

Inteligentny agent – teoria i perspektywy

MARCIN SKŁADANEK

W tekście tym chciałbym przybliżyć czytelnikowi pojęcie inteligentnego agenta (*intelligent agent*). Pojęcie to od przeszło dekady na dobre już zdomowione w badaniach na sztuczną inteligencję (AI – *Artificial Intelligence*), staje się dzisiaj równie istotne dla wielu dziedzin i obszarów badawczych związanych z informatyką (*computer science*), które na różnych poziomach podejmują refleksję nad tym nowym korpusem zagadnień. Co więcej, systematycznie zwiększa się liczba aplikacji wykorzystujących podejście agentowe.

Czym jest inteligentny agent? Zanim nieco dokładniej zdefiniuję to pojęcie, odwołam się do intuicji. Wyobraźmy sobie, że po włączeniu naszego domowego komputera wita nas przyjemny i ciepły głos naszego osobistego cyfrowego asystenta/ki (PDA – *personal digital assistant*), informujący nas o obecnych w naszej skrzynce pocztowej e-mailach, przypominający o wyznaczonych na dany dzień spotkaniach. Nie wymaga to oczywiście szczególnie wybujałej wyobraźni, ale pójdźmy dalej. Asystent, znając nasze zajęcia, zainteresowania, upodobania, przekonania polityczne, zaserwuje nam odpowiednie artykuły, wskaże linki do stron WWW, gdzie możemy pogłębić dane zagadnienie. Mało tego, jeśli wyrazimy taką potrzebę, agent skontaktuje się z innymi asystentami i umówi nas na wizytę u stomatologa, zamówi stolik w restauracji, a jeżeli będziemy późno w domu, przypomni naszym dzieciom o odrobieniu lekcji, a może nawet sprawdzi ich zadania domowe i poprawi błędy. Czy jest to niewybredna fantazja, daleka przyszłość, czy rzeczywistość dnia jutrzejszego? Mając na uwadze ogromne zainteresowanie technologią agentową wykazywane ze strony przemysłu, skłaniałbym się ku ostatniemu scenariuszowi. Tym bardziej że z agentami (bardziej lub mniej inteligentnymi) mamy do czynienia już dziś, a asystent programu Microsoft Word jest właśnie przykładem – mniej inteligentnego – agenta. Jako przykłady błyskotliwszych agentów wymienia się często systemy kontroli ruchu lotniczego czy złożone aplikacje internetowe. Trwają obecnie intensywne prace nad agentami mobilnymi oraz agentami konwersacyjnymi zaopatrzonymi w interfejs rozpoznawania mowy. Prawdziwym wyzwaniem są systemy składające się z wielu agentów połączonych interaktywną więzią – tzw. systemy multiagentowe (*multi-agent system*). Tu bowiem dokonuje się zrozumienia fenomenu interaktywności (wiedza pojmowana nie jako struktura, ale jako efekt złożonych interakcji) i sztucznej inteligencji rozproszonej (DAI – *Distributed Artificial Intelligence*).

Zaznaczyłem już wyżej, że pojęcie inteligentnego agenta pochodzi przede wszystkim z AI, choć od niedawna staje się również istotne w szerszej dziedzi-

nie, tzn. w ogólnie pojętej informatyce czy komputerystyce (*computer science*). W niniejszym tekście chciałbym – między innymi – podać kilka podstawowych argumentów, które przemawiają za tym, iż pojęcie to jest godne zainteresowania także teoretyków komunikacji czy teoretyków mediów. Dodajmy jednak szybko, że wielu z nich już wykazuje takie zainteresowanie, dlatego – uściślając – będę się starał wskazać, jakie obszary badań, jakie perspektywy teoretyczne implikuje zjawisko inteligentnego agenta. Warto jednak wyjść od tej definicji pojęcia, która funkcjonuje na gruncie nauki dla niego macierzystej.

Definicja agenta

Michael Wooldridge i Nicholas R. Jennings¹ rozróżniają słabe i mocne znaczenia pojęcia. To pierwsze nie budzi większych sporów – w słabszym znaczeniu agent byłby programem lub systemem komputerowym, który ma następujące własności:

- *autonomiczność* (*autonomy*) – agent jest zdolny do samodzielnego działania (sekwencja jego działań nie jest prostym schematem reakcji na impulsy płynące z zewnątrz); idzie tu przede wszystkim o samodzielne rozwiązywanie problemów;

- *zdolności społeczne* (*social ability*) – agent jest w stanie interaktywnie komunikować się z człowiekiem lub innymi agentami;

- *reaktywność* (*reactivity*) – agent reaguje na zmianę postrzeganego przez siebie „otoczenia” („otoczeniem” może być świat rzeczywisty, interfejs użytkownika, Internet, inni agenci etc.);

- *aktywność* (*pro-activeness*) – agent może wykonywać sekwencję działań ze względu na jakiś cel, wykazując przy tym inicjatywę.

