod względem badania poznania i zachowań inteligentnych jako transformacji sensomotorycznej a prace w dziedzinie robotyki adaptacyjnej są w dużej mierze zgodne z radykalnie konstruktywistycznym poglądem na powstawanie inteligencji poprzez interakcję sensomotoryczną z otoczeniem w celu osiągnięcia pewnego dopasowania.

Jednym z problemów związanym z nowymi badaniami sztucznej inteligencji jest to, że pomimo twierdzeń przeciwnych a także nacisku na ucieleśnienie i osadzenie systemów w środowisku, wielu badaczy wyznaje funkcjonalistyczną ideę, że inteligentne działanie jest w dużej mierze niezależne od materialnych cech systemu³⁷⁶. Wiele wysiłku badawczego poświęca się mechanizmom kontrolnym lub sztucznej sieci neuronowej, lecz porównanie ich z naturalnym układem nerwowym zwierzęcia, sugeruje, że związek między zachowaniem a sztucznym układem nerwowym jest niezależny od ciała. Takie podejście wskazuje, że działanie układu nerwowego jest obliczeniowe a ciało zostaje zredukowane do interfejsu sensomotorycznego systemu kontroli obliczeniowej z otoczeniem.

Zarówno Uexküll jak i Maturana i Varela twierdzili, że w żywych organizmach ciało i układ nerwowy nie oddzielnymi częściami, lecz wszystkie są zintegrowane jako składniki organizmu³⁷⁷. Podobnie w trakcie dyskusji na temat procesów znaczenia w organizmie ciało jest uważane za źródło powstawania poczucia subiektywnej rzeczywistości jako korelatu neuronowego przeciw-ciała, które powstaje i jest aktualizowane w mózgu w wyniku ciągłego przepływu informacji o objawach proprioceptywnych z mięśni, stawów i innych części kończyn³⁷⁸. Postrzeganie świata łączy w sobie zarówno działanie umysłu i układu, ponieważ tworzą strukturę przestrzenną, w obrębie której organizm się orientuje.

376 C. Emmeche, *Life as an abstract phenomenon: Is artificial life possible*, [w:] *Toward a practice of autonomous systems - Proceedings of the First European Conference on Artificial Life*, F. Varela, P. Bourguin (red.), Cambridge 1992, s. 471.

377 H.R. Maturana, F.J. Varela, *The tree of knowledge...*, op. cit., s. 34.

378 T. von Uexküll, *Introduction: Meaning and science in Jakob von Uexküll's concept of biology...*, op. cit.

W elementach fizycznych sztucznych systemów nie zachodzą procesy endosemiozy (czyli biologiczne procesy znaczenia takie jak np.: propriocepcja)³⁷⁹. Zatem nie ma żadnej integracji, komunikacji ani wzajemnego wpływu pomiędzy częściami ciała, z wyjątkiem ich czysto mechanicznych interakcji. Ponadto nie ma znaczącej integracji sztucznego układu nerwowego z ciałem fizycznym, poza tym, że niektóre części zapewniają systemowi dostęp do danych sensomotorycznych, który z kolei uruchamia ruch niektórych innych części. Tak więc nowa SI, choć uznaje rolę fizyczną, nadal w dużej mierze opiera się na w starym rozróżnieniu między sprzętem i oprogramowaniem, które było kluczowe dla tradycyjnej SI.

Podobnie krytyka Searle'a dotycząca powstania silnej sztucznej inteligencji jest wciąż aktualna. Tak jak w przypadku argumentu Chińskiego Pokoju, ciało sztucznego agenta jest zredukowane do urządzeń wejść i wyjść, które nie mają dużego wpływu na rdzeń oprogramowania mechanizmu sterującego. Działanie systemu składa się z niezależnych od fizycznego sprzętu procesów obliczeniowych w sensie funkcjonalistycznym. Zatem z punktu widzenia sztucznej inteligencji i kognitywistyki badania nad robotami adaptacyjnymi utknęły w punkcie, który można określić jako pułapkę oprogramowania³⁸⁰. Jest zatem wysoce wątpliwe, jeśli ten podział zostanie utrzymany, aby badania nad sztuczną inteligencją mogły kiedykolwiek wyjść poza ograniczenia wskazane argumentem Chińskiego Pokoju. Można argumentować, że prace nad adaptacyjną i ewolucyjną robotyką w pewnym stopniu zacierają to rozróżnienie, jednak argumenty Uexküll'a oraz Maturany i Vareli dowodzą, że nawet takie systemy pozostają allopoetyczne i heteronomiczne, bez względu na to, czy mogą się same konfigurować przy pomocy sztucznych metod ewolucyjnych. Pomimo nacisku na ucieleśnienie, umiejscowienie i autonomię, sztuczna inteligencja i robotyka adaptacyjna w obecnej formie nie jest i nie może być silną sztuczną inteligencją. Dopóki badacze SI będą głównie angażować się w rozwój technik obliczeniowych i rozwój inteligencji oprogramowania, ich modele zawsze będą znajdowały się w świetle krytyki, którą zapoczątkował Searle. Z drugiej strony, jeśli badacze SI zaczęliby rozwijać inteligencję

379 J. Hoffmeyer, *The Natural History of Intentionality. A Biosemiotic Approach*, [w:] *The Symbolic Species Evolved*, T. Schilhab, F. Stjernfelt, T. Deacon (red.), t. 6, Dordrecht 2012, ss. 97–116.

380 T. Ziemke, *Cybernetics and embodied cognition...*, *op. cit.*

sprzętowa, musieliby zbudować sprzęt autopoetyczny, czyli samokonstruujący się, czyli to, co przy obecnej technologii wydaje się niemożliwe.

W czasach Uexküll'a rozróżnienie między organizmami i mechanizmami było proste. Maszyny były przedmiotem ludzkiej konstrukcji i oddziaływały z otoczeniem w bierny i mechaniczny sposób. Natomiast organizmy zawsze samodzielnie konstruuja się zgodnie z własnym planem budowy, zarówno w sensie biologicznym, jak i psychologicznym oraz odciskają swoje działania we wszystkich interakcjach z ich fizycznym środowiskiem. Tradycyjna sztuczna inteligencja wybrała rozwiązania, które nie wymagały bezpośredniej interakcji z żadnym środowiskiem ani autonomii jakiegokolwiek rodzaju. Zamiast tego skupiła się na skonstruowanych przez człowieka formalnie zdefiniowanych systemach obliczeniowych, które można interpretować jako modele poznania i inteligentnego zachowania tylko w opinii obserwatora. Z drugiej strony nowa biologicznie inspirowana sztuczna inteligencja ma na celu zmniejszenie zależności obserwatora poprzez fizyczne umieszczenie sztucznych agentów w środowisku, tak aby mogli z nim bezpośrednio oddziaływać oraz wyposażyc ich w pewną zdolność do samoorganizacji. Ma to na celu zapewnienie sztucznemu systemowi możliwości do autonomicznego konstruowania własnego subiektywnego spojrzenia na świat i tego, co można nazwać fenomenalnym położeniem.

Pomimo tego, że zastosowanie technik uczenia maszynowego (*machine learning*) zapewniło sztucznym agentom pewne aspekty konstruktywnych procesów, które Uexküll uważał za specyficzne dla organizmów, ostatecznie te roboty pozostają heteronomiczne. Ich inteligentne i autonomiczne działanie pozostaje cechą widzianą tylko przez obserwatora. Sztuczne systemy adaptacyjne mogą okazać się przydatnymi narzędziami do badania i do modelowania procesów poznawczych, lecz pozostają słabą wersją sztucznej inteligencji lub jak określa to Ziemke – syntetycznego radykalnego konstruktywizmu³⁸¹. Zawsze będzie im brakować głęboko zakorzonego, biologicznie powstałego ucieleśnienia i osadzenia w środowisku typowym dla żywych organizmów.

381 N.E. Sharkey, T. Ziemke, *Mechanistic versus phenomenal embodiment: Can robot embodiment lead to strong AI?*, op. cit.

4.3 Sztuczne życie jako najbardziej perspektywiczny program dla rozwoju sztucznej inteligencji?

Biosemiotyka i radykalny konstruktywizm mogą przyczynić się nie tylko do rozwoju nauk poznawczych i sztucznej inteligencji, ale również do badań nad sztucznym życiem (*Artificial Life*, *ALife*, *AL*)³⁸². Ucieleśnienie sztucznych systemów jest zazwyczaj krytykowane z powodu braku ich naturalnego przystosowania do środowiska żywych organizmów. Sztuczne życie odnosi się do biologii w analogiczny sposób, jak sztuczna inteligencja odnosi się do psychologii. *ALife* proponuje umieszczenie sztucznych systemów w syntetycznym środowisku, które zapewni im adekwatną możliwość adaptacji. W programie tym podejmuje się próby zrozumienia złożonych zjawisk życia, dzięki rozłożeniu go na prostsze czynniki. *ALife* oferuje syntetyczną perspektywę: zaczyna się od prostych reguł i pojęć, następnie łączy je, aby sprawdzić, jak złożone zjawiska powstaną. Modele komputerowe opracowane dzięki tej metodologii odzwierciedlają często bardzo dokładnie obserwacje życia biologicznego. Sztuczne życie ma duży potencjał, by wyjaśnić, jak łączenie prostych reguł w symulacji komputerowej prowadzi do osiągnięcia realistycznych inteligentnych zachowań.

Niektórzy badacze tej dziedziny, jak Jean-Arcady Meyer, uważają, że sztuczne systemy w sztucznym środowisku, wykazują zachowania charakterystyczne dla naturalnych systemów życia. Dodatkowo mogą uzupełnić dziedzinę sztucznej inteligencji, poprzez próbę syntezy zachowań w sztucznych ośrodkach, które przypominają inteligentne zachowania żywych organizmów³⁸³. Program badawczy *ALife* zakłada, że w sztucznym środowisku, powstają zachowania, wywołane przez zbiór oddziałujących ze sobą bytów (będą tu nazywane animatami), których nie można sprowadzić do sumy właściwości elementów oddziałujących na niższych poziomach. Meyer twierdzi, że całość jest większa

382 Badania nad sztucznym życiem (oficjalna nazwa powstała w 1987 roku) wywodzą się z cybernetyki, teorii automatów i ogólnych badań systemów. Są w silny sposób połączone ze sztuczną inteligencją. Główna różnica polega na próbie stworzenia w wirtualnym środowisku syntetycznych systemów, które wykazują niektóre z cech biologicznych systemów ożywionych.

383 V.V. Zavertanyy, A.S. Makarenko, *Genotype dynamic for agent neuroevolution in artificial life model*, „System Research and Information Technologies”, nr 1, 2017, ss. 75–87.

niż suma jej części³⁸⁴. Metody ALife są też oddolne, sztuczne życie rozpoczyna się od stworzenia zbioru podmiotów, których zachowania są bardzo proste i zrozumiałe. Wraz z rozwojem powstają coraz bardziej złożone systemy, w których bioty wchodzą w interakcję w nieaddytywny sposób i powodują powstanie realistycznych właściwości i inteligentnych zachowań. Ponieważ podejście takie pozwala również na wysoki stopień kontroli i powtarzalność eksperymentów, może okazać się przydatnym uzupełnieniem w stosunku do biologicznie inspirowanego podejścia ucieleśnionej sztucznej inteligencji.

Podejście ALife próbuje wyjaśniać, w jaki sposób powstają zdolności żywych organizmów; od zachowań adaptacyjnych aż po skomplikowane inteligentne działania. Animaty okazują się być zdolne do aktywnego wyszukiwania przydatnych informacji i wybierania zachowań, które pozwalają im na korzystne interakcje z otoczeniem. Często są w stanie poprawić swoje zdolności adaptacyjne dzięki indywidualnemu doskonaleniu się i procesom ewolucyjnym. Metody ALife w dużej mierze opierają się na wynikach badań nad mechanizmami poznawczymi zwierząt oraz na tworzeniu narzędzi obliczeniowych inspirowanych biologicznie, *jak sieci neuronowe i algorytmy genetyczne*. Działanie animatów ma na celu powstanie adaptacyjnego i inteligentnego zachowania, pojawiającego się jako zwiększając się funkcjonalności powstające w wyniku interakcji między prostymi mechanizmami behawioralnymi³⁸⁵. Zamiast systemów robotycznych, które wykazują wąskie kompetencje na wysokim poziomie w naturalnych środowiskach, podejście ALife ma na celu modelowanie organizmów, które są proste, lecz kompletne, a ich interakcje w sztucznych środowiskach są tak realistyczne, jak to możliwe. Środowisko animatów jest dostosowane do jego możliwości i potrzeb, pozwala im na interakcje i na takie działania jak pożywanie się, rozmnażanie się, uciekanie przed drapieżnikami, polowanie i atakowanie ofiar itd.

Podstawowym twierdzeniem jest tutaj, że inteligencja animatów może powstać poprzez użycie metod oddolnych, które zakładają istnienie minimalnej architektury systemu i jego środowiska, które stopniowo zwiększają swoją złożoność. Architektura uwzględnia tylko te funkcje, które są niezbędne do osiągnięcia zdolności percepcji, kategoryzacji i autonomii. Zgadza się to ze stwierdzeniem Brooksa, który uważał, że

384 J.-A. Meyer, *Artificial life and the animat approach to artificial intelligence*, [w:] *Artificial intelligence*, Elsevier 1996, ss. 325–354.

385 L. Steels, R. Brooks, *The artificial life route to artificial intelligence: Building embodied, situated agents*, London 2018, s. 4.

powstawanie inteligentnych zachowań, języka, wiedzy eksperckiej i jej zastosowania oraz uzasadnienia, są stosunkowo proste, kiedy możliwe jest uzyskanie osadzonego w środowisku i reagującego systemu³⁸⁶. Zakłada się tutaj, że inteligencja powstaje i wzrasta w ciągu życia podmiotu, dzięki nauce i rozwojowi, zarówno indywidualnemu jak i grupowemu oraz w czasie ewolucyjnych zmian między kolejnymi pokoleniami³⁸⁷. Inteligencja w kontekście tych badań nie jest stanem statycznym, ale dynamicznym procesem i zwiększa się dzięki interakcjom ze środowiskiem. Celem inteligencji jest zapewnienie działania mechanizmów generujących korzystne zachowania, który dadzą możliwość przetrwania informacji genetycznej. Inteligencja jest nierozzerwalnie związana z życiem. Posiada cel i łatwo ją zrozumieć w świetle tego celu.

Studia sztucznego życia mogą uzupełnić badania sztucznej inteligencji. Mimo, że wiele programów SI silnie się rozwija, zakładana w nich koncepcja inteligencji jest mocno ograniczona. Dziedzina SI zajmuje się głównie izolowanymi kompetencjami systemów, wciąż ignorując fakt, że inteligentne żywe organizmy zawsze znajdują się i rozwijają swoje umiejętności w środowisku, które można nazwać kompatybilnym. Ponadto, badania głównie koncentrują się na rozwijaniu procesów algorytmicznych, w mniejszym stopniu przywiązując wagę do takich podstawowych zdolności naturalnych, jak postrzeganie, działanie i koadaptacja. Natomiast podejście ALife kładzie nacisk na relacje między animatem i jego otoczeniem, a szczególnie podkreśla zdolność animatów do przetrwania w nieoczekiwanych warunkach środowiskowych. Skupienie się na badaniu zachowań koadaptacyjnych animatów poszerza wiedzę o domenach, których badanie jest mocno zaniedbane przez SI.