Natomiast w silniejszym znaczeniu, które jednoznacznie identyfikujemy z badaniami nad sztuczną inteligencją, agent jest definiowany jako system komputerowy, który poza własnościami wymienionymi wyżej legitymuje się innymi, bardziej złożonymi cechami, które – co najważniejsze – jednoznacznie przypisujemy inteligentnym podmiotom ludzkim. Zatem agent byłby złożonym interaktywnym programem lub systemem komputerowym, któremu jesteśmy skłonni przypisać pewne przekonania, pragnienia, dążenia, działania skierowane na dany cel i środki, za pomocą których ten cel jest realizowany etc. Innymi słowy, jesteśmy skłonni przypisać mu stany intencjonalne. Wielu badaczy idzie nawet dalej, przypisując agentom również takie cechy, jak mobilność, prawdomówność i niesprzeczność (agent nie może przekazywać fałszywych informacji i podejmować sprzecznych ze sobą celów) czy wreszcie racjonalność.

Powstaje jednak pytanie: jak można przypisać stany intencjonalne, tradycyjnie przypisywane jedynie podmiotowi ludzkiemu, czemuś, co jest jedynie złożonym układem elektronicznym i w rzeczywistości przedstawia jedynie ciągi zer i jedynek, nie mając nawet pojęcia, co one oznaczają? Dodatkowe problemy rodzą takie predykaty, jak „inteligentny” czy „racjonalny”.

Założenie o intencjonalności

Kluczem do powyższych wątpliwości oraz elementem szczególnie ważnym w kontekście zagadnienia komunikacji człowiek – komputer jest tzw. nastawienie intencjonalne. Jest to strategia powszechnie przez nas stosowana, z której

wnioski filozoficzne wyprowadził Daniel C. Dennett – współczesny filozof identyfikowany z interdyscyplinarnym nurtem *cognitive science*.

Jego wyjściowe konstatacje są oczywiste, niezwykle proste i w zasadzie nie budzą poważniejszych polemik. Strategia „nastawienia intencjonalnego” to po pierwsze użycie języka mentalistycznego (języka „psychologii potocznej”) do opisu zachowania dowolnych obiektów. Dennett pisze: *Nastawienie intencjonalne jest strategią interpretacji zachowania jakiegoś bytu (osoby, zwierzęcia, wytworu, człowieka, czegośkolwiek) polegającą na traktowaniu go, jak gdyby był racjonalnym podmiotem, który „wybiera” takie a nie inne „działanie”, „biorąc pod uwagę” swoje „przekonania i chęci”. (...) Nastawienie intencjonalne jest postawą czy perspektywą, którą rutynowo przyjmujemy wobec siebie nawzajem, zatem przyjęcie tego nastawienia w stosunku do czegoś innego wydaje się zamierzonym antropomorfizowaniem*².

O pożyteczności i skuteczności stosowania nastawienia intencjonalnego w celu przewidywania działania dowolnego obiektu łatwo się przekonać, zestawiając ową strategię z innymi strategiami predykcji czy opisu – nastawieniem fizycznym oraz nastawieniem konstrukcyjnym. *Pierwsze jest po prostu standardową, pracochłonną metodą nauk fizycznych, w której wykorzystujemy wszystko to, co wiemy na temat praw fizyki i budowy interesujących nas problemów, do sformułowania naszych przewidywań. Kiedy twierdzę, że kamień wypuszczony z ręki spadnie na ziemię, przyjmuję nastawienie fizyczne. Nie przypisuję kamieniowi przekonani i chęci, lecz masę czy energię, a przewidywania opieram na prawie grawitacji. W przypadku rzeczy, które nie są ani ożywione, ani nie są wytworami człowieka, nastawienie fizyczne jest jedyną dostępną strategią przewidywań. (...) Budzik, który jest pewną konstrukcją (inaczej niż kamień), poddaje się bardziej wymyślnemu rodzajowi przewidywań – przewidywaniom z nastawienia konstrukcyjnego. Nastawienie konstrukcyjne jest wspaniałym skrótem, z którego wszyscy cały czas korzystamy. (...) Nie potrzeba rozkładać budzika, ważyć jego części i mierzyć elektrycznych napięć. Po prostu zakładam, że stanowi on pewną konstrukcję – konstrukcję, którą nazywam budzikiem – i że będzie zgodnie z nią działał. (...) Przewidywania z nastawienia konstrukcyjnego są ryzykowne w porównaniu z fizycznym (które są bezpieczne, ale żmudne), lecz jeszcze bardziej ryzykowne i jeszcze szybsze jest nastawienie intencjonalne. Może być ono postrzegane jako podgatunek nastawienia konstrukcyjnego, w którym analizowaną konstrukcją jest swego rodzaju podmiot. (...) Przyjęcie nastawienia intencjonalnego jest bardziej użyteczne – a w istocie, praktycznie nieuniknione – gdy wytwór, z jakim mamy do czynienia, jest znacznie bardziej skomplikowany niż budzik. Moim ulubionym przykładem jest komputer z programem szachowym... – pomyśl o nich jak o racjonalnych podmiotach, które chcą wygrać, znają zasady gry w szachy i wiedzą, jaka jest pozycja figur na szachownicy. W jednej chwili twój problem przewidywania i interpretacji ich zachowania staje się zdecydowanie łatwiejszy, niż byłby wtedy, gdybyś próbował skorzystać z nastawienia fizycznego lub konstrukcyjnego*³.

Widać więc, że cała kwestia sprowadza się do kilku bardzo prostych konstatacji. Dysponujemy trzema językami przewidywania i opisu: (1) językiem fizycznym, (2) językiem funkcjonalnym (konstrukcyjnym), (3) językiem mentalistycznym. Języki te stosujemy w zależności od natury opisywanego obiektu, a dokładniej – w zależności od aspektu danego przedmiotu. Wiedzę, przekonania i pragnienia czło-

wieka teoretycznie można opisać za pomocą każdego z powyższych języków, ale oczywiście najskuteczniejszy jest język mentalistyczny, przynajmniej dopóki neurofizjologia, posługująca się językiem fizycznym i funkcjonalnym, nie poczyni daleko posuniętego postępu.

Dlaczego jednak powyższe uwagi są tak istotne dla badań nad AI oraz dla komunikacji człowiek – komputer? Musi coś w tym być, skoro Daniel C. Dennett stał się pierwszoplanową postacią „filozoficznego zaplecza” AI, a sam Marvin Minsky mówił o nim, że jest naszym najlepszym współczesnym filozofem. Przede wszystkim chodzi o kwestię następującą: skoro najłatwiejszą i najbardziej ekonomiczną strategią przewidywania działań i opisu istot, przedmiotów, systemów jest zastosowanie wobec nich nastawienia intencjonalnego, tzn. przypisanie im przekonań, pragnień, celów etc., a ponadto współczesne systemy komputerowe przekroczyły już ów stopień złożoności, przy którym laik może stosować nastawienie konstrukcyjne, to dlaczego nie pójść dalej i nie podjąć się rzeczywiście reprezentacji i symulowania owych przekonań, pragnień, czy celów.

Zatem, podsumowując: inteligentny agent to system lub program, wobec którego stosujemy nastawienie intencjonalne i ze względu na złożoność tego układu jest to jedyna (o ile nie jesteśmy programistami lub elektronikami) strategia predykcji. Dlatego „wewnętrzne stany mentalne” inteligentnego agenta traktujemy jako heurystyczne konstrukcje, których zachodzenie zakładamy jedynie w celu przewidywania jego zachowania. Skuteczność owego przewidywania jest tu wartością nadrzędną. Jest to więc po prostu pożyteczna strategia, której zasadność musi być każdorazowo testowana.