Badania ALife są bliskie dziedzinie biosemiotyki, ponieważ w pewnym stopniu wraz z animatami konstruują dla nich znaczący *Umwelt*, w którym systemy mogą być obserwowane. Życie w *Umwelcie*, wraz z jego różnymi formami można nazwać złożonym, ponieważ zawiera w sobie bogaty repertuar zachowań i zależnych relacji. Inteligencję również można postrzegać jako pewien próg złożoności. Sztuczne systemy w syntetycznym środowisku, są selekcjonowane wedle zachowań. Gdy rośnie liczba zależności między nimi – rośnie również złożoność zachowań, które zależą od

386 R.A. Brooks, *Elephants don't play chess*, „Robotics and Autonomous Systems”, t. 6, nr 1–2, 1990, ss. 3–15.

387 A. Riegler, *Constructivist artificial life, and beyond*, „Proceedings of the Workshop on Autopoiesis and Perception”, Dublin, 1992, ss. 121–136.

wcześniejszego doświadczenia jak i otrzymanych bodźców. W naturze budowa organizmu jest oparta na kodzie genetycznym, a jego interpretacja, wraz z interakcjami środowiskowymi, powoduje samoorganizację organizmu³⁸⁸.

Animaty w swoim środowisku zaczynają wykazywać inteligentne zachowania, dzięki zdolności do autonomicznego konstruowania nowych zachowań przez cały okres ich istnienia. W kategoriach inżynierii komputerowej, w przeciwieństwie do tego, co robi większość systemów uczenia maszynowego, sprowadza się to do pozyskiwania nowego kodu wykonywalnego zamiast danych (parametrów, symboli, wag w sieciach neuronowych). Terry Winograd i Fernando Flores argumentowali, że gdy początkowy system nie ma struktury bezpośrednio związanej z zadaniem, które widzi jego projektant, może skonstruować nową strukturę ukształtowaną przez interakcję z jego środowiskiem. Dotyczy to systemów, których elastyczne oprogramowanie wykracza poza ograniczenia predefiniowanych programów i może podlegać sprzężeniu strukturalnemu ze środowiskiem. Sprzężenie strukturalne może wtedy kreować własne cele, które są nowością w stosunku do początkowego systemu zaprogramowanego przez człowieka³⁸⁹.

Wiele modeli ALife zostało opracowanych w celu badania życia społecznego, ekologii systemu, inteligencji roju itd. Pomimo rozwoju tych modeli, wzajemne powiązanie dynamiki genotypu i fenotypu nadal pozostają słabo zbadanym zagadnieniem. Jednak poważny problem leży w idei postrzegania istoty życia animatów. Christopher Langton, jeden z najbardziej uznanych badaczy dziedziny sztucznego życia twierdził, że można oddzielić „formę logiczną” organizmu od jej materialnej podstawy. „Żywotność” tej formy (przez co Langton rozumie zdolność do życia i rozmnażania się) miałyby być własnością formy, a nie materii³⁹⁰. Zgodnie z jego argumentami można syntetyzować prawdziwe życie w komputerze. Z biologicznej i chemicznej wiedzy wystarczy wyodrębnić logiczne zasady, które dotyczą życia. Następnie zasady muszą zostać sformalizowane, aby były realizowane przez program komputerowy. W tym kontekście wykorzystuje się metody oddolne, które charakteryzują się stosunkowo niewielką ilością prostych zasad stosowanych w modelowaniu układów żywych. Po uruchomieniu programu

388 F. Varela, *Autopoiesis and a biology of intentionality*, [w:] *Proceedings of a Workshop on Autopoiesis and Percetion*, Dublin 1992, ss. 4–14.

389 A. Riegler, *Constructivist artificial life, and beyond...*, *op. cit.*

390 C.G. Langton i in., *Artificial Life*, [w:] *Proceedings of an Interdisciplinary Workshop on the Synthesis and Simulation of Living Systems*, t. 6, New York 1989, s. 11.

sztuczne komórki prezentują, ciekawe wzory i dynamiczne formy, które wyłaniają się z interakcji między wieloma symulowanymi komórkami. Według Langtona powstanie zbiorowego zachowania jest właśnie sztucznym życiem. Perspektywę Langtona, podzielaną przez innych badaczy, można porównać z silną wersją sztucznej inteligencji, ponieważ zakłada, że życie może być badane w sposób abstrakcyjny, bez zwracania uwagi na poziomy materialne. W tym ujęciu życie jest zestawem funkcji, które można realizować za pomocą komputera.

Z perspektywy biologii komórka naturalna jednak znacznie różni się od symulowanej komórki, a animaty istnieją tylko jako obiekty matematyczne w komputerze. Jednak perspektywa ALife zakłada, że na poziomie zbioru, zachowanie wielu wzajemnie oddziaływujących naturalnych komórek w organizmie i interakcje sztucznych komórek to dwa przypadki tego samego zjawiska. W pewnym sensie to podejście może być właściwe, ponieważ możliwe jest rozszerzenie perspektywy biologii teoretycznej w taki sposób, aby jej przedmiotem było nie tylko życie takie, jakie jest, ale także życie, jakie mogłoby być – teoria subiektywnego *Umweltu* Uexküll’a pozwala na przyjęcie tej możliwości. Z powodu przypadkowych zdarzeń, życie mogłoby ewoluować nawet na Ziemi w zupełnie inny sposób. Z drugiej strony jednak to podejście ukrywa bardzo silny dualizm, a mianowicie twierdzenie, że można opisać życie funkcjonalnie, z dala od jego poziomu materialnego. Jest szczególnie ważne, aby zrozumieć ograniczenia, które wykorzystuje ewolucja biologiczna, oraz fakt, że procesy mózgu są tylko jednymi z wielu przejawów życia. Życie w żadnym wypadku nie istnieje jako czysta forma, którą można odtworzyć w sztucznym systemie. Wzory i złożone zachowania pojawiające się w obliczeniach modeli sztucznego życia mogą być równie dobrze postrzegane jako interesujące symulacje prawdziwych systemów biologicznych (co odpowiada słabej wersji SI)³⁹¹. Komputer z programem do opracowywania samoreprodukujących się danych jest bardzo przydatnym narzędziem do wdrażania formalnych teorii życia, jednak środowisko, które wytwarza dla animatów nie posiada cech własnego życia.

Oderwane od ekologii i ewolucji aspekty budowy organizmu rozważać można tylko częściowo. Powstanie obserwowanej formy organicznej odbywa się w obrębie historii i kontekstu, w którym forma i treść są należną do tej samej

391 C. Emmeche, *Life as an abstract phenomenon: Is artificial life possible...*, op. cit.

rzeczywistości³⁹². Abstrakcyjna forma nie wyjaśnia mechanizmów poznawczych ani celu organizmu tworzącego niszę dla swojego przetrwania i reprodukcji lub też przekształcania dostępnych zasobów w zorganizowaną strukturę. Konstruktivistyczne podejście badań ALife może spowodować lepsze zrozumienie ogólnych zasad biologicznych, ale nie pomoże zrozumieć różnych poziomów życia ani ich wewnętrznych połączeń. Projekt ALife można zatem skrytykować również pod względem biologicznym i semiotycznym z powodów metafizycznych. Podobnie jak silna wersja sztucznej inteligencji lub rozumienie mechanizmów poznawczych jako reguł manipulacji symbolami (obliczeniami) w głowie, projekt ten też wskazuje na koncepcje dualistyczne.

4.4 Integracja sztucznej i ludzkiej inteligencji

Sztuczna inteligencja posiada pewne zalety, które sprawiają, że niektóre z jej zastosowań przewyższają możliwości naturalnej inteligencji. Jej procesy przewyższają zwykle procesy naturalne pod względem jakościowym i ilościowym, np.: tranzystory, które są podstawowym elementem do budowy współczesnych systemów komputerowych, działają z prędkością 10 miliardów operacji na sekundę. Tymczasem kora wzrokowa człowieka nie jest w stanie uchwycić, że film na ekranie telewizora to seria nieruchomych obrazów wyświetlanych z częstotliwością 50 razy na sekundę. Maszyny są znacznie szybsze. Maszyny mają większą możliwość szybszego poprawiania błędów. Ich cyfrowe działanie zapewnia im możliwość płynnego i automatycznego pobierania ogromnych ilości informacji. Modułowa architektura sztucznych systemów pozwala na dostosowywanie ich w zależności od potrzeb i celu. Sztuczna inteligencja ma więc wiele zalet, które zmuszają do zastanawiania się, dlaczego nie radzi sobie lepiej. Nie ma możliwości samonaprawiania się, a jej elementy są drogie i delikatne. Błędy sztucznej inteligencji często powodują konieczność wyłączenia i ponownego uruchomienia komputera, co nie wydarza się w mózgach żywych istot.

Chociaż obecnie działanie modułów wejściowych i wyjściowych jest imponujące, nie zbliża to sztucznej inteligencji do jej fizycznych granic. Każdy sztuczny system może

392 C. Emmeche, *A biosemiotic note on organisms, animals, machines, cyborgs, and the quasi-autonomy of robots*, „Pragmatics & Cognition”, t. 15, nr 3, 2007, ss. 455–483.

działać szybciej, lepiej i wydajniej, podczas gdy naturalne umiejętności żywych organizmów posiadają fizyczne ograniczenia. Ciało człowieka jest delikatne i podatne na uszkodzenia, jest wrażliwe na choroby i promieniowanie, potrzebuje wody i pokarmu. Połączenie ciała ludzkiego i sztucznych inteligencji wydaje się być nieuniknione w przyszłości. Chociaż powstanie osobliwości, którą opisuje Raymond Kurzweil³⁹³ jest mocno wątpliwe, ponieważ brak jeszcze odpowiedniego podłoża inżynieryjnego, to perspektywa połączenia naturalnej i sztucznej inteligencji już się dokonuje.

Istnieje olbrzymia tendencja do zwiększania ludzkiej inteligencji. Ludzie starają się ulepszyć na różne sposoby i często są to względnie trywialne umiejętności osiągnane dzięki smartfonom i dostępowi do Internetu. Ta integracja z czasem stanie się dużo bardziej intymna. Młodszy ludzie, którzy od dzieciństwa korzystają z tych urządzeń, od samego początku są w pewnym sensie cyborgami³⁹⁴. Nauki humanistyczne starają się odkryć niezbadane dotychczas struktury ludzkiego umysłu. Inteligentne sztuczne systemy pozwalają rozszerzyć tę wiedzę. Badania psychologii poznawczej, neuroobrazowanie, badania chronometryczne i psychofizjologiczne pozwalają uzyskać coraz pełniejszy obraz działania ludzkiego systemu kognitywnego, ponieważ jego funkcjonowanie zależy od funkcji neurobiologicznych podstaw, czyli od działania mózgu. Istnieje możliwość analizy tych danych i badania mechanizmów, które powstają w trakcie przetwarzania informacji³⁹⁵.

Działania inteligentne, co starano się wykazać w pierwszym rozdziale, są warunkowane przez cielesność. Ucieleśnienie i osadzenie w świecie wpływają na to, w jaki sposób kształtuje się zachowanie. Układ sensomotoryczny pomaga w nadawaniu odpowiedniego znaczenia i ma nieprzeceniony wpływ na inteligentne działanie. Biologiczna teoria znaczenia Uexküll'a pozwala uzyskać odpowiedzi na pytania o istotę przepływu informacji między organizmem a środowiskiem. Obiekty świata są składowymi tworzonego modelu świata. Model ten zawiera w sobie to, co ma znaczenie dla podmiotu. Równie ważne w tej perspektywie jest to, że umysł i obiekty percepowane przez podmiot poznający pozostają nierozdzielne we własnym *Umwelcie*. Biologiczna teoria znaczenia i

393 R. Kurzweil, *The singularity is near: When humans transcend biology*, Penguin 2005.

394 A. Clark, *Natural-born cyborgs. Minds, technologies, and the future of human intelligence...*, *op. cit.*

395 Patrz. S. Pinker, *Jak działa umysł*, M. Koraszewska (tłum.), Książka i Wiedza 2002; J. Tooby, L. Cosmides, *The psychological foundations of culture*, „The adapted mind: Evolutionary psychology and the generation of culture”, 1995, ss. 19–136.

radikalny konstruktywizm wyjaśniają jak powstają powiązania inteligentnego działania w wyniku powiązań umysłu i świata. Każdy organizm potrzebuje środowiska i jego obiektów, które dają mu informację, jak działać. Przedstawiana przez Uexküll'a i Piageta perspektywa jest poparta wynikami badań biologii ewolucyjnej i psychologii ewolucyjnej, które wskazują, w jaki sposób czynniki zewnętrzne wpływają na modyfikacje gatunkowe i jak zmiany środowiskowe wpływają na kształtowanie się mechanizmów poznawczych³⁹⁶. Organizm pozostaje zanurzony w świecie, odbiera informacje, przetwarza je i tworzy nowe. Świat kształtują umysł organizmu, a organizm kształtuje świat i jego obiekty. Ciało przestaje pełnić funkcję graniczną, która oddziela żywą istotę od środowiska.

Piaget dowodził, że istnieją mechanizmy, za pomocą których informacje ze środowiska i pomysły jednostki oddziałują na siebie, tworząc zinternalizowane struktury w umyśle podmiotu. Zidentyfikował procesy asymilacji i akomodacji, które są kluczowe w tej interakcji, gdy ludzie budują nową wiedzę na podstawie swoich doświadczeń. Przyswajanie nowych informacji, jest włączaniem ich do istniejących ram (bez zmiany tych ram). Może więc wystąpić brak zmiany błędnego zrozumienia: ktoś nie zauważa wydarzeń, źle rozumie opinie innych lub uznaje, że wydarzenie jest przypadkiem i nie jest istotną informacją. Często doświadczenia jednostek są sprzeczne z ich wewnętrznymi reprezentacjami, zatem mogą zmienić sposób postrzegania tych doświadczeń, aby dopasować je do swoich wewnętrznych reprezentacji.

Rozwój technologiczny modyfikuje sposób, w jaki człowiek istnieje i działa. Użycie narzędzi zmienia metody funkcjonowania w środowisku, ponieważ przekształca i definiuje perspektywę postrzegania świata. Użycie technologii daje nowe możliwości zbierania doświadczeń, budowania wiedzy, kształtowania opinii. Z czasem nowe narzędzia zewnętrzne stają się integralną częścią ludzkiego systemu poznawczego. Technologia sprawiła, że status biologiczny człowieka jest uzupełniany przez zdolność do tworzenia nowych połączeń między umysłem a artefaktami. Powiązania te stały się nowymi wzorcami wzajemnych i przyczynowych zależności. Wykonanie działania zależy w dużej mierze od wcześniej przetworzonych (przynajmniej częściowo) przez narzędzie informacji. Ludzka inteligencja rozwija się na dzięki nowym technologicznym artefaktom, które dostarczają koewolucyjnej informacji zwrotnej, a tym samym powodują nowe formy adaptacji do środowiska.

396 D.M. Buss, *Psychologia ewolucyjna*, M. Orski (tłum.), Gdańskie Wydawnictwo Psychologiczne 2001.

Peirce dowodził, że umysł rozciąga się poza ludzkie ciało. Ten rozszerzony umysł to jedność ludzkiego umysłu i quasi-umysłu zewnętrznego przedmiotu, którego używa człowiek³⁹⁷. Użył w tym celu przykładu kałamarza, bez którego zapisanie myśli jest niemożliwe: „...myśli same nie przychodzą do mnie. Więc mój obszar dyskusji jest zlokalizowany także w moim kałamarzu”³⁹⁸. Możliwe jest to dzięki quasi-umysłowi narzędzia, który pozwala odnaleźć w dyskursie myśl, zinterpretować znak i przedstawić go w materialnej formie. Zapisany znak jest fizycznym i semiotycznym fundamentem myśli. Drugi umysł – umysł artefaktu pozwala na materializację zewnętrznego znaku interpretowanego przez ludzki umysł. Działanie mózgu i jego efekt są połączone i jednoczą się w procesie przetwarzania znaku.