Poziomy refleksji nad agentami

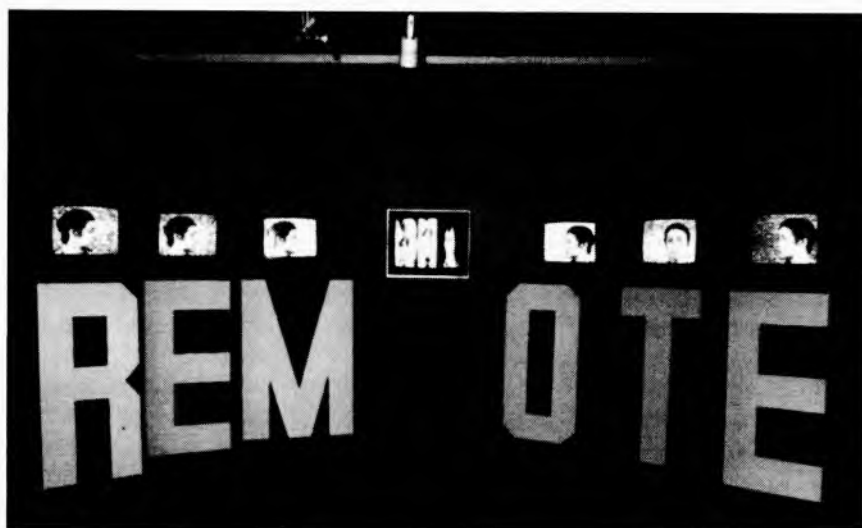
Powróćmy do sposobów definiowania i zakresu problematyki inteligentnego agenta. Standardowo na gruncie AI i szerzej informatyki problem ten sprowadza się do trzech grup zagadnień:

- *teoria* – sprecyzowanie, konceptualizacja i formalizacja problemu,
- *architektura* – metodologia budowy systemu,
- *języki programowania (agent languages)* – narzędzia programowania systemów tworzonych w technologii agentowej.

Powyższe trzy grupy zagadnień wyznaczają trzy nierozdzielnie ze sobą związane obszary badań. Architektura systemowa oraz języki programowania są obszarami ściśle inżynierskimi. Wysiłki konstruktorów sprowadzają się do „ucieleśnienia”, realizacji postulatów sformułowanych przez teoretyków. Chodzi więc przede wszystkim o metodologię budowania hardware’u i software’u oraz konkretne narzędzia stojące za taką metodologią. W te problemy, z oczywistych względów, nie będę się zagłębiał. Warto jednak wspomnieć, że coraz większa liczba badaczy postrzega właśnie języki agentowe (*agent languages*) oraz programowanie zorientowane agentowo (*agent-oriented programming*) jako nowy paradygmat, który jest rozwinięciem programowania obiektowego. W tym kontekście „agent” to rozwinięcie pojęcia „klasy” (*class*) – podstawowego atrybutu programowania obiektowego. Warto tutaj zauważyć, że takie rozumienie agenta jest właśnie rozumieniem „słabym”, podczas gdy np. PDA odpowiada rozumieniu mocne.



The Way, T. Waliczky (1994)



Remote Faces, M. Rogala (1986)

Szczególne zainteresowanie, zwłaszcza ze strony teoretyków komunikacji, mediów, a dalej psychologów, filozofów, językoznawców, kognitywistów, należy się jednak – w moim przekonaniu – przede wszystkim teorii inteligentnego agenta. Na tym bowiem obszarze dochodzi do ciekawych rozstrzygnięć implikujących nowe przestrzenie badań, wymuszających na nas zmianę dotychczasowych paradygmatów. Aby jednak przyjrzeć się nieco bliżej teorii agenta, warto – na poziomie jak najbardziej ogólnym i dalekim od wszelkich formalnych zawiłości – powiedzieć, czym jest i co zakłada modelowanie i symulacja komputerowa. Otóż modelowanie, czyli tworzenie modelu będącego podstawą formalnego zapisu danej rzeczywistości, sprowadza się do następujących czynności:

- dokonanie wyboru aplikowanej dziedziny (*application domain*), czyli świata, fragmentu rzeczywistości (prawdziwej lub nie), którą zamierzamy poddać modelowaniu;

- dojście do pewnych nieformalnych, semantycznych intuicji interpretujących ową rzeczywistość (*informal semantic*) – jest to *de facto* proces tworzenia pojęciowego (ale nie formalnego) obrazu danego świata (semantyka, interpretacja, model dla systemu formalnego);

- skonstruowanie formalnej notacji dla naszych intuicji (*formal representation*) – logiczny system formalny.

Podczas gdy proces modelowania, jak zresztą doskonale to widać, jest domeną teorii, to symulacja jest już dziedziną ściśle praktyczną i poddaną wymogom projektowania i inżynierii, zarówno hardware'u, jak i software'u. W procesie symulacji formalny zapis naszych intuicji, będący efektem modelowania, należy przełożyć na język „rozumiały” dla komputera (język programowania) oraz zaprojektować maszynę, która będzie zdolna działać według tego programu (architektura systemowa).