David J. Bolter przekonuje, że produkt zmienia swojego twórcę³⁹⁹. Używanie sztucznych systemów, robotów, procesorów itd. staje się przyczyną modyfikacji interakcji człowieka ze środowiskiem. Szczególnie inteligentne systemy zapewniają łatwiejsze i bardziej efektywne funkcjonowanie. Sprawiają również, że poprzednie sposoby działania okazują się bardziej czasochłonne, wyczerpujące i przynoszą znacznie mniejszy zysk. Zmodyfikowana ludzka aktywność szybko wypiera starsze modele działania. Nowe sposoby działania są przekazywane potomkom i zapewniają ciągły rozwój, a także różne sposoby postrzegania rzeczywistości. Podczas rozważań nad możliwościami poszerzania ludzkiej inteligencji o sztuczną inteligencję, nie można odrzucić konsekwencji wynikających z rozważań Peirce’a. Umysł osoby nadzorującej działanie sztucznego systemu mieści się nie tylko w jego głowie, lecz również w quasi-umyśle zewnętrznego obiektu. Szczególnie widoczne jest to w robotyce, gdy przestrzeń myślowa jest wzmocniona, poprzez rozszerzenie ludzkiego systemu poznawczego poprzez układy sensomotoryczne robota: czujniki, mikrofony, kamery. Sztuczny układ zmysłowy poszerza możliwości ludzkich zmysłów poprzez odpowiedniki oczu, słuchu, dotyku.

Rewolucja technologiczna pozwoliła na zastosowanie sztucznej inteligencji każdemu człowiekowi. Komputery osobiste, telefony komórkowe, powszechnie

397 Skagestad, *10 Peirce's Semeiotic Model of the Mind*, [w:] *The Cambridge Companion to Peirce*, C. J. Misak (red.), Cambridge University Press 2004, s. 253.

398 C.S. Peirce, *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 7: Science and Philosophy*, Charlottesville, Va 1994, akap. 365.

399 J.D. Bolter, *Człowiek Turinga: kultura Zachodu w wieku komputera*, T. Goban-Klas (tłum.), Państwowy Instytut Wydawniczy 1990, s. 308.

dostępne automaty stały się elementami naturalnego środowiska człowieka. Komputery domowe podłączone do sieci uzyskują wszystkie potrzebne informacje w ciągu kilku sekund. To nowe wielopoziomowe medium między człowiekiem a światem zewnętrznym, mające istotny wpływ na jego postrzeganie rzeczywistości i działania. Komputery zapewniają szybkie i atrakcyjne gromadzenie wiedzy w oparciu o interakcję wideo, audio i systemu odniesień, dostarczają gotowych opinii, aktualizują i modyfikują wiedzę o świecie. Wszystko to wpływa na zachowanie przeciętnego użytkownika, narzucając mu w pewnym stopniu popularne modele zachowania. Komputery wybierają określone informacje, kształtują zachowanie i wpływ na podejmowanie decyzji przez użytkowników. To najprostszy przykład integracji naturalnej i sztucznej inteligencji. Komputer osobisty można tu zdefiniować nie tylko jako przedmiot fizyczny lub przydatne narzędzie, ale także jako symbiotyczny zewnętrzny system poznawczy. Może modelować i symulować pewne funkcje poznawcze, które w różnych warunkach wymagałyby pracochłonnej i ponadczasowej aktywności.

Andy Clark w swojej książce *Natural-born cyborgs. Minds, technologies, and the future of human* twierdzi, że ludzie od zawsze funkcjonują jako cyborgi⁴⁰⁰. Termin cyborg został użyty tam jako metafora, ponieważ Clark dowodzi, że człowiek ma naturalne predyspozycje do rozszerzenia swoich struktur poznawczych poprzez instrumenty zewnętrzne. Umysł, ciało a i środowisko technologiczne budują symbiotyczne połączenia poprzez systemy myślenia i rozumowania, łącząc biologiczne obwody mózgowe i sztuczne interfejsy. Ludzki umysł jest połączony ze światem relacją, która odnosi się do zewnętrznych obiektów i wydarzeń. Systemy biologiczne i technologiczne współpracują ze sobą, aby zminimalizować problemy związane z poznaniem.

W tym miejscu należy powrócić do wspomnianych wcześniej rusztowań kognitywnych, które odciążają układ poznawczy i minimalizują energochłonne obliczenia. Rusztowania wykorzystują stabilne regularności w środowisku, używając dostępnych i rzeczywistych przedmiotów jako najlepszych modeli świata. Dla ludzi znaczna część rusztowań składa się z instytucji kultury, języka i technologii. Notatniki, kalendarze lub zdjęcia w smartfonie stają się realną częścią pamięci

400 A. Clark, *Natural-born cyborgs. Minds, technologies, and the future of human intelligence...*, op. cit.

użytkownika. Elastyczność poznawcza systemu jest zjawiskiem naturalnym. Ta zdolność pozwala na wzmocnienie systemu poznawczego, pozwalającą na wykorzystanie repozytoriów zewnętrznych w celu szybkiego i skutecznego uzyskania informacji i wiedzy. Rusztowania kognitywne stają się elementami procesów poznawczych i wpływają na percepcję. Ich pośrednictwo ułatwia rozpoznanie wzorów, zapamiętywanie, dostosowywanie się, przewidywanie i wykonanie zadań.

W konstruktywizmie wiedza jest postrzegana jako funkcja tworzenia sensu na podstawie własnych doświadczeń⁴⁰¹. Zdobywanie wiedzy jest znacznie łatwiejsze w minimalnie ukierunkowanym środowisku, w którym podmiot poznawczy sam konstruuje własne działania w oparciu o dostępne informacje. Zastosowanie zasad promujących konstruktywistyczne poznanie pozwala na rozwiązanie autentycznych problemów⁴⁰².

Powyższe refleksje wskazują, że można próbować opisać fenomenalny świat człowieka połączonego w ten sposób ze sztuczną inteligencją. Uexküll badając relacje między żywymi organizmami a zewnętrznym światem obiektów, postulował, by próbować budować nieantropomorficzną psychologię. Dzięki niej podmiotowość organizmu ukazuje się jako zespół mechanizmów łączący go ze środowiskiem. Taka perspektywa koncentruje się na roli podmiotu, który poprzez postrzeganie i działanie konstruuje znaczenie. W przypadku połączenia świata percypowanego przez człowieka i poprzez sztuczny system można badać rozszerzony świat systemu poznawczego, ponieważ w dużym stopniu jest zależny od ludzkiego mózgu i ciała. Metody badań musiałyby skupiać się na psychologicznym działaniu czynników poznawczych związanych z urządzeniami. Ta część umysłu nie jest odizolowana ani od osoby, ani od kontekstu. Z drugiej strony jednak rozszerzony na sztuczną inteligencję umysł ludzki dociera do obiektów, które nie występują w ludzkim *Umwelcie*. Zwrócenie uwagi na sposoby wykorzystania tych nowych przedmiotów percepcji do ekspansji

401 D.H. Jonassen, *Objectivism versus constructivism: Do we need a new philosophical paradigm?*, „Educational technology research and development”, t. 39, nr 3, 1991, ss. 5–14.

402 Przykładem zastosowania konstruktywistycznego uczenia się jest nauczanie przedmiotów ścisłych, w którym studenci proszeni są o odkrywanie zasad nauki poprzez naśladowanie kroków i działań badaczy. Patrz. W.R. van Joolingen i in., *Co-Lab: research and development of an online learning environment for collaborative scientific discovery learning*, „Computers in human behavior”, t. 21, nr 4, 2005, ss. 671–688.

ludzkiego systemu poznawczego i możliwości zwiększenia potencjału człowieka w rozwijaniu inteligencji może dużo powiedzieć o samym fenomenie inteligencji.

4.5 Radykalny konstruktywizm a inteligencja

Konstruktywizm powstał się jako postmodernistyczna krytyka realistycznej epistemologii. Jego nacisk na paradygmatyczną naturę ludzkiej wiedzy jest często przedstawiany jako „nowa epistemologia”⁴⁰³, ponieważ zaprzecza istnieniu jakiegokolwiek obiektywnej prawdy. Konstruktywizm jest niewątpliwie idealistycznym i relatywistycznym sceptycyzmem, który zakłada, że nawet gdyby istniała rzeczywistość niezależna od umysłu, ludzki umysł nie byłby w stanie uzyskać do niej dostępu. Wiedza i inteligencja są zawsze subiektywne, powstają w wyniku działania podmiotu, a nie niezależnie poza umysłem obserwatora. W świetle konstruktywizmu, inteligencja jest osadzona w działaniach, kontekstach i umiejętności wykorzystania doświadczeń. Istnienie inteligencji zarówno naturalnej jak i sztucznej jest nie do pomyślenia poza subiektywnymi interpretacjami⁴⁰⁴. Zatem zachowanie inteligentne jest lokalne i posiada własny wymiar. Konstruktywiści zaprzeczają, by rozstrzyganie między konkurującymi perspektywami lub teoriami inteligencji było możliwe, ponieważ nie istnieją obiektywne naukowe oceny takich zachowań, a jedynie konstruowane społecznie konwencje, zależne zazwyczaj od widzianego przez obserwatora kontekstu i jego przekonań⁴⁰⁵. Konstruktywiści, którzy zajmowaliby się sztuczną inteligencją, zapewne przyjęliby stanowisko, że rozwój tej dziedziny powinien poprawiać życie człowieka i uwzględniać możliwość powstania SI, zamiast dążyć do obiektywnej teorii inteligencji.

Istnieje pewne podobieństwo między konstrukcjami a schematami poznawczymi w działaniach inteligentnych maszyn i żywych organizmów, szczególnie

403 G. Heath, *A constructivist attempts to talk to the field*, „International Journal of Psychotherapy”, t. 5, nr 1, 2000, s. 15; M.R. Matthews, *Constructivism in science education: A philosophical examination*, Springer Science & Business Media 1998, ss. 11–12; E. Von Glasersfeld, *Cognition, construction of knowledge, and teaching*, „Synthese”, t. 80, nr 1, 1989, ss. 121–140.

404 G. Heath, *A constructivist attempts to talk to the field...*, *op. cit.*, s. 15.

405 J. Sandoval, *Constructivism, consultee-centered consultation, and conceptual change*, „Journal of Educational and Psychological Consultation”, t. 7, nr 1, 1996, ss. 89–97.

gdy dotyczą człowieka, który buduje sztuczne systemy wedle swoich wyobrażeń o inteligencji. Mimo tego, inteligencja człowieka nie może być aplikowana w sztucznym systemie, ponieważ powstaje ona w wyniku procesu samodzielnie konstruowanego przez aktywny podmiot. Jest to proces, w którym doświadczenia powodują, że proces poznawczy podmiotu i związane z nim działania dostosowują się do wymagań jego fenomenalnego świata i samoorganizują się. Ponieważ każdy podmiot poznawczy sam konstruuje swój własny świat, to ustalenie, czy jego zachowania inteligentne są porównywalne, jest niemożliwe do zbadania. Ludzie mogą rozpoznać tylko własne reprezentacje mentalne i struktury poznawcze, lecz nie są w stanie uzyskać dostępu do jakiegokolwiek rzeczywistości niezależnej od ich własnego umysłu⁴⁰⁶.

Inteligentne działanie jest wynikiem aktywności samoorganizującego się podmiotu, jest poddawane selekcji i ewoluuje. Zachowania systemu są dopasowywane do kontekstu sytuacji i środowiska. Jeśli stare reguły i zasady działania nie są w stanie zasymilować nowych doświadczeń, stają się zbędne i zostają zmienione. Dopasowanie zachowania wynika zarówno z interakcji między podmiotem a światem zewnętrznym, jak również z korespondencji pojęć i doświadczeń⁴⁰⁷. Inteligencja powstaje więc jako konstrukt podmiotu, nie może zostać przekazana ani oceniona przez żaden obiektywny standard. Programista sztucznej inteligencji może być postrzegany jako uczestnik początkujący proces. Obiektywność nie jest możliwa, ponieważ inteligencja jest dyskursywna – można ją określić jako zestaw odpowiednich i opłacalnych zachowań, lecz określenie prawdziwa mija się z celem. Inteligentny, autonomiczny i adaptacyjny sztuczny system musi posiadać swój własny sposób postrzegania świata, który nie jest epistemicznie słabszy lub lepszy niż inne⁴⁰⁸.

406 C.W. Joldersma, *Ernst von Glasersfeld's radical constructivism and truth as disclosure*, „Educational Theory”, t. 61, nr 3, 2011, ss. 275–293.

407 *Ibid.*

408 M.R. Matthews, *Introductory comments on philosophy and constructivism in science education*, „Science & Education”, t. 6, nr 1–2, 1997, ss. 5–14.

4.6 Czy sztuczna inteligencja powstaje?

Badania Uexküll'a jak i teoria konstruktywistyczna zmuszają do zrozumienia fenomenalnego świata podmiotu jako czegoś, co obejmuje więcej niż znaki zapośredniczone przez język, czyli znaki, które pojawiają się zarówno wewnątrz, jak i na zewnątrz podmiotu i pozostają związane z bogactwem ikonicznych, indeksowych i innych procesów semiotycznych w naturze. Krytyka SI pozostaje osadzona w czystym, subiektywistycznym humanizmie. Nie da się sprowadzić logiki do psychologii lub do biologii (lub odwrotnie). Filozofia jednak potrafi przeprowadzić podział na subiektywne i obiektywne koncepcje informacji, co swoją drogą odzwierciedla dualizm przenikający ludzką kulturę. Aby wyjść poza ten dualizm należy zaakceptować znaki i symbole, jako informacje, które nie są do końca mentalne ani materialne. Jeśli zaś chodzi o znaczenie zawarte w *Umweltie*, nie można jednoznacznie stwierdzić, że sztuczne systemy go nie posiadają. Orzeczenie takie byłoby niczym innym jak subiektywną opinią obserwatora, który uznaje swoje postrzeganie rzeczywistości jako obiektywną formę percepcji. Tymczasem nie mamy pewności, czy w strukturach obliczeniowych nie znajduje się coś więcej, jakaś forma znaczenia, której nie jesteśmy w stanie dostrzec albo zrozumieć.

Należy zwrócić uwagę na dwie rzeczy. Po pierwsze, pojęcie symbolu w dziedzinie sztucznej inteligencji wydaje się zbyt wąskie i zbyt dynamiczne. Symbole są postrzegane jako zwykłe znaki fizyczne, których zdolność do reprezentacji wynika wyłącznie z działania umysłu człowieka: programisty, architekta systemu lub użytkownika. Jednak symbole działają kompleksowo w całym systemie i są interpretowane przez inne struktury symboliczne – przez żywe organizmy a także sztuczne systemy, które są nie tylko strukturami czysto formalnymi, ale każda z nich posiada własną specyficzną materialność i które w końcu mogą nauczyć się reprezentować swoje ja i swój stosunek do świata zewnętrznego za pomocą symbolicznych lub niesymbolicznych relacji znakowych. Nie można z góry wykluczyć, że triadyczność semiozy w sztucznym systemie nie przebiega w pełnym znaczeniu, ponieważ nie posiadamy pełnego dostępu do *Umweltu* maszyny.