Widać więc, że o ile języki agentowe i architektura agenta wiążą się z procesami symulacji, to teoria jest procesem modelowania. W ścisłym znaczeniu teoria agenta ma na celu sformułowanie formalnej notacji wszelkich działań i właściwości, jakie przypisujemy agentowi, bo one właśnie są aplikowaną dziedziną, która poddaje się procesowi modelowania. Podstawowym paradygmatem formalizacji przekonań, pragnień, chęci etc. jest semantyka światów możliwych – standardowy model logik modalnych oraz będących ich pochodną logik epistemicznych, logik percepcji, logik deontycznych, logik temporalnych etc. Nie jest tutaj możliwe rozwinięcie powyższych wątków. Zaznaczę jedynie, iż sprowadzają się one – szczególnie jeśli idzie o inteligentnego agenta – do tzw. inżynierii wiedzy (*Knowledge Engineering*). Jej podstawowy korpus zagadnień to m.in.: analiza struktury wiedzy (hierarchia, schematy, reguły, dane, wiedza jawna i ukryta); metody rozumowania; rozwiązywanie problemów (różnice pomiędzy algorytmami a heurystykami); procesy decyzyjne; innowacja etc. Skala wyzwań jest tu ogromna, tak jeśli idzie o samo zrozumienie, czym jest wiedza, jak i skonstruowanie systemów zdolnych ją symulować. Problem polega przede wszystkim na tym, że kluczem nie jest tutaj sama logika czy inżynieria oprogramowania, ale komunikacja, a nauka o tej ostatniej dopiero się tworzy. Symulowanie złożonych procesów komunikacyjnych, uwzględniające przy tym tak podstawowe jej aspekty, jak znaczenie, kontekst, wiedza kulturowa, struktura konwersacji, wymaga sięgnięcia po metody i narzędzia wykraczające poza te, których dostarcza sama

algorytmika, matematyka, logika czy elektronika. Dzisiejszej informatyce potrzeba czegoś więcej. Owym „więcej” jest kognitywistyka, psychologia, filozofia, socjologia, teoria mediów, estetyka i *design* etc. Dlatego też coraz częściej mówi się o „miękkiej matematyce”, „informatyce, jako nauce empirycznej”. Bodaj najbardziej czytelnym świadectwem zmian, jakie stają się udziałem nauki o komputerach, jest – obok wyraźnie interdyscyplinarnego charakteru – szczególnie nacisk na zrozumienie fenomenu interaktywności i budowę systemów rozproszonych.

Aplikacje i perspektywy

Uznawszy to za niezbędne wprowadzenie do problematyki, starałem się przybliżyć samo pojęcie inteligentnego agenta oraz – pozostając jedynie na poziomie rudymentów – nakreślić obszar, w jakim dokonuje się jego eksploracji. Pozostałem przy tym na poziomie elementarnym i abstrakcyjnym.

Oczywiście tym, co jest najbardziej interesujące, szczególnie dla teoretyków komunikacji, mediów czy cyberkultury, są aplikacje technologii agentowej. Dziedzina owych aplikacji jest znaczna i różnorodna, toteż trudno dokonać w niej jakiejś wyczerpującej i ostatecznej kategoryzacji. Wydaje się jednak, iż sensowne jest – szczególnie na użytek niniejszego szkicu – oddzielenie od całej reszty aplikacji tych, które określa się jako „agent użytkownika” (*user agent*). Sądzę, iż właśnie tego rodzaju zastosowania technologii agentowej stanowią rzeczywiste wyzwanie dla programistów oraz właśnie w ich konstruowaniu dochodzi do najbardziej inspirujących interakcji pomiędzy inżynierią oprogramowania, inżynierią wiedzy, kognitywistyką, psychologią, filozofią, teorią komunikacji czy teorią mediów. Tutaj też – w stopniu bodaj największym – objawiają swoją użyteczność, jeśli nie niezbędność – kontrowersyjne, bądź co bądź, badania nad sztuczną inteligencją. Przede wszystkim dlatego, że wykorzystanie ich w konstruowaniu inteligentnych interfejsów użytkownika pozwala spojrzeć na AI nie jak na dziedzinę owładniętą maniakałnym dążeniem do stworzenia sztucznego człowieka i wykazania tym samym, że ludzka inteligencja daje się zamknąć w modelu obliczeniowym, ale jak na obszar, na którym dokonuje się badań i eksperymentów wydatnie przyczyniających się do rozwoju technologii informatycznej oraz zrozumienia ludzkich procesów poznawczych.

Teza, której obrony chciałbym się tutaj podjąć, brzmi więc następująco: w konstruowaniu interfejsów podług nadrzędnych kryteriów wydajności i efektywności niezbędne są nie tylko technologia, badania nad procesami kognitywnymi człowieka, wizualizacja danych, multimediami, estetyką i *designem*, ale również AI. Czyli, *de facto* – dla usprawnienia komunikacji człowiek – komputer – potrzebne jest podejście agentowe, ponieważ przy szerokim rozumieniu samego pojęcia AI jest poddyscypliną nauki o komputerach, która zajmuje się konstruowaniem agentów.

Nie zgadzam się tym samym ze zdaniem Jarona Laniera, który w artykule pod znamienym tytułem *Agents of Alienation*⁴ twierdzi, iż idea inteligentnego agenta jest błędna i zła (*wrong and evil*). Lanier postrzega w niej czynnik, który zakłóca nadrzędną – jego zdaniem – funkcję rewolucji informatycznej, jaką jest łączenie ludzkości, tworzenie ściślejszych więzi komunikacyjnych pomiędzy



Legible City, J. Show (1988)

jednostkami ludzkimi. Ponadto, broniąc wyjątkowości człowieka, uznaje, że agent nie wspomaga użytkownika, ale go ogranicza i zmusza do redefiniowania swojej pozycji i statusu. Otóż nawet jeśli istnieją pewne zagrożenia związane z rozpowszechnianiem się inteligentnych agentów, to – uważam – że nie powinny one przesłaniać oczywistych zalet, jakie niesie obecność wśród nas tego rodzaju istot. To, że samochody przyczyniają się do zanieczyszczenia środowiska, nie znaczy przecież, że mamy przesiąść się na rowery, a jedynie to, że powinniśmy użyć wszelkich dostępnych środków, aby owe zanieczyszczenia minimalizować. Jestem skłonny przyznać Lanierowi rację, kiedy twierdzi, że *napisanie dobrego interfejsu użytkownika dla tak skomplikowanych zadań, jak znajdowanie i filtrowanie informacji, jest znacznie trudniejsze od skonstruowania inteligentnego agenta*. Nie mogę się jednak z nim zgodzić, że *agenci są dziełem leniwych programistów*, ponieważ są to często ci sami ludzie, którzy robią dobre interfejsy. Rzecz w tym, że prace nad agentami i prace nad interfejsami użytkownika nie są ze sobą sprzeczne, ale – i na to postaram się dalej położyć nacisk – są to często prace jak najbardziej komplementarne. Uogólniając – skupię się dalej na statusie, roli oraz znaczeniu inteligentnych agentów w niezwykle dynamicznie ostatnio rozwijającej się i stwarzającej niezwykle perspektywy dziedzinie tzw. interakcji człowiek –

komputer (*HCI – Human-Computer Interaction*), a w szczególności na tym, w jaki sposób technologia agentowa wspomaga badania nad udoskonaleniem efektywności i wydajności interfejsów.

Oczywiste jest, że nie wszystkie programy i systemy muszą legitymować się interfejsem inteligentnym i zdolnym do samodzielnego działania. Taka potrzeba dotyczy przede wszystkim programów o określonych zadaniach. Najczęściej, jako te aplikacje, w których inteligentny, agentowy interfejs użytkownika najlepiej się sprawdza, wymienia się następujące programy:

- aplikacje służące do wyszukiwania, filtrowania i monitorowania informacji
- jest to szeroka gama inteligentnych wyszukiwarek WWW, specjalistycznych serwisów informacyjnych czy wreszcie osobistych asystentów informacyjnych (PDA); agent mający dostęp do wielu źródeł informacji chroni użytkownika przed zgubieniem się w *wirtualnych przestrzeniach sieci WWW*, kierując jego nawigacją lub wręcz – ponieważ zna jego zapotrzebowania – wyluskując zbiór wartych uwagi informacji; ważnym zastosowaniem tego typu agentów są inteligentne interfejsy bogatych i złożonych systemów baz danych;

- programy edukacyjne – wszelkie multimedialne encyklopedie, kursy, poradniki; agent może tu występować w roli nauczyciela czy trenera, ułatwiając znacznie przyswajanie wiedzy;

- programy eksperckie i diagnostyczne – aplikacje z wbudowanym systemem wiedzy z wąskiej dziedziny;