Wydaje się, że rozwiązaniem pozwalającym na uznanie inteligencji maszynowej, mogłoby być powstanie precyzyjnej, wspólnej dla artefaktów i organizmów definicji inteligencji. O ile możliwe byłoby stworzenie nowego terminu,

możliwa byłaby dyskusja o podobieństwach i różnicach umiejętności tych odmiennych systemów. Do tej pory możemy jednak mówić tylko i wyłącznie o inteligencji maszyn, wpisując ten termin w cudzysłów, lub idąc za przykładem Peirce'a, nazywać ją quasi-inteligencją. Inteligencja jest pojęciem niejednorodnym i składa się na nią zbyt wiele aspektów, by uznać bezsprzecznie, że roboty, systemy decyzyjne a nawet zwierzęta są inteligentne. Z tego też powodu pojęcie inteligencji powinno być całkowicie oddzielone od pojęcia świadomości, ponieważ istnieje możliwość, że ta właściwość przynależy do niewielkiej klasy istot biologicznych wyposażonych w centralny układ nerwowy i wysoko rozwiniętą korę mózgową. Nie można ponad wszelką wątpliwość stwierdzić, że prymitywne zwierzęta posiadają jakąś protoświadomość. Przedmiotem dyskusji pozostaje, czy wysoka inteligencja wymaga świadomości, czy też wszystkie inteligentne zadania mogą być w zasadzie wypełniane przez nieświadome artefakty⁴⁰⁹. Rozsądne wydaje się stosowanie pojęcia inteligencji w związku ze sposobem funkcjonowania w *Umwelcie*. Z tego powodu biosemiotyczne pojęcia znaku, znaczenia itp. należy badać w odniesieniu do funkcji i formy. Tym bardziej, że oczywiste umiejętności wykazywane przez gatunki biologiczne w radzeniu sobie z ich *Umweltem* sprawiają, że pojęcie wyższego lub niższego stopnia inteligencji jest *a priori* koncepcją biosemiotyki.

Stworzenie silnej wersji sztucznej inteligencji wydaje się w tym momencie niemożliwe. Oczywiście chodzi o inteligencję, która dorównywałaby lub przewyższała inteligencję człowieka. Jest to niemożliwe z wielu względów, które zostały opisane w tej pracy. Nie ma wątpliwości, że sztuczna inteligencja w słabej wersji istnieje i często jest nią to, co kilka dekad temu było uważane za ogólną SI, a już dziś nie jest tak definiowane. Słabą (lub wąską) sztuczną inteligencję są technologie służące np. do rozpoznawania obrazu i mowy, chatboty, samochody samokierujące, programy eksperckie. Technologia wciąż rozwija nasze definicje i oczekiwania wobec inteligentnych systemów. Słaba sztuczna inteligencja oznacza jednak zastosowanie, które mogą być użyte do wykonywania określonych zadań dzięki prostemu programowaniu i regułom. Przykładem użycia słabej wersji SI w życiu codziennym jest wspomaganie głosowe w telefonach. Tu zauważyć można również różnicę między nadzorowanym a nienadzorowanym programowaniem, ponieważ wspomaganie

409 D.J. Chalmers, S. Wolf, *The conscious mind: in search of a fundamental theory*, „Integrative Physiological and Behavioral Science”, t. 32, nr 3, 1997, ss. 291–292.

głosowe zwykle ma zaprogramowaną odpowiedź. To, co robi sztuczna inteligencja w smartfonach, to przeszukuje repozytoria danych w poszukiwaniu informacji podobnych do tych, które już posiada i odpowiednio je klasyfikuje. Jest to cecha inteligencji podobnej do inteligencji człowieka, ale na tym kończą się podobieństwa, ponieważ słaba SI jest po prostu naśladownictwem pewnych elementarnych umiejętności. Reaguje jednak tylko na to, co zostało zaprogramowane. Nie rozumie ani nie wywodzi znaczenia z posiadanych informacji.

Należy przyjąć możliwość, że może jednak powstać inteligencja innego rodzaju – inteligencja maszynowa. Można stwierdzić nawet, że już istnieje, a brak akceptacji tego stanu rzeczy wynika stąd, że obarczeni jesteśmy antropocentrycznym spojrzeniem na odmienność innych systemów niż organiczne, a w szczególności ludzkie. Dodatkowo przymiotnik sztuczna wprowadza w błąd, ponieważ zmusza teoretyków i praktyków do szukania porównań z naturalną inteligencją. Jest wysoce możliwe, że sztuczne systemy wytwarzają własny *Umwelt*, jednak tak znacząco różni się od ludzkiego *Umweltu*, że nie jesteśmy w stanie zauważyć jego przejawów. Na poziomie przetwarzania ciągów liczbowych, może dziać się coś, czego nie jesteśmy w stanie zinterpretować inaczej niż jako nieprzezroczystość epistemiczną. Niestety nikt w tym momencie nie jest w stanie udzielić odpowiedzi, czy w tej nieprzezroczystości, często uważanej za szum, coś się sobie kryje i czy może być emanacją sztucznego fenomenalnego świata. Odpowiedź na to pytanie wymaga zupełnie odmiennych badań pozwalających lepiej zrozumieć czym mógłby być i jak się przejawia *Umwelt* sztucznego systemu. Z pomocą dla zrozumienia tego problemu przyjść może człowiekowi połączenie naturalnej i sztucznej inteligencji, ponieważ hybryda takich możliwości tych odmiennych systemów będzie miała ułatwiony dostęp do obu światów, a dzięki temu transkrypcja znaczenia będzie znacznie ułatwiona.

Równie ważną rolę w dziedzinie sztucznej inteligencji może odegrać radykalny konstruktywizm, który spina klamrą biosemiotyczne podejście do dziedziny SI, teorii obliczeniowej, kształtowania się zdolności poznawczych sztucznych systemów oraz efekty połączenia odmiennych rodzajów inteligencji. Metody tego podejścia mogą ocenić filozoficzne i koncepcyjne znaczenie podejścia obliczeniowego, np. czy możliwe jest sformułowanie modeli obliczeniowych procesów poznawczych i czy modele obliczeniowe mogą kiedykolwiek przyczynić się do rozwoju sztucznej inteligencji. W szczególności jednak konstruktywizm może zbadać narzędzie, jakim są

modele komputerowe, by odpowiedzieć na pytanie, czy ich zastosowanie może pomóc sztucznym systemom w inteligentnym konstruowaniu rzeczywistości.

Zakończenie

Praca ta miała na celu zbadanie możliwości użycia biosemiotycznych koncepcji inteligencji w rozwoju inteligencji maszyn. Przyglądając się badaniom SI można bez wątplenia stwierdzić, że borykają się one z problemami, które nie mogą zostać rozwiązane jedynie przy użyciu dostępnych technik obliczeniowych. Biosemiotyczne podejście do procesu kształtowania się inteligentnych zachowań przedstawia takie własności inteligencji, które nie były dotychczas analizowane w ani w badaniach sztucznej inteligencji, ani nawet w badaniach dotyczących ogólnej inteligencji. Wyniki rozważań zawartych w tej pracy pokazują, że zastosowanie koncepcji biosemiotyki może wskazać nie tylko nowe kierunki rozwoju badań nad sztuczną inteligencją, lecz przede wszystkim umożliwia spojrzenie na procesy powstawania inteligentnych zachowań w zupełnie nowym świetle. Nowatorski opis semantyki, związany z biosemiotyką, której badania opierają się na aktach indykatywnych, może być bardzo cenny dla orientacyjnego podejścia SI.

W pierwszym rozdziale pracy ukazano problem pojęcia inteligencji w kontekście pytania Turinga o możliwość stworzenia myślących maszyn. Okazuje się, że jednym z poważnych wyzwań dla osób zaangażowanych w rozwój SI jest to, że nie istnieje standardowa definicja tego czym jest inteligencja. W wyniku analizy wskazano, że badania w tej dziedzinie kierują się różnorodnymi pojęciami. Warto zatem podkreślić, że w przypadku sztucznej inteligencji, nie mamy do czynienia z jednym przedmiotem badań, lecz zbiorem praktyk i cech, które ludzie uważają za inteligentne. Większość badaczy zgadza się, że podstawowym celem sztucznej inteligencji jest umożliwienie maszynom poznania, postrzegania i podejmowania decyzji, które wcześniej posiadali tylko ludzie lub inne zwierzęta. Ponieważ nie do końca wiemy, jak jednoznacznie zdefiniować inteligencję biologiczną, próba jej sztucznego naśladowania jawi się jako źle określone zadanie. Warto przypomnieć, że rozważania dotyczyły silnej wersji SI, która nieodmiennie stanowi niedościgniony cel badań nad inteligencją sztucznych systemów, w założeniu mających przypominać inteligencję typu ludzkiego.

Celem pierwszego rozdziału było również przedstawienie historycznego zarysu rozwoju tej dyscypliny, ze szczególnym uwzględnieniem jej ograniczeń, aby pokazać jej teoretyczne korzenie i problemy tradycyjnego podejścia obliczeniowego. Z

naukowego i praktycznego punktu widzenia, SI jako dziedzina zajmująca się konstruowaniem modeli naturalnego poznania i inteligencji osiąga duże sukcesy. Mimo tego, symboliczne podejście ukazało istotne problemy zastosowania sztucznych systemów w rozwiązywaniu specyficznych problemów związanych z semantyką leżącą u podstaw inteligentnych zachowań. Modele obliczeniowe ignorowały ograniczenia systemów symbolicznych, takich jak problem semantycznego ugruntowania symboli a często również problem fizycznego ucieleśnienia i usytuowania systemów. W latach osiemdziesiątych XX wieku starano się rozwiązać te problemy, szczególnie na polu robotyki, niestety z niewystarczającym skutkiem. W nowym podejściu koncepcja ucieleśnienia leży u podstaw modelowania inteligentnych systemów. Pomimo tego, ucieleśniona sztuczna inteligencja była zbyt wąsko pojmowana, ponieważ każda perspektywa ucieleśnienia i semantycznego zaangażowania sztucznych systemów była rozważana tylko częściowo. W pierwszym rozdziale poruszony został również problem antropocentrycznej perspektywy przyjętej w badaniach SI. Wewnętrzna dynamika badań nad SI musi ostatecznie doprowadzić do zmiany tego nastawienia. Badacze kładą coraz większy nacisk na doświadczenie, ugruntowaną wiedzę i inteligencję usytuowaną. W efekcie zmusza do skorygowania podstawowych założeń i intuicji na temat natury inteligencji w kierunku nieantropocentrycznym.

W związku z powyższymi problemami, w drugim rozdziale próbowano zbadać istotę inteligencji w świetle teorii biosemiotycznych. Teoretyczne ramy systemowe inspirowane biosemiotyką mogą zostać wykorzystane do badania użycia znaków przez sztucznych agentów. Dzięki wykorzystaniu etologicznego podejścia do analizy interakcji między zwierzęciem a środowiskiem, starano się przedstawić, w jaki sposób rozwiązywane są analogiczne semantyczne problemy pojawiające się w świecie zwierząt. Dlatego też zaproponowano skupienie się na rozważaniu procesu przetwarzania znaków – semiozy, która prowadzi do powstania pierwszorzędnej semantyki, gdzie znaczenie jest definiowane jako wynik przetwarzania znaków. W szczególności skupiono się na pracach Jakoba von Uexküll'a, którego teoria znaczenia szczególnie wpłynęła na powstanie tej pracy. Pojęcie *Umweltu* – fenomenalnego świata podmiotu poznającego, może okazać się szczególnie przydatne dla współczesnych dyskusji toczących się w ramach sztucznej inteligencji, sztucznego życia i badania sztucznych systemów autonomicznych. Biosemiotycy jak i większość

biologów odczuwają niechęć do antropocentrycznych i antropomorficznych wyjaśnień. Uexküll wpłynął na opracowanie biologicznych podstaw badania zachowań zwierząt i wysunął teorię na temat interakcji organizm-środowisko w kategoriach subiektywnych światów percepcyjnych i efektorowych, a tym samym zaprzeczał wyjaśnieniom antropomorficznym i czysto mechanistycznym. Jego celem było znalezienie sposobu wyjaśnienia na płaszczyźnie biologicznej zachowania organizmów i ich osadzenia w środowisku.

Wskazano również, że zaproponowany przez Uexkülla model kręgu funkcjonalnego zawiera wszystkie elementy, które są częścią procesu przetwarzania znaków. Przedstawiono w związku z tym koncepcję semiozy, stworzonej przez Charlesa Peirce'a, lecz prezentowanej w perspektywie biologii teoretycznej. Interpretacja działania kręgu funkcjonalnego wskazuje na istnienie biologicznych procesów semiozy. Organizm jest podmiotem, który otrzymuje i tłumaczy pewne sygnały środowiskowe, które odgrywają rolę znaków, a tym samym determinują biologiczny stan organizmu oraz jego zachowanie. Organizmy są osadzone w swoim środowisku za pomocą kręgów funkcjonalnych, a ich inteligentne zachowanie postrzega się jako wynik ciągłej interakcji między nimi a środowiskiem w celu wywołania skutecznego zachowania. Znaki zaś są postrzegane jako odgrywające funkcjonalną rolę w tej interakcji. Sztuczne systemy symboliczne można zatem spróbować powiązać z semantyką Uexkülla, w której znaki i reprezentacje są postrzegane jako osadzone w kręgach funkcjonalnych, według których zorganizowane jest oddziaływanie organizmów i środowiska.

Z działania kręgów funkcjonalnych wynika możliwość autonomii podmiotu. W kontekście teorii *Umweltu* Uexkülla każdy podmiot może być autonomiczny, o ile posiada znaczenie i używa go wchodząc w interakcje z postrzeganym *Umweltem*. Wśród najważniejszych cech autonomicznego systemu znajduje się zdolność do postrzegania środowiska przez narządy zmysłowe i działanie na niego poprzez narządy efektorowe, posiadanie własnego postrzegania wykonywania zadań, interpretacji percepcji, wchodzenia w interakcje z innymi agentami oraz umiejętność ukierunkowania zachowania na osiągnięcie celu, poprzez wykazywanie inicjatywy. Aspekt percepcyjny i operacyjny autonomii daje się opisać w kategoriach kręgu funkcjonalnego Uexkülla, co ułatwić może konstruowanie modelu procesu tworzenia znaczeń i procesu decyzyjnego u sztucznych agentów. Dużą rolę w powstawaniu

autonomicznych i inteligentnych zachowań sprawują zewnętrzne rusztowania poznawcze, pełniące funkcję narzędzi bądź zewnętrznych repozytoriów danych. Semioza stanowi aktywny proces poszukiwania znaczeń, który często skutkuje powstaniem względnie statycznych lub nawet stałych struktur, które same w sobie zawierają znaczenie. Pozwala ona zanalizować dalsze zachowania i ukształtować nowe nawyki. Rusztowania semiotyczne nie tylko ułatwiają uzyskanie stabilności poznawczej, lecz także zapewniają przestrzenne elementy, na których podmiot może oprzeć swoją semantyczną topologię systemu znaków. Rusztowania umożliwiają użytkownikom symboli inteligentne omawianie abstrakcyjnych, kontrfaktycznych znaczeń a tym samym samo rusztowanie może być przekształćane w inne, ułatwiając w ten sposób proces, w którym inteligencja się rozwija.