- programy monitorujące działanie systemów i innych programów, systemy eksperckie – tutaj inteligentny agent zastępuje człowieka w nadzorowaniu działania danego systemu, informuje o przebiegu jego działania, czy jest pośrednikiem pomiędzy nami a systemem; pole do popisu mają tu konstruktorzy różnorodnych „agentów domowych” – znany entuzjasta inteligentnych, cyfrowych agentów Nicholas Negroponte pisze: *Jeżeli lodówka zauważy, że kończy się mleko, może „poprosić” samochód, aby nam przypominał, że należy je kupić po drodze do domu. Obecny sprzęt domowy ma za mało możliwości obliczeniowych. (...) Dom współczesnego Amerykanina ma prawdopodobnie ponad sto mikroprocesorów. Ale nie są one ze sobą połączone. (...) Brak komunikacji między różnego rodzaju urządzeniami jest m.in. wynikiem bardzo prymitywnego interfejsu każdego z tych urządzeń. (...) Instrukcja obsługi powinna zniknąć. Najlepszym nauczycielem, jak używać urządzenia, powinno być samo urządzenie. (...) Maszyna wyposażona w pewną wiedzę o użytkowniku (jest lewo- czy praworęczny, dobrze słyszy, czy ma kłopoty ze słuchem, nie ma cierpliwości do urządzeń) może być znacznie lepszym instruktorem własnych operacji niż jakikolwiek dokument* ⁵.

CSCW (*Computer-Supported Cooperative Work*) – komputerowo wspomagane współdziałanie jest dziedziną nową, ale niezwykle dynamicznie się rozwijającą i mającą już na swym koncie pierwsze sukcesy – oparte na podejściu agentowym systemy kontroli ruchu lotniczego. CSCW jest doskonałym przykładem interdyscyplinarności, o jakiej wspominałem wyżej, a jaką wymuszają badania nad systemami inteligentnymi (inżynieria wiedzy): *CSCW była – i jest nadal – jedną z najbardziej eklektycznych dziedzin badań, jakie kiedykolwiek istniały. Wielkie letnie konferencje międzynarodowe CSCW przyciągają przedstawicieli socjologii, lingwistyki, kognitywistyki, psychologii, komputerystyki, projektowania systemów, inżynierii systemów, filozofii, zarządzania i zapewne jeszcze innych* ⁶.

Dlaczego inteligentny i autonomiczny (tj. agentowy) interfejs jest nam potrzebny? Odpowiedź na to pytanie ograniczę do trzech – w moim przekonaniu – rozstrzygających argumentów: po pierwsze, zarówno dzisiejsze programy, jak i ich tradycyjne interfejsy stają się zbyt skomplikowane, a podejście agentowe daje bodaj najwięcej szans na rozwiązanie tego problemu; po drugie, jedynie inteligentne interfejsy agentowe są zdolne do elastyczności co do przyzwyczajęń, potrzeb i preferencji użytkownika; wreszcie – po trzecie, badania pokazują, iż interakcja człowiek – komputer jest ufundowana na tej samej dynamice, co społeczne interakcje międzyludzkie ⁷.

Wszystkie powyższe argumenty przemawiają za tym, że podejście agentowe, będąc dużym wyzwaniem dla programistów, jest jednocześnie paradygmatem, który istotnie przyczynia się do postępu w rozwoju interfejsów. Jest to najlepszy dowód na to, że oparcie interakcji człowiek – komputer na technologii agentowej nie przekreśla wypracowanych już strategii projektowania interfejsów, gdyż nadal wyznacznikami „dobrego” interfejsu pozostają efektywność i użyteczność, ale istotnie owe wymogi wspomaga. Inteligentny interfejs agentowy rzeczywiście sprawia – i sądzę, iż jest to duża wartość – że interfejs jest pomiędzy mną a oprogramem. Jest autentycznym pośrednikiem pomiędzy użytkownikami i aplikacjami. Jedynie bowiem interfejs, który legitymuje się pewnym stopniem autonomii i potrafi podejmować i realizować postawione mu zadania, jest w stanie spełnić oczekiwania, jakie zwykle w interfejsie pokładamy – chciałbym, aby mój interfejs uczył się moich nawyków i zapotrzebowań, co więcej, chciałbym, aby podpowiadał mi możliwe rozwiązania, a często wręcz wyręczał mnie w wykonywaniu pewnych uciążliwych działań. Raz jeszcze odwołam się do cennych uwag Nicholasa Negroponte: *Najlepsze porównanie, jakie przychodzi mi do głowy, gdy myślę o środkach porozumiewania się człowieka z komputerem, to dobry angielski „lokaj”. Ten „agent” odpowiada na telefony, rozpoznaje rozmówców, przeszkadza tylko w odpowiednim momencie i może nawet kłamać w żywe oczy w naszym imieniu. (...) nie chodzi tu o współczynnik inteligencji, ale o praktykę używania inteligencji dla naszego dobra. (...) Koncepcja „agenta” zawiera ideę człowieka pomagającego innemu człowiekowi, gdyż często wiedza połączona jest ze znajomością naszych upodobań. (...) Agent interfejsu musi się uczyć i rozwijać w czasie, podobnie jak sekretarki czy asystenci. (...) Trzeba także zdawać sobie sprawę, że agentowe podejście do interfejsu jest całkowicie różne od obecnej mody nawigowania po Internecie za pomocą przeglądarki Mosaic lub Natescape. Maniacy Internetu mogą się w nim poruszać, wynajdować ogromne informacje i zdobywać wiedzę oraz wchodzić w różne grupy socjalne. To szczególnie rozpowszechnione zjawisko nie zaniknie, ale jest to tylko jeden rodzaj zachowania – raczej bezpośrednie manipulowanie niż oddawanie uprawnień* ⁸.