W drugim rozdziale starano się również trudne pojęcia filozoficzne, jak podmiotowość i intencjonalność, które odgrywają nieprzecenioną rolę w debacie poświęconej inteligencji. Perspektywa biosemiotyczna ukazuje podejście, które opiera się na relacji między podmiotem a obiektem. Starano się wyjaśnić warunki stosowania terminów jak podmiot i podmiotowość zarówno w teorii *Umweltu* jak i we współczesnym dyskursie biosemiotycznym. Podejście biosemiotyczne stara się zniwelować ontologiczną lukę między człowiekiem a zwierzęciem poprzez przekształcenie pojęcia podmiotu. Podmiot w tej perspektywie jest wynikiem procesów semiotycznych, posiadającym pierwszoosobową perspektywę życia empirycznego w pojęciu podmiotowości. W odniesieniu do filozoficznych rozważań opartych na odwołaniach do jaźni, biosemiotyczne podejście zawarte w tej pracy ma wielką zaletę, ponieważ stanowi alternatywę dla dualistycznych ontologii opartych na niemożliwym do pokonania podziale między umysłem a materią, który wywodzi się z filozofii Kartezjusza. Umożliwia to ponowne rozważenie problemu intencjonalności, odwołując się do relacyjnego charakteru procesu przetwarzania znaków. W zakresie filozofii analitycznej intencjonalność jest przedmiotem teorii, które oferują wyższy poziom samoświadomości, natomiast biosemiotyczne podejście nie kontynuuje tego podejścia i nie odnosi się do „trudnego problemu świadomości”. Intencjonalność jest zawarta w semiozie. Wychodząc z semiotycznego realizmu Peirce’a, który ugruntował intencjonalność w uogólnionym działaniu semiozy, triadyczność działania znaków wskazuje na intencjonalność, ponieważ dla podmiotu znak jest „o czymś”. Tłumacz znaku i interpretator nie muszą być z konieczności istotą ludzką. Stany wewnętrzne

systemu i jego środowisko są ze sobą powiązane przez tworzenie znaczenia łączącego go ze światem zewnętrznym poprzez akt interpretacyjny.

W rozdziale trzecim podjęto pytanie, czy sztuczny system może być aktywnym uczestnikiem semiozy i użytkownikiem własnego fenomenalnego *Umweltu*. Przedstawiono zatem głębsze zrozumienie prawdziwie ucieleśnionych i usytuowanych systemów autonomicznych jako złożonych samoorganizujących się agentów semiotycznych, charakteryzujących się wyłaniającymi się właściwościami jakościowymi. W tej części analizowano pojęcia autopoezy i autonomii, które są związane z jakościowym aspektem organizmu. Wskazano, że pojęcie symbolu wewnątrz sztucznej inteligencji jest zbyt wąskie i zbyt diadyczne. Symbole są traktowane jak zwykłe znaki fizyczne, których zdolność do reprezentowania wynika wyłącznie z odniesień ustalonych przez programistę. Należy jednak zauważyć, że symbole interpretowane przez inne struktury symboliczne w maszynach, mogą przedstawiać swój własny stosunek do świata zewnętrznego za pomocą symbolicznych lub nie-symbolicznych relacji znakowych. Nie można wykluczyć, że znaki w procesach przetwarzania symboli w sztucznych systemach są interpretowane jako znaczące w pełnym tego słowa znaczeniu. Istnieje jednak poważny problem z pierwotną semantyką, którego nie można rozwiązać wyłącznie poprzez formalizację, ponieważ wykracza ona poza opis symboliczny. Sztuczny system opisywany jest zazwyczaj zwykle poprzez język formalny, można go więc pod pewnymi względami postrzegać jako uproszczoną, diadyczną wersję ludzkiej semantyki. W niektórych przypadkach język symboliczny dostarcza jednak dowodów na istnienie logicznej formy zdań, które są wywnioskowane jedynie pośrednio w języku człowieka. System symboliczny może nie tylko realizować zmienne logiczne, ale pozwala przyjrzeć się rzeczywistości zmiennych, mechanizmów oraz jego dynamice. Niewątpliwie, semantyka sztucznego systemu symbolicznego posiada bogate elementy ikoniczne, które rzadko są brane pod uwagę w formalnych badaniach języka mówionego.

Pragmatycznym problemem związanym z ucieleśnioną sztuczną inteligencją, jest to, że wielu badaczy SI odrzuca ideę, że inteligencję i poznanie można wyjaśnić w kategoriach czysto obliczeniowych, lecz nie przedstawia żadnego alternatywnego wyjaśnienia. Stwierdzenie, że inteligencja wymaga fizycznej implementacji i ucieleśnienia prowadzi do stwierdzeń, że bycie fizycznym może co najwyżej być niezbędnym warunkiem inteligencji. Ucieleśniona SI dotyczy więc robotycznych

modeli poznania, co pozwala nam odróżnić podejście ucieleśnionej sztucznej inteligencji od jej tradycyjnego odpowiednika. Użycie fizycznie ucieleśnionych robotów nie jest jednak samo w sobie, wyróżniającą cechą ucieleśnionej sztucznej inteligencji, gdyż zastosowane systemy sztucznej inteligencji mogą bazować na tradycyjnym podejściu symbolicznym. Stąd też pojawia się propozycja, by uwzględnić stopniowalność ucieleśnienia. W takim wypadku, ucieleśniona sztuczna inteligencja nie posiada spójnej definicji ani specyficznych ram teoretycznych. W praktyce naukowej pozostaje pluralistycznym obszarem badań, odnosząc się w pewnym stopniu zarówno do obliczeniowych, jak i fizycznych modeli.

W dalszej części trzeciego rozdziału badano możliwości uzyskania autonomii sztucznych systemów, która jest postrzegana jako cecha wyższego poziomu, wynikająca z ucieleśnienia, samorozwoju i adaptacji sztucznego systemu. Dzięki przyjęciu perspektywy zakładającej istnienie poziomów autonomii, można badać strukturę systemu, który może zrealizować autonomiczne zachowanie, nie rozważając sposobu jego powstania i rozwoju. Uzasadniono, że w dyskusjach na temat autonomii, użyteczne może być dokonanie kantowskiego rozróżnienia między autonomią jako właściwością rzeczy samych w sobie a autonomią jako właściwością, którą ludzie obserwatorzy, przypisują rzeczom, które uważają za autonomiczne. Do pewnego stopnia to rozróżnienie pozwala określić autonomię jako domenę, w której istnieją różne systemy. Wydaje się, że nawet wśród zwolenników teorii autopoezy istnieje pewien spór co do tego, czy autonomia biologiczna jest rzeczywiście koniecznym wymogiem powstawania zachowań inteligentnych.

W rozdziale trzecim przedstawione zostały również modele, których zadaniem było opisanie możliwości aplikacji pojęć biosemiotycznych w SI. Przedstawiono rozwiązania na polu architektury sztucznych systemów oraz możliwości rozbudowania dotychczasowych technologii SI o właściwości takie jak umiejętności adaptacyjne, rozpoznanie semiotyczne i modelowanie intencjonalności agentów. Starano się również zaprezentować modele rozwiązania problemu ugruntowania symboli oraz możliwości użycia i działania rusztowań poznawczych. Część z tych badań została już przetestowana i zweryfikowana, inne zaś powstały, by wskazać, że biosemiotycznie inspirowana SI jest możliwa do wykonania. Wielu badaczy sztucznej inteligencji zgadza się, że układ nerwowy żywego organizmu może być modelem obliczeniowym, którego sztuczny odpowiednik odwzorowuje wejścia sensoryczne na wyjścia silników

w sztucznych systemach. Mimo tego, oddolne biosemiotyczne podejście mocno dystansuje się od symbolicznego reprezentacyjnego podejścia i skupia się na interaktywności, widzianej jako struktura kształtująca zachowania inteligentne. Sztuczny układ nerwowy w połączeniu z technikami uczenia się i metodami ewolucyjnymi jest wykorzystywany w próbach tworzenia sztucznych organizmów, umożliwiając im samoorganizację procesów semiotycznych. Kilka omówionych przykładów pokazało, w jaki sposób takie techniki pozwalają robotom znaleźć własny sposób zorganizowania ich kręgów funkcjonalnych, czyli wewnętrznego sposobu wykorzystania znaków i odpowiedzi na bodźce pochodzące ze środowiska. Wykorzystanie takiej samoorganizacji czyni sztuczne systemy wyjątkowym rodzajem mechanizmu, który powinny stać się przedmiotem badań, jako systemy semiotyczne.

Podejście do autonomicznych agentów i koncepcja ucieleśnionego poznania, są w dużej mierze kompatybilne, lecz istnieje zbyt mało wspólnych opracowań, które mogłyby wnieść znaczący wkład w rozwój SI. Sztuczne systemy, mimo implementacji biologicznych inspiracji, są wciąż dalekie od posiadania inteligencji żywych systemów. Wynika to głównie z faktu, że sztuczne systemy składają się głównie z części mechanicznych i programów sterujących, które nie powstały w wyniku samorozwoju ani ewolucji w celu naturalnego przystosowania się do środowiska. W prezentowanych badaniach, starano się odsunąć zarzut, że sztucznym systemom z konieczności brakuje semantyki, wewnętrznego znaczenia i własnych celów. Dlatego należałoby coraz bardziej minimalizować rolę ludzkiego programisty w tworzeniu semantyki sztucznych systemów i badać alternatywne podejścia.

Czwarty rozdział przedstawia potencjał wykorzystania radykalnego konstruktywizmu jako paradygmatu rozwoju badań nad SI. W tej części wyjaśniono czym jest konstruktywizm w SI i dlaczego wydaje się być najlepszym podejściem. Stanowisko konstruktywistyczne stara się ograniczyć dostęp sztucznego systemu do wiedzy o świecie i możliwych działaniach. System powinien mieć możliwość introspekcji własnej wiedzy, nie tylko stosować reguły, ale być w stanie wyszukiwać je i manipulować nimi jako własnymi przedmiotami. Takie podejście pozwala sztucznemu systemowi na własne rozpoczęcie nauki, zdobywanie doświadczeń i gromadzenie wiedzy. Dzięki zastosowaniu konstruktywistycznej metodologii badawczej, możliwe jest modelowanie ścieżki rozwoju systemu jako cel sam w sobie,

w miejsce osiągania pewnych szczególnych kompetencji. Wydaje się, że takie podejście może doprowadzić do powstania sztucznych systemów, które są bardziej biegłe w zadaniach ogólnych, w porównaniu z klasycznymi systemami SI, które osiągają kompetencje w ograniczonych zadaniach.

Przedstawionych zostało kilka możliwych ścieżek rozwoju, które mogłyby znacząco wpłynąć na rozwój SI. Między innymi przedstawiono możliwość powstania sztucznych *Umweltów*. Takie podejście podąża za konstrukcją przestrzeni działania sztucznych systemów przez analogię i znajdowało już swoje częściowe realizacje w pracach Rodneya Brooksa. Mimo, że podejście radykalnego konstruktywizmu wydaje się być pracochłonnym procesem, ponieważ oczekiwany jest samorozwój systemu, w miejsce stosowania gotowych reguł i procedur, musi ono przynieść pewne korzyści, które należy uznać za wartościowe. Modelowanie sztucznej inteligencji w sposób konstruktywistyczny może okazać się bardzo przydatne dla sztucznych systemów, które muszą działać w rzeczywistym świecie. Mimo, że postrzeganie i działanie sztucznego systemu w świecie wydaje się być w takim przypadku zbudowane z bardziej prymitywnych mechanizmów, pozwala jednak na autonomiczny rozwój inteligentnych zachowań.

Szczególną rolę odegrać może program badawczy ALife, gdzie w symulowanym obliczeniowo środowisku można uzyskać cechy agentów trudne lub niemożliwe do uzyskania w typowych systemach robotycznych. Dzięki temu możliwa powinna być implementacja funkcji semiotycznych w tego typu systemach. Dzięki tworzeniu sztucznych agentów w sztucznych środowiskach można rozpocząć badanie procesów semiotycznych od najniższego możliwego poziomu, przez co unika się wielu założeń *a priori* i podejścia antropocentrycznego. Ze względu na charakter badań ALife, w przypadku odpowiedniej implementacji, możemy testować uzupełniające się związki między epistemologią a metodologią. W modelach ALife znaczenia tworzone w obrębie sztucznego systemu mogą zostać wzbogacone o różnorakie konteksty, m.in. ugruntowanie społeczne czy środowiskowe. Celem modeli ALife jest wykazanie za pomocą symulacji, w jaki sposób wyższe struktury poznawcze mogą się wyłonić z tworzenia zasad i nawyków przez proste istoty sensomotoryczne. Zaletą systemów ALife w badaniach nad inteligencją jest możliwość wprowadzenia realistycznych mechanizmów uczenia się i ewolucji odpowiadających cechom żywych organizmów. Celem konstruktywistycznego modelu ALife może być programowanie robota, które

ma tę samą metodologię pracy, np. wybór odpowiednich schematów w celu reagowania na niektóre bodźce ze swojego otoczenia. Wszystkie schematy są produktem „*quasi*-piagetowskiego” mechanizmu uczenia się, który jest analogiczny do rozwoju zwierząt w ich okresie sensomotorycznym.

Kolejną ścieżką wyznaczoną przez radykalny konstrukttywizm może stać się rozwój inteligentnych systemów wspierających inteligencję człowieka. Taka perspektywa pozwala przyjrzeć się połączeniu ludzkiej i sztucznej inteligencji poprzez mechanizmy wspierające reprezentację wiedzy, rozumowanie i podejmowanie decyzji w inteligentnych systemach. Inteligencja zostaje wzmocniona poprzez kontekst nabywania wiedzy; uczenie się wymaga aktywnej interakcji i ma większe znaczenie jako proces niż produkt. Modele łączenia inteligencji odmiennych systemów pozwalają na rozwój działań inteligentnych, poprzez integrację nowych aspektów doświadczeń z istniejącą uprzednio wiedzą. Wchodzenie w interakcje w konkretnych sytuacjach, łączenie informacji z kilku źródeł, w tym z poprzednich doświadczeń, muszą zostać opracowane i zintegrowane przez agenta w znaczący sposób, aby nowe informacje, które są generowane i interpretowane, mogły być powiązane z istniejącą wiedzą, która może również zostać ponownie zinterpretowana w świetle nowego doświadczenia. Ostatecznie prowadzi to do konstruktywnego procesu, w którym generowana jest nowa wiedza i związana z elementami poprzednich doświadczeń.

Mimo iż biosemiotyka zastosowana w ramach podejścia konstruktivistycznego daje nam pewien dostęp do zrozumienia kształtowania się inteligentnych zachowań, niesie ze sobą również pewne ograniczenia. Dlatego w badaniach nad sztuczną inteligencją niezbędne jest podejście interdyscyplinarne. Biosemiotyka wydaje się być nieodłącznie związana z konstruktywnymi podejściami w aspekcie tworzenia inteligencji sztucznych agentów, uwzględniającej autonomię, autopoezę, interaktywność, intencjonalność i procesy powstawania znaczenia. Związek ten poszerza naukowe znaczenie biosemiotyki jako integrującego podejścia teoretycznego, zapewniając nieredukcyjny sposób użycia pojęć semiotycznych. Biosemiotyka uznaje twórczą rolę podmiotu poznawczego w budowaniu wiedzy, ale unika sceptycznego podejścia do możliwości poznania rzeczywistości. Nauka ta więc dobrze rokuje jako dyscyplina pośrednicząca w integracji semiotycznych i konstruktywnych podejść do sztucznej inteligencji. Wciąż jednak brakuje powszechnego zrozumienia zarówno

procesu poznania u żywych organizmach, jak i wiedzy na temat budowy ucieleśnionych, usytuowanych sztucznych agentów.

W pracy wskazano, w jaki sposób biosemiotyczne podejście może pomóc w zrozumieniu i modelowaniu inteligencji sztucznych systemów, poprzez zrozumiałe i proste koncepcje dotyczące powstawania inteligencji, mechanizmów poznania i nabywania wiedzy oraz w jaki sposób teoria ta może inspirować nowe modele SI. Inteligencja podmiotu poznawczego musi zostać traktowana jako wynik interakcji ze światem, w którym poprzez procesy powstawania znaczenia staje się istotnym czynnikiem jej rozwoju. Przyjmując jednak założenia biosemiotyki w kształtowaniu sztucznej inteligencji, musimy odejść od typowo ludzkiej perspektywy inteligencji, która wydaje się nieosiągalnym i niekoniecznym wyzwaniem dla sztucznych systemów. Analogicznie do dokonanego przez biosemiotykę wykazania odrębności oraz samoistnej wartości form inteligencji wytworzonych w różnych gatunkach organizmów żywych, od programu sztucznej inteligencji oczekiwać należy raczej sformułowania specyficznych form inteligencji. Można pokusić się o stwierdzenie, że inteligencja sztucznych systemów powinna być traktowana jako inteligencja sztucznych gatunków i z takiej perspektywy powinna być oceniana. Świadomość możliwości i ograniczeń biosemiotycznych powinny skłonić do porzucenia nierealistycznych dążeń do odtworzenia sztucznej kopii ludzkiej inteligencji, a zamiast tego skupić się na tym, aby sztuczni agenci znaleźli odpowiednią niszę w świecie naszych interakcji i by relacje z nimi miały charakter symbiotyczny. Zamiast humanizacji maszyn rozumianej jako dopasowywanie ich do ludzkiego wzorca inteligencji bardziej wartościowe praktycznie jak i poznawczo wydaje się skierowanie wysiłków na to, aby inteligentne maszyny stały się częścią naszego środowiska, z którym żyjemy w symbiozie. Taka „ekologiczna” perspektywa rozwoju sztucznej inteligencji jest jednym z najważniejszych wątków przyszłych badań rysujących się w perspektywie zaprezentowanych tu rozważań.

Niniejsza dysertacja proponując istotną zmianę paradygmatu myślenia o celach sztucznej inteligencji nie neguje dotychczasowych osiągnięć na tym polu, ale zwraca uwagę na konieczność rozwinięcia nowych badań w innym kierunku. Zapewniając sztuczным agentom większą swobodę i elastyczność oraz podążanie za ich rozwojem możemy uzyskać dostęp do wyników działań inteligentnych niedostępnych dla ludzkich podmiotów (np. z uwagi na ograniczenia czasowe lub z uwagi na

konieczność uwzględnienia zbyt wielu czynników naraz). Warto podkreślić, że omawiane tu podejście umożliwia powstanie inteligentnych systemów, które nie będą potrzebowały zwiększonej mocy obliczeniowej. Umożliwienie sztucznym systemom tworzenia i wykorzystywania struktur zewnętrznych poprzez znaczące relacje ze światem i innymi agentami powinno rzucić nieco więcej światła na to specyficzne zdolności poznawcze i specyficzne kształtowanie się inteligencji, a z drugiej strony powinno wskazać drogi do lepszego wykorzystania zasobów obliczeniowych jakie posiada nasza cywilizacja. Oszczędności zasobów obliczeniowych, na jakie pozwala analizowane podejście biosemiotyczne jest również godną uwagi konsekwencją opisywanego paradygmatu, co w dzisiejszych czasach nabiera szczególnego znaczenia. Od strony teoretycznej natomiast wiąże się to z lepszym zrozumieniem koncepcji obliczeń naturalnych (ang. *natural computing*), co rzuca również interesujące światło na koncepcję matematyczności rzeczywistości. Wskazane perspektywy mam nadzieję uzasadniają celowość dalszych badań w kierunku zastosowania perspektywy biosemiotycznej w badaniach sztucznej inteligencji.

Bibliografia

- Anderson M., Deely J., Krampen M., i in., *A semiotic perspective on the sciences: Steps toward a new paradigm*, „Semiotica”, t. 52, nr 1/2, 1984, ss. 7–47.
- Arkin R.C., Arkin R.C., *Behavior-based robotics*, MIT press 1998.
- Barbieri M., *Biosemiotics: a new understanding of life*, „Naturwissenschaften”, t. 95, nr 7, 2008, ss. 577–599.
- Barbieri M., *Life is “artifact-making”*, „Journal of Biosemiotics”, t. 1, nr 1, 2005, ss. 81–101.
- Bateson G., *Information, codification, and metacommunication*, „Communication and Culture, Holt, Rinehart and Winston, New York”, 1966, ss. 412–426.
- Beer R.D., *The Dynamics of Active Categorical Perception in an Evolved Model Agent*, „Adaptive Behavior”, t. 11, nr 4, 2003, ss. 209–243.
- Bennett R.J., Chorley R.J., *Environmental Systems: Philosophy, Analysis and Control*, Princeton University Press 2015.
- Berne E., Izdebski P., Wydawnictwo Naukowe PWN, *W co grają ludzie: psychologia stosunków międzyludzkich*, Warszawa 2017.
- Bickhard M.H., *Autonomy, function, and representation*, „Communication and Cognition-Artificial Intelligence”, t. 17, nr 3–4, 2000, ss. 111–131.
- Bickhard M.H., *Robots and representations*, [w:] *From animals to animats 5-Proceedings of the Fifth International Conference on Simulation of Adaptive Behavior*, Cambridge 1998, ss. 58–63.
- Bickhard M.H., *The interactivist model*, „Synthese”, t. 166, nr 3, 2009, ss. 547–591.
- Bickhard M.H., Terveen L., *Foundational issues in artificial intelligence and cognitive science: Impasse and solution*, Elsevier 1996, t. 109.
- Billard A., Dautenhahn K., *Grounding communication in situated, social robots*, [w:] 1997.
- Bolter J.D., *Człowiek Turinga: kultura Zachodu w wieku komputera*, T. Goban-Klas (tłum.), Państwowy Instytut Wydawniczy 1990.

- Boruszewski J., *Fizyczne systemy symboliczne i ich magia*, „Filo-Sofija”, t. 17, nr 36, 2017.
- Brentano F., *Psychology from an empirical standpoint*, Londyn 1874.
- Brier S., *Cybersemiotics: Why information is not enough!*, University of Toronto Press 2008.
- Brier S., *The paradigm of Peircean biosemiotics*, „Signs-International Journal of Semiotics”, t. 2, 2008, ss. 30–81.
- Brooks R.A., *A robust layered control system for a mobile robot*, „Robotics and Automation”, t. 2, nr 1, 1986, ss. 14–23.
- Brooks R.A., *Achieving Artificial Intelligence through Building Robots*, 1986.
- Brooks R.A., *Artificial life and real robots*, [w:] *Toward a practice of autonomous systems: Proc. of the 1st Europ. Conf. on Artificial Life*, 1992, s. 3.
- Brooks R.A., *Elephants don't play chess*, „Robotics and Autonomous Systems”, t. 6, nr 1–2, 1990, ss. 3–15.
- Brooks R.A., *Elephants don't play chess*, „Robotics and Autonomous Systems”, t. 6, nr 1–2, 1990, ss. 3–15.
- Brooks R.A., *Flesh and machines: How robots will change us*, Vintage 2003.
- Brooks R.A., *Intelligence without representation*, „Artificial Intelligence”, t. 47, nr 1–3, 1991, ss. 139–159.
- Brooks R.A., *New approaches to robotics*, „Science”, t. 253, nr 5025, 1991, ss. 1227–1232.
- Brooks R.A., *The Bitter Lesson*, [w:] *Rodney Brooks Robots, AI, and other stuff*.
- Bunge M., *System boundary*, „International Journal of General System”, t. 20, nr 3, 1992, ss. 215–219.
- Buss D.M., *Psychologia ewolucyjna*, tłum. M. Orski (tłum.), Gdańskie Wydawnictwo Psychologiczne 2001.
- Cangelosi A., Greco A., Harnad S., *Symbol Grounding and the Symbolic Theft Hypothesis*, [w:] *Simulating the Evolution of Language*, A. Cangelosi, D. Parisi (red.), Springer, London 2002.
- Cariani P., *Towards an evolutionary semiotics: the emergence of new sign-functions in organisms and devices*, [w:] *Evolutionary systems*, Springer 1998, ss. 359–376.
- Cassirer E., *The philosophy of symbolic forms. The metaphysics of symbolic forms*, J. M. Krois, D. P. Verene (red.), Yale 1996, t. 4.

- Chalmers D.J., Wolf S., *The conscious mind: in search of a fundamental theory*, „Integrative Physiological and Behavioral Science”, t. 32, nr 3, 1997, ss. 291–292.
- Chandrasekharan S., Nersessian N.J., *Building Cognition: The Construction of Computational Representations for Scientific Discovery*, „Cognitive Science”, t. 39, nr 8, 2015, ss. 1727–1763.
- Christensen W., Hooker C., *Anticipation in autonomous systems: foundations for a theory of embodied agents*, „Int J Comput Anticip Syst”, t. 5, 2000, ss. 135–154.
- Clark A., *Being there: Putting brain, body, and world together again*, MIT Press 1997.
- Clark A., *Natural-born cyborgs. Minds, technologies, and the future of human intelligence*, Oxford University Press 2004.
- Clark A., *Reasons, robots and the extended mind*, „Mind & Language”, t. 16, nr 2, 2001, ss. 121–145.
- Coradeschi S., Saffiotti A., *An introduction to the anchoring problem*, „Robotics and Autonomous Systems”, t. 43, nr 2–3, 2003, ss. 85–96.
- Craik K.J.W., *The nature of explanation*, Cambridge 1967.
- Czarnocka M., *Podmiotowość w neokantyzmie*, „Filozofia i nauka. Studia filozoficzne i interdyscyplinarne”, t. 3, 2015, ss. 119–139.
- D. Wilson A., Golonka S., *Ucieleśnienie poznania to nie to, co myślisz*, „Avant”, t. 5, 2014, ss. 21–56.
- Damasio A., *Descartes' error: emotion, reason, and the human brain*, New York 1994.
- Darwin C., *The Origin of Species*, Alachua 2009.
- De Waal F., *Are we smart enough to know how smart animals are?*, WW Norton & Company 2016.
- Deacon T., *The symbolic species*, New York 1997.
- Deely J., *Four ages of understanding*, University of Toronto Press 2001.
- Deely J., *Pars Pro Toto from Culture to Nature: An Overview of Semiotics as a Postmodern Development, with an Anticipation of Developments to Come*, „The American Journal of Semiotics”, t. 25, nr 1, 2009, ss. 167–192.
- Deely J.N., *Basics of semiotics*, Indiana University Press 1990.
- Deep Blue*, [w:] <http://www-03.ibm.com/ibm/history/ibm100/us/en/icons/deepblue/> (15.10.2017).
- Denning P.J., *The science of computing: Is thinking computable?*, „American Scientist”, t. 78, nr 2, 1990, ss. 100–102.

- Descartes R., *Medytacje o pierwszej filozofii*, M. Ajdukiewicz, K. Ajdukiewicz (tłum.), Kęty 2001.
- Descartes R., *Rozprawa o metodzie*, T. Żeleński (tłum.), Warszawa 1952.
- Dorffner G., *Radical connectionism-a neural bottom-up approach to AI*, „Neural Networks and a New Artificial Intelligence”, 1997, ss. 93–132.
- Dorffner G., *Radical connectionism-a neural bottom-up approach to AI*, „Neural Networks and a New Artificial Intelligence”, 1997, ss. 93–132.
- Dreyfus H., Dreyfus S., *Making a Mind Versus Modelling the Brain: Artificial Intelligence Back at a Branchpoint*, „Daedalus”, t. 117, nr 1, 1988, ss. 185–197.
- Dreyfus H.L., *What computers can't do: The limits of artificial intelligence*, Harper & Row New York 1979.
- Dreyfus H.L., *What computers still can't do: a critique of artificial reason*, MIT press 1992.
- Dreyfus H.L., *Why Heideggerian AI failed and how fixing it would require making it more Heideggerian*, „Artificial Intelligence”, t. 171, nr 18, 2007, ss. 1137–1160.
- Elman J., *Generalization, simple recurrent networks, and the emergence of structure*, [w:] Mahwah, NJ: Lawrence Erlbaum Associates 1998, s. 6.
- Elman J.L., *Distributed representations, simple recurrent networks, and grammatical structure*, „Machine learning”, t. 7, nr 2–3, 1991, ss. 195–225.
- Emmeche C., *A biosemiotic note on organisms, animals, machines, cyborgs, and the quasi-autonomy of robots*, „Pragmatics & Cognition”, t. 15, nr 3, 2007, ss. 455–483.
- Emmeche C., *Does a robot have an Umwelt? Reflections on the qualitative biosemiotics of Jakob von Uexkull*, „Semiotica”, t. 134, nr 1/4, 2001, ss. 653–694.
- Emmeche C., *Life as an abstract phenomenon: Is artificial life possible*, [w:] *Toward a practice of autonomous systems - Proceedings of the First European Conference on Artificial Life*, F. Varela, P. Bourgine (red.), Cambridge 1992, ss. 466–474.
- Emmeche C., *Modeling life: A note on the semiotics of emergence and computation in artificial and natural living systems*, [w:] *Biosemiotics. The Semiotic Web*, T. A. Sebeok (red.), Berlin 1992.
- Emmeche C., *Organicism and qualitative aspects of self-organization*, „Revue internationale de philosophie”, nr 2, 2004, ss. 205–217.
- Emmeche C., *Organism and body: The semiotics of emergent levels of life*, [w:] *Towards a Semiotic Biology: Life is the action of signs*, World Scientific 2011, ss. 91–111.

- Emmeche C., *The emergence of signs of living feeling: Reverberations from the first Gatherings in Biosemiotics*, „Σημειωτική - Sign Systems Studies”, t. 29, nr 1, 2001, ss. 369–376.
- Emmeche C., Hoffmeyer J., *From language to nature: The semiotic metaphor in biology*, „Semiotica”, t. 84, nr 1–2, 1991, ss. 1–42.
- Emmeche C., Kull K., *Towards a semiotic biology: Life is the action of signs*, World Scientific 2011.
- Engelkamp J., *Organisation and recall in verbal tasks and in subject-performed tasks*, „European Journal of Cognitive Psychology”, t. 8, nr 3, 1996, ss. 257–274.
- Epley N., Waytz A., Cacioppo J.T., *On seeing human: A three-factor theory of anthropomorphism.*, „Psychological Review”, t. 114, nr 4, 2007, ss. 864–886.
- Favareau D., *Beyond self and other: On the neurosemiotic emergence of intersubjectivity*, „Sign Systems Studies”, t. 30, nr 1, 2002, ss. 57–99.
- Favareau D., *Essential Readings in Biosemiotics: Anthology and Commentary*, Springer Science & Business Media 2010.
- Ferreira M.I.A., *On Meaning: A Biosemiotic Approach*, „Biosemiotics”, t. 3, nr 1, 2010, ss. 107–130.
- Ferreira M.I.A., Caldas M.G., *Modelling Artificial Cognition in Biosemiotic Terms*, „Biosemiotics”, t. 6, nr 2, 2012, ss. 245–252.
- Flasiński M., *Wstęp do sztucznej inteligencji*, Warszawa 2011.
- Floreano D., Husbands P., Nolfi S., *Evolutionary Robotics*, [w:] *Springer Handbook of Robotics*, B. Siciliano, O. Khatib (red.), Berlin, Heidelberg 2008, ss. 1423–1451.
- Franklin S., *Autonomous agents as embodied AI*, „Cybernetics & Systems”, t. 28, nr 6, 1997, ss. 499–520.
- Franklin S., Graesser A., *Is It an agent, or just a program?, A taxonomy for autonomous agents*, [w:] *Intelligent Agents III Agent Theories, Architectures, and Languages*, Springer, Berlin, Heidelberg 1996, ss. 21–35.
- Franklin S., Wolpert S., McKay S.R., Christian W., *Artificial minds*, „Computers in Physics”, t. 11, nr 3, 1997, ss. 258–259.
- Fredslund J., Mataric M.J., *A general algorithm for robot formations using local sensing and minimal communication*, „IEEE transactions on robotics and automation”, t. 18, nr 5, 2002, ss. 837–846.

- Gilbert S.F., Epel D., *Ecological developmental biology: integrating epigenetics, medicine, and evolution*, Sunderland 2009.
- Glaserfeld E. von, *The Radical Constructivist View of Science*, „Foundations of Science”, t. 6, nr 1, 2001, ss. 31–43.
- Graubard S.R., *The artificial intelligence debate*, MIT press Cambridge (Ma) 1988.
- Hall W.P., *Biological nature of knowledge in the learning organisation*, „The Learning Organization”, t. 12, nr 2, 2005, ss. 169–188.
- Hammond D., *The Science of Synthesis: Exploring the Social Implications of General Systems Theory*, University Press of Colorado 2010.
- Harnad S., *Computation is just interpretable symbol manipulation; cognition isn't*, „Minds and Machines”, t. 4, nr 4, 1994, ss. 379–390.
- Harnad S., *The symbol grounding problem*, „Physica D: Nonlinear Phenomena”, t. 42, nr 1, 1990, ss. 335–346.
- Harter S., *The construction of the self: A developmental perspective.*, Guilford Press 1999.
- Haugeland J., *Artificial Intelligence: The Very Idea*, MIT Press 1989.
- Hawkins D.R., *Reality and Subjectivity*, Hay House, Inc 2013.
- Hayes P., *What the frame problem is and isn't*, 1987.
- Heath G., *A constructivist attempts to talk to the field*, „International Journal of Psychotherapy”, t. 5, nr 1, 2000, ss. 11–35.
- Heidegger M., *Bycie i czas*, Warszawa 1994.
- Heidegger M., *Die Grundbegriffe der Metaphysik: Welt, Endlichkeit, Einsamkeit*, Frankfurt am Main 1983.
- Heinrich B., Kober H., *Mind of the raven: investigations and adventures with wolf-birds*, Cliff Street Books New York 1999.
- Hendriks-Jansen H., *Catching ourselves in the act: Situated activity, interactive emergence, evolution, and human thought*, Cambridge 1996.
- Hinton G.E., Sejnowski T.J. (red.), *Unsupervised learning: foundations of neural computation*, Cambridge, Mass 1999.
- Hoffmeyer J., *Biosemiotics: an examination into the signs of life and the life of signs*, Scranton 2008.
- Hoffmeyer J., *From Thing to Relation. On Bateson's Bioanthropology*, [w:] *A Legacy for Living Systems*, t. 2, 2008, ss. 27–44.

- Hoffmeyer J., *Semiotic Scaffolding of Living Systems*, [w:] *Introduction to Biosemiotics*, M. Barbieri (red.), Springer Netherlands 2008, ss. 149–166.
- Hoffmeyer J., *Signs of Meaning in the Universe*, Indiana University Press 1997.
- Hoffmeyer J., *Surfaces inside surfaces. On the origin of agency and life*, „Cybernetics & Human Knowing”, t. 5, nr 1, 1998, ss. 33–42.
- Hoffmeyer J., *The Natural History of Intentionality. A Biosemiotic Approach*, [w:] *The Symbolic Species Evolved*, T. Schilhab, F. Stjernfelt, T. Deacon (red.), t. 6, Dordrecht 2012, ss. 97–116.
- Hoffmeyer J., *The semiome: From genetic to semiotic scaffolding*, „Semiotica”, t. 2014, nr 198, 2014, ss. 11–31.
- Hoffmeyer J., *The Unfolding Semiosphere*, [w:] *Evolutionary Systems*, Springer, Dordrecht 1998, ss. 281–293.
- Hoffmeyer J., Kull K., *Baldwin and biosemiotics: What intelligence is for*, [w:] *Evolution and learning: The Baldwin effect reconsidered*, B. Weber, D. Depew (red.), Cambridge 2003, ss. 253–272.
- Hoffmeyer J.N., *Biosemiotics: Towards a new synthesis in biology*, „European Journal for Semiotic Studies”, t. 9, nr 2, 1997, ss. 355–376.
- Holland J.H., *Hidden order: how adaptation builds complexity*, Cambridge, Mass. 2003.
- Husserl E., *Idee czystej fenomenologii i fenomenologicznej filozofii*, D. Gierulanka (tłum.), Państwowe Wydawnictwo Naukowe 1975.
- Inteligencja*, [w:] <https://encyklopedia.pwn.pl/haslo/inteligencja;3915042.html>.
- Inverno M., Luck M., *Creativity through autonomy and interaction*, „Cognitive Computation”, t. 4, nr 3, 2012, ss. 332–346.
- Jacob P., *Intentionality*, [w:] *The Stanford Encyclopedia of Philosophy*, E. N. Zalta (red.), Metaphysics Research Lab, Stanford University 2014.
- Jensen A.R., *The g factor and the design of education*, „Intelligence, instruction, and assessment: Theory into practice”, 1998, ss. 111–131.
- Joldersma C.W., *Ernst von Glasersfeld's radical constructivism and truth as disclosure*, „Educational Theory”, t. 61, nr 3, 2011, ss. 275–293.
- Jonassen D.H., *Objectivism versus constructivism: Do we need a new philosophical paradigm?*, „Educational technology research and development”, t. 39, nr 3, 1991, ss. 5–14.

- Joolingen W.R. van, Jong T. de, Lazonder A.W., i in., *Co-Lab: research and development of an online learning environment for collaborative scientific discovery learning*, „Computers in human behavior”, t. 21, nr 4, 2005, ss. 671–688.
- Kampis G., *Self-modifying systems*, „Biology and Cognitive Science: A New Framework for Dynamics, Information, and Complexity”, 1991.
- Kant I., *Krytyka czystego rozumu*, R. Ingarden (tłum.), Warszawa 1957.
- Kaveh K., Bui M.D., Rutschmann P., *Development of an Artificial-Neural-Network-based concept for hydro-morphodynamic modelling in rivers*, [w:] *New Challenges in hydraulic research and engineering. Proc. of the 5th IAHR-Europe Congress*, 2018, ss. 721–722.
- Kawade Y., *The two foci of biology: Matter and sign*, „Semiotica”, t. 127, nr 1–4, 1999, ss. 369–384.
- Kirsh D., *Adapting the Environment Instead of Oneself*, „Adaptive Behavior”, t. 4, nr 3–4, 1996, ss. 415–452.
- Koriat A., Pearlman-Avnion S., *Memory organization of action events and its relationship to memory performance.*, „Journal of Experimental Psychology: General”, t. 132, nr 3, 2003, s. 435.
- Kull K., *Catalysis and scaffolding in semiosis*, [w:] *The Catalyzing Mind*, Springer 2014, ss. 111–121.
- Kull K., *Evolution, choice, and scaffolding: semiosis is changing its own building*, „Biosemiotics”, t. 8, nr 2, 2015, ss. 223–234.
- Kull K., *On Semiosis, Umwelt, and Semiosphere*, „Semiotica”, t. 120, nr 3–4, 1998, ss. 299–310.
- Kull K., *Semiosis stems from logical incompatibility in organic nature: Why biophysics does not see meaning, while biosemiotics does*, „Progress in Biophysics and Molecular Biology”, t. 119, nr 3, 2015, ss. 616–621.
- Kull K., *Towards biosemiotics with Yuri Lotman*, „Semiotica”, t. 127, nr 1–4, 1999, ss. 115–132.
- Kull K., Deacon T., Emmeche C., i in., *Theses on biosemiotics: Prolegomena to a theoretical biology*, [w:] *Towards a Semiotic Biology: Life is the Action of Signs*, World Scientific 2011, ss. 25–41.
- Kull K., Torop P., *Biotranslation: Translation between umwelten*, „Readings in Zoosemiotics.”, t. 8, 2011, ss. 411–425.

- Kurowski M.M., *Niklasa Luhmanna radykalizacja projektu fenomenologii*, „Fenomenologia”, nr 1, 2003, ss. 63–74.
- Kurzweil R., *The singularity is near: When humans transcend biology*, Penguin 2005.
- Lahti D.C., *The limits of artificial stimuli in behavioral research: the umwelt gamble*, „Ethology”, t. 121, nr 6, 2015, ss. 529–537.
- Lakoff G., *Smolensky, semantics, and the sensorimotor system*, „Behavioral and Brain Sciences”, t. 11, nr 1, 1988, ss. 39–40.
- Lakoff G., Johnson M., *Metaphors we live by*, University of Chicago press 2008.
- Lakoff G., Johnson M., *Philosophy in the Flesh*, New York 1999.
- Lakoff G., Johnson M., *The metaphorical structure of the human conceptual system*, „Cognitive science”, t. 4, nr 2, 1980, ss. 195–208.
- Lakoff G., Núñez R., *Where Mathematics Comes From: How the Embodied Mind Brings Mathematics into Being*, New York 2000.
- Langton C.G., Taylor C., Farmer J., Rassmussen S., *Artificial Life*, [w:] *Proceedings of an Interdisciplinary Workshop on the Synthesis and Simulation of Living Systems*, t. 6, New York 1989.
- Laughlin R.B., *A Different Universe: Reinventing Physics from the Bottom Down*, New York 2005.
- Lefebvre V., *Conflicting Structures*, Los Angeles 2015.
- Lektorskii V.A., *Realism, Antirealism, Constructivism, and Constructive Realism in Contemporary Epistemology and Science*, „Journal of Russian & East European Psychology”, t. 48, nr 6, 2010, ss. 5–44.
- Loeb J., *Concerning the theory of tropisms*, „Journal of Experimental Zoology”, t. 4, nr 1, 1907, ss. 151–156.
- Lorenz K., *The Foundations of Ethology*, Springer Science & Business Media 1981.
- Lotman Y., *Universe of the Mind. A Semiotic Theory of Culture*, Indiana 1990.
- MacDorman K., *Cognitive Robotics. Grounding symbols through sensorimotor integration.*, „Journal of the Robotics Society of Japan”, t. 17, nr 1, 1999, ss. 20–24.
- Magnus R., *Time-plans of the organisms: Jakob von Uexküll's explorations into the temporal constitution of living beings*, „Σημειωτική-Sign Systems Studies”, t. 39, nr 2–4, 2011, ss. 37–57.
- Maran T., Kleisner K., *Towards an evolutionary biosemiotics: semiotic selection and semiotic co-option*, „Biosemiotics”, t. 3, nr 2, 2010, ss. 189–200.

- Martinez M.E., *Future Bright: A Transforming Vision of Human Intelligence*, OUP USA 2013.
- Mataric M.J., Brooks R.A., *Learning a distributed map representation based on navigation behaviors*, [w:] *Proceedings of 1990 USA Japan Symposium on Flexible Automation*, 1990, ss. 499–506.
- Matthews M.R., *Constructivism in science education: A philosophical examination*, Springer Science & Business Media 1998.
- Matthews M.R., *Introductory comments on philosophy and constructivism in science education*, „Science & Education”, t. 6, nr 1–2, 1997, ss. 5–14.
- Maturana F., Shen W., Norrie D.H., *MetaMorph: an adaptive agent-based architecture for intelligent manufacturing*, „International Journal of Production Research”, t. 37, nr 10, 1999, ss. 2159–2173.
- Maturana H.R., *Autopoiesis and Cognition: The Realization of the Living*, Springer Science & Business Media 1980.
- Maturana H.R., Guilloff G.D., *The quest for the intelligence of intelligence*, „Journal of Social and Biological Structures”, t. 3, nr 2, 1980, ss. 135–148.
- Maturana H.R., Varela F.J., *The tree of knowledge: The biological roots of human understanding*, New Science Library/Shambhala Publications 1987.
- Mazur M., *Cybernetyka i charakter*, Warszawa 1999.
- McCarthy J., Hayes P.J., *Some philosophical problems from the standpoint of artificial intelligence*, [w:] *Readings in artificial intelligence*, Elsevier 1981, ss. 431–450.
- McCarthy J., Minsky M.L., Rochester N., *A Proposal for The Dartmouth Summer Research Project on Artificial Intelligence*, 1955.
- McCorduck P., *Machines who think: a personal inquiry into the history and prospects of artificial intelligence*, Natick, Mass. 2004.
- McCulloch W.S., Pitts W., *A logical calculus of the ideas immanent in nervous activity*, „The bulletin of mathematical biophysics”, t. 5, nr 4, 1943, ss. 115–133.
- Menzel R., Fischer J., *Animal thinking: contemporary issues in comparative cognition*, MIT press 2011, t. 8.
- Merleau-Ponty M., *Fenomenologia percepcji*, M. Kowalska, J. Migasiński (tłum.), Fundacja Aletheia 2001.
- Merleau-Ponty M., *Nature: Course Notes from the Collège de France*, Northwestern University Press 2003.

- Meyer J.-A., *Artificial life and the animat approach to artificial intelligence*, [w:] *Artificial intelligence*, Elsevier 1996, ss. 325–354.
- Meyer J.-A., *From natural to artificial life: Biomimetic mechanisms in animat designs*, „Robotics and Autonomous Systems”, t. 22, nr 1, 1997, ss. 3–21.
- Miles L.K., Nind L.K., Macrae C.N., *Moving Through Time*, „Psychological Science”, t. 21, nr 2, 2010, ss. 222–223.
- Mingers J., *Self-producing systems: implications and applications of autopoiesis*, New York 1995.
- Moravec H., *When will computer hardware match the human brain*, „Journal of evolution and technology”, t. 1, nr 1, 1998, s. 10.
- Morris C., *Signs, language and behavior.*, Oxford 1946.
- Nagel T., *Jak to jest być nietoperzem*, [w:] *Pytania ostateczne*, A. Romaniuk (tłum.), Warszawa 1997, ss. 203–219.
- Newell A., *Physical symbol systems*, „Cognitive science”, t. 4, nr 2, 1980, ss. 135–183.
- Newell A., *Physical symbol systems*, „Cognitive science”, t. 4, nr 2, 1980, ss. 135–183.
- Newell A., Simon H.A., *Computer science as empirical inquiry: Symbols and search*, „Communications of the ACM”, t. 19, nr 3, 1976, ss. 113–126.
- Nolfi S., Bongard J., Husbands P., Floreano D., *Evolutionary robotics*, [w:] *Springer Handbook of Robotics*, Springer 2016, ss. 2035–2068.
- Nolfi S., Floreano D., *Coevolving Predator and Prey Robots: Do “Arms Races” Arise in Artificial Evolution?*, „Artificial Life”, t. 4, nr 4, 1998, ss. 311–335.
- Nöth W., *Semiotic machines*, „Cybernetics & Human Knowing”, t. 9, nr 1, 2002, ss. 5–21.
- Nöth W., *Sign machines in the framework of Semiotics Unbounded*, „Semiotica”, t. 2008, nr 169, 2008, ss. 319–341.
- Patee H., *Evolving Self-Reference: Matter, Symbols and Semantic Closure*, „Communication and Cognition - Artificial Intelligence”, t. 12, 1995, ss. 9–27.
- Pattee H.H., *Simulations, Realizations, and Theories of Life.*, [w:] 1987, ss. 63–78.
- Pavlov I.P., *Conditioned reflexes: An investigation of the physiological activity of the cerebral cortex*, G. V. Anrep (tłum.), Oxford University Press London 1928.
- Peirce C.S., *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 1: Principles of Philosophy*, Charlottesville, Va 1994.
- Peirce C.S., *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 2: Elements of Logic*, Charlottesville, Va 1994.

- Peirce C.S., *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 4: The Simplest Mathematics*, Charlottesville, Va 1994.
- Peirce C.S., *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 5: Pragmatism and Pragmaticism*, Charlottesville, Va 1994.
- Peirce C.S., *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 7: Science and Philosophy*, Charlottesville, Va 1994.
- Peirce C.S., *The Collected Papers of Charles Sanders Peirce. Electronic edition. Volume 8: Reviews, Correspondence, and Bibliography*, Charlottesville, Va 1994.
- Peirce C.S., *The essential Peirce: Selected philosophical writings*, N. Houser (red.), Indiana University Press 1998.
- Pelc J., *Wstęp do semiotyki*, Warszawa 1984.
- Peschl M.F., Riegler A., *Does Representation Need Reality?*, [w:] *Understanding Representation in the Cognitive Sciences: Does Representation Need Reality?*, A. Riegler, M. Peschl, A. von Stein (red.), Boston, MA 1999, ss. 9–17.
- Pezzulo G., Barsalou L.W., Cangelosi A., i in., *The Mechanics of Embodiment: A Dialog on Embodiment and Computational Modeling*, „Frontiers in Psychology”, t. 2, 2011.
- Pfeifer R., Bongard J., *How the Body Shapes the Way We Think: A New View of Intelligence*, Cambridge 2006.
- Pfeifer R., Gómez G., *Interacting with the real world: design principles for intelligent systems*, „Artificial Life and Robotics”, t. 9, nr 1, 2005, ss. 1–6.
- Pfeifer R., Scheier C., *Understanding intelligence*, MIT Press 2001.
- Piaget J., *La construction du réel chez l'enfant*, Neuchâtel; Paris 1937.
- Piaget J., *The construction of reality in the child.*, M. Cook (tłum.), New York 1954.
- Pinker S., *Jak działa umysł*, M. Koraszewska (tłum.), Książka i Wiedza 2002.
- Place U.T., *Is consciousness a brain process?*, [w:] *The Mind-Brain Identity Theory*, C. V. Borst (red.), London 1970, ss. 42–51.
- Prem E., *A Biosemiotic Framework for Artificial Autonomous Sign Users*, [w:] <https://www.aaai.org/Papers/Workshops/2004/WS-04-03/WS04-03-007.pdf>.
- Purves D., *Body and brain: a trophic theory of neural connections*, Cambridge 1988.
- Rafieian S., *A biosemiotic approach to the problem of structure and agency*, „Bisemiotics”, t. 5(1), 2012, ss. 83–93.
- Richter C.P., *Animal Behavior and Internal Drives*, „The Quarterly Review of Biology”, t. 2, nr 3, 1927, ss. 307–343.

- Riegler A., *Constructivist artificial life, and beyond*, „Proceedings of the workshop on autopoiesis and perception”, 1992, ss. 121–136.
- Rocha L.M., *Evolution with material symbol systems*, „Biosystems”, t. 60, nr 1, 2001, ss. 95–121.
- Rosch E., *Principles of categorization*, „Concepts: core readings”, t. 189, 1999.
- Rosen R., *Fundamentals of measurement and representation of natural systems*, Elsevier Science Ltd 1978, t. 1.
- Rosen R., Pattee H.H., Somorjai R.L., *A symposium in theoretical biology*, „A question of physics: Conversations in physics and biology”, 1979, ss. 84–123.
- Rothschild F.S., *Creation and Evolution: A Biosemiotic Approach*, Transaction Publishers 2000.
- Rothschild F.S., *Laws of symbolic mediation in the dynamics of self and personality*, „Annals of the New York Academy of Sciences”, t. 96, nr 3, 1962, ss. 774–784.
- Rudolph L., *A Unified Topological Approach to Umwelts and Life Spaces Part II: Constructing Life Spaces from an Umwelt*, [w:] *The Observation of Human Systems*, Routledge 2017, ss. 117–140.
- Rumelhart D.E., McClelland J.L., *Psychological and Biological Models*, MIT Press 1986.
- Ryan R.M., Deci E.L., *Self-determination theory: Basic psychological needs in motivation, development, and wellness*, Guilford Publications 2017.
- Salge C., Guckelsberger C., *Does empowerment maximisation allow for enactive artificial agents?*, [w:] MIT Press 2016, ss. 704–711.
- Sandoval J., *Constructivism, consultee-centered consultation, and conceptual change*, „Journal of Educational and Psychological Consultation”, t. 7, nr 1, 1996, ss. 89–97.
- Sarosiek, A., *Kręgi funkcjonalne jako wzór dla modelowania procesów autonomicznych sztucznych systemów*, [w:] *Studia z filozofii informatyki*, K. Sołoducha, P. Stacewicz (red.), Warszawa 2018, ss. 171–190.
- Schlosser M., *Agency*, [w:] *The Stanford Encyclopedia of Philosophy*, E. N. Zalta (red.), Metaphysics Research Lab, Stanford University 2015.
- Schmidt C.T.A., Kraemer F., *Robots, Dennett and the Autonomous: A Terminological Investigation*, „Minds and Machines”, t. 16, nr 1, 2006, ss. 73–80.
- Searle J.R., *Minds, brains, and programs*, „Behavioral and brain sciences”, t. 3, nr 03, 1980, ss. 417–424.
- Searle J.R., *The rediscovery of the mind*, Cambridge 1992.

- Sebeok T.A., *Animal.Biological and semiotic perspective*, [w:] *What is an Animal*, T. Ingold (red.), 1988, ss. 63–76.
- Sebeok T.A., *Communication, Language and Speech: Evolutionary Considerations*, [w:] *Semiotic Theory and Practice: Proceedings of the Third International Congress of the IASS Palermo, 1984*, M. Herzfeld, L. Melazzo (red.), t. II, Berlin, Boston 1988, ss. 1083–1091.
- Sebeok T.A., *Contributions to the Doctrine of Signs*, University Press of America 1976, t. 4.
- Sebeok T.A., *Semiotics and ethology*, [w:] *Approaches to Animal Communication*, T. A. Sebeok, A. Ramsay (red.), Haga 1969, ss. 200–231.
- Sebeok T.A., *The semiotic self*, [w:] *A Sign is Just a Sign*, Bloomington 1979, ss. 263–267.
- Sebeok T.A., *The semiotic self revisited*, [w:] *A sign is just a sign*, 1991, ss. 41–48.
- Sebeok Thomas A., „Tell me, where is fancy bred?” *The biosemiotic self*, [w:] *Biosemiotics: The Semiotic Web 1991*, J. Umiker-Sebeok (red.), Berlin 1992.
- Shanahan M., *The Frame Problem*, [w:] *The Stanford Encyclopedia of Philosophy*, E. N. Zalta (red.), Metaphysics Research Lab, Stanford University 2016.
- Shannon C.E., *A mathematical theory of communication*, „Bell system technical journal”, t. 27, nr 3, 1948, ss. 379–423.
- Sharkey N.E., Ziemke T., *Mechanistic versus phenomenal embodiment: Can robot embodiment lead to strong AI?*, „Cognitive Systems Research”, t. 2, nr 4, 2001, ss. 251–262.
- Sharov A., Maran T., Tønnessen M., *Comprehending the Semiosis of Evolution*, „Biosemiotics”, t. 9, nr 1, 2016, ss. 1–6.
- Sharov A.A., *Biosemiotics: A functional-evolutionary approach to the analysis of the sense of information*, [w:] *Biosemiotics: The semiotic web*, 1991, ss. 345–373.
- Sharov A.A., *Pragmatics and biosemiotics.*, „Sign Systems Studies”, t. 30, nr 1, 2002.
- Sheets-Johnstone M., *Consciousness: A natural history*, „Journal of Consciousness Studies”, t. 5, nr 3, 1998, ss. 260–294.
- Shelton J., *The role of observation and simplicity in Einstein’s epistemology*, „Studies in History and Philosophy of Science Part A”, t. 19, nr 1, 1988, ss. 103–118.
- Shepard R.N., *Toward a universal law of generalization for psychological science*, „Science”, t. 237, nr 4820, 1987, ss. 1317–1323.

- Spearman C., „ *General Intelligence, Objectively Determined and Measured*, „The American Journal of Psychology”, t. 15, nr 2, 1904, ss. 201–292.
- Steels L., *The symbol grounding problem has been solved. So what's next*, „Symbols and embodiment: Debates on meaning and cognition”, 2008, ss. 223–244.
- Steels L., Brooks R., *The artificial life route to artificial intelligence: Building embodied, situated agents*, London 2018.
- Steels L., Vogt P., *Grounding adaptive language games in robotic agents*, [w:] *Proceedings of the fourth european conference on artificial life*, t. 97, MIT Press 1997, ss. 474–482.
- Steffens M.C., Buchner A., Wender K.F., *Quite ordinary retrieval cues may determine free recall of actions*, „Journal of Memory and Language”, t. 48, nr 2, 2003, ss. 399–415.
- Stewart I., *Self-organization in evolution: a mathematical perspective*, „Philosophical Transactions of the Royal Society of London. Series A: Mathematical, Physical and Engineering Sciences”, t. 361, nr 1807, 2003, ss. 1101–1123.
- Stjernfelt F., *Tractatus Hoffmeyerensis: Biosemiotics as expressed in 22 basic hypotheses*, „Sign Systems Studies”, t. 30, nr 1, 2002, ss. 337–345.
- Stjernfelt F., Emmeche C., Kull K., *Reading Hoffmeyer, rethinking biology*, Tartu University Press 2002.
- Strazzoni A., *Subjectivity and individuality: Two strands in early modern philosophy: Introduction*, „Societate si Politica”, t. 9, nr 1, 2015, ss. 5–9.
- Susi T., Ziemke T., *On the subject of objects: Four views on object perception and tool use*, „tripleC”, t. 3, nr 2, 2005, ss. 6–19.
- System autopojetyczny*, [w:] *Wikipedia, wolna encyklopedia*, 2018.
- Taddeo M., Floridi L., *Solving the symbol grounding problem: a critical review of fifteen years of research*, „Journal of Experimental & Theoretical Artificial Intelligence”, t. 17, nr 4, 2005, ss. 419–445.
- Tadeusiewicz R., *Sieci neuronowe*, Akademicka Oficyna Wydawnicza 1993.
- Thompson E., *Empathy and Consciousness*, „Journal of Consciousness Studies”, t. 8, nr 5–7, 2001, ss. 1–32.
- Thorndike E., *Animal intelligence: Experimental studies*, Routledge 2017.
- Tønnessen M., *Biosemosis: Questionnaire from Biosemiotics - and info about the biosemiotic glossary project*, [w:] <http://biosemiosis.blogspot.com/2013/11/questionnaire-from-biosemitics-and.html> (4.12.2018).

- Tønnessen M., *The biosemiotic glossary project: Agent, agency*, „Biosemiotics”, t. 8, nr 1, 2015, ss. 125–143.
- Tooby J., Cosmides L., *The psychological foundations of culture*, „The adapted mind: Evolutionary psychology and the generation of culture”, 1995, ss. 19–136.
- Touretzky D.S., Pomerleau D.A., *Reconstructing physical symbol systems*, „Cognitive Science”, t. 18, nr 2, 1994, ss. 345–353.
- Turing A.M., *Computing machinery and intelligence*, „Mind”, 1950, ss. 433–460.
- Turkle S., Breazeal C., Dasté O., Scassellati B., *Encounters with kismet and cog: Children respond to relational artifacts*, „Digital media: Transformations in human communication”, t. 120, 2006, ss. 313–330.
- Turner J.S., *Semiotics of a superorganism*, „Biosemiotics”, t. 9, nr 1, 2016, ss. 85–102.
- Uexküll J. von, *Bedeutungslehre*, Leipzig 1940.
- Uexküll J. von, *Streifzüge durch die Umwelten von Tieren und Menschen. Ein Bilderbuch unsichtbarer Welten.*, Rohwohlt Verlag, Hamburg 1934.
- Uexküll J. von, *The Theory of Meaning*, „Semiotica”, t. 42, nr 1, 2009, ss. 25–79.
- Uexküll J. von, *Theoretische biologie*, Berlin 1928.
- Uexküll J. von, *Umwelt und Innenwelt der Tiere*, Berlin 1921.
- Uexküll T., *Semiotics and the problem of the observer*, „Semiotica”, t. 48, nr 3/4, 1984, ss. 187–195.
- Uexküll T. von, *Glossary*, „Semiotica”, t. 42, nr 1, 2009, ss. 83–87.
- Uexküll T. von, *Introduction: Meaning and science in Jakob von Uexküll's concept of biology*, „Semiotica”, t. 42, nr 1, 1982, ss. 1–24.
- Uexküll T. von, *The Sign Theory of Jakob von Uexküll*, [w:] *Classics of Semiotics*, M. Krampen, K. Oehler, R. Posner, i in. (red.), Berlin 1987, ss. 147–179.
- Varela F., *Autopoiesis and a biology of intentionality*, [w:] *Proceedings of a workshop on Autopoiesis and Percetion*, Dublin 1992, ss. 4–14.
- Varela F.J., Thompson E., Rosch E., *The embodied mind*, MIT Press 1993.
- Von Glasersfeld E., *Cognition, construction of knowledge, and teaching*, „Synthese”, t. 80, nr 1, 1989, ss. 121–140.
- Von Glasersfeld E., *Radical Constructivism: A Way of Knowing and Learning*. London, Wasington, DC, London 1995.
- Vygotsky L.S., *Thought and language*, „Annals of Dyslexia”, t. 14, nr 1, 1964, ss. 97–98.

- Weijermars R., *Building Corporate IQ – Moving the Energy Business from Smart to Genius: Executive Guide to Preventing Costly Crises*, Springer Science & Business Media 2011.
- Whitley D., *Next Generation Genetic Algorithms: A User's Guide and Tutorial*, [w:] *Handbook of Metaheuristics*, M. Gendreau, J.-Y. Potvin (red.), Cham 2019, ss. 245–274.
- Wiener N., *Cybernetics or Control and Communication in the Animal and the Machine*, MIT press 1965, t. 25.
- Wittgenstein L., *Dociekania filozoficzne*, B. Wolniewicz (tłum.), Państwowe Wydawnictwo Naukowe 1972.
- Zavertanyy V.V., Makarenko A.S., *Genotype dynamic for agent neuroevolution in artificial life model*, „System research and information technologies”, t. 0, nr 1, 2017, ss. 75–87.
- Ziemke T., *Cybernetics and embodied cognition: on the construction of realities in organisms and robots*, „Kybernetes”, t. 34, nr 1/2, 2005, ss. 118–128.
- Ziemke T., *On the role of emotion in biological and robotic autonomy*, „Biosystems”, t. 91, nr 2, 2008, ss. 401–408.
- Ziemke T., *On 'parts' and 'wholes' of adaptive behavior: Functional modularity and diachronic structure in recurrent neural robot controllers*, [w:] *From animals to animats*, t. 6, 2000, ss. 171–180.
- Ziemke T., *The construction of 'reality' in the robot: Constructivist perspectives on situated artificial intelligence and adaptive robotics*, „Foundations of Science”, t. 6, nr 1–3, 2001, ss. 163–233.
- Ziemke T., Sharkey N.E., *A stroll through the worlds of robots and animals: Applying Jakob von Uexkull's theory of meaning to adaptive robots and artificial life*, „Semiotica”, t. 134, nr 1/4, 2001, ss. 701–746.

Oświadczenie

Niniejszym wyrażam zgodę na udostępnianie przez Archiwum Uniwersytetu Papieskiego Jana Pawła II w Krakowie mojej pracy doktorskiej.

Kraków, dnia

czytelny podpis autora

Oświadczenie kierującego pracą

Niniejsza praca została przygotowana pod moim kierunkiem i może być podstawą postępowania o nadanie jej autorowi (autorce) tytułu zawodowego.

Data

Podpis kierującego pracą

Oświadczenie autora pracy

Świadom odpowiedzialności prawnej oświadczam, że niniejsza praca dyplomowa została napisana przeze mnie samodzielnie i nie zawiera treści uzyskanych w sposób niezgodny z obowiązującymi przepisami.

Oświadczam również, że przedstawiona praca nie była wcześniej przedmiotem procedur związanych z uzyskaniem tytułu zawodowego w wyższej uczelni.

Oświadczam ponadto, że niniejsza wersja pracy (przedstawiona do obrony) jest identyczna z załączoną wersją elektroniczną

Data

Własnoręczny podpis autora pracy