Na zakończenie warto poświęcić parę słów ostatniemu argumentowi, kto wie czy nie najistotniejszemu. Sprowadza się do wagi antropomorfizmu w projektowaniu inteligentnych interfejsów – czyli nadawania im cech ludzkich czy ostrożniej – cech postrzeganych przez konstruktorów jako ludzkie ⁹. Antropomorfizm w odniesieniu do agenta jest implikowany przez mocną definicję agenta oraz nastawienie intencjonalne i jest czymś naturalnym dla każdego agenta spełniającego ową definicję. Zalety antropomorficznego podejścia w projektowaniu interfejsów, obok tych, które mają charakter funkcjonalny, a o których pisałem wyżej,

mają przede wszystkim podłoże psychologiczne. Zasada jest prosta: wzorem dla komunikacji człowiek – komputer jest komunikacja interpersonalna. Już samo zastąpienie edycji tekstowej animowaną postacią sprawia, że komunikacja przebiega sprawniej. Celem jest więc skonstruowanie takiej wirtualnej reprezentacji człowieka, która jest zdolna imitować konwersację na wielu poziomach – począwszy od mowy, gestów, mimiki etc.

Podsumowanie

W „osobie” inteligentnego agenta – zarówno sama informatyka, jak i stale powiększająca się dziedzina jej aplikacji – otrzymuje potężne narzędzie, o którego mocy już mamy sposobność się przekonać. Oczywiście, nie sposób tutaj nie dodać, że pierwsze lata XXI wieku to dopiero początek naszej „przygody z agentami”. Zakres możliwych eksploracji tego zagadnienia jest ogromny i wiele pozostało jeszcze do zrobienia. Ale jednocześnie, i to według mnie jest podstawowa zaleta owego pojęcia i tego, co ono implikuje, skupia ono wszystkie najważniejsze wyzwania, jakie stoją zarówno przed samą informatyką (*computer science*), jak i obszarem jej zastosowań. Tymi najistotniejszymi, powiązаныmi wszak ze sobą, wyzwaniami są w moim przekonaniu:

– *interdyscyplinarność* – przełamanie ponewtonowskiego paradygmatu podziału nauk, na nauki przyrodnicze, formalne i społeczne, w celu zrozumienia fenomenu poznania, języka i wiedzy oraz tworzenia przyjaznego i funkcjonalnego cyfrowego otoczenia człowieka;

– *interaktywność* – postrzeganie interaktywności jako podstawowego czynnika stymulującego postęp informatyki oraz rozwój teorii i praktyki modelowania komputerowego (sieci, systemy o architekturze rozproszony, systemy multiagentowe)¹⁰;

– *interakcja człowiek – komputer (HCI)* – projektowanie wydajnych, efektywnych i „przyjaznych” interfejsów użytkownika.

MARCIN SKŁADANEK

¹ M. Wooldridge, N. R. Jennings, *Intelligent Agents: Theory and Practice*, „Knowledge Engineering Review”, October 1994.

² D. C. Dennett, *Natura umysłów*, tłum. W. Turpeński, Warszawa 1997, s. 39-40, por. także: D. C. Dennett, *Consciousness Explained*, Boston 1991.

³ D. C. Dennett, dz. cyt., s. 40-44.

⁴ J. Lanier, *Agents of Alienation*, artykuł dostępny pod adresem: www.well.com/user/jaron/agentalien.html

⁵ N. Negroponte, *Cyfrowe życie*, tłum. M. Łakomy, Warszawa 1997, s. 176-177.

⁶ K. Devlin, *Zegnaj Kartezjuszu*, tłum. B. Stanosz, Warszawa 1999, s. 253.

⁷ Y. Takeuchi, Y. Katagiri Clifford Nass, B. J. Fogg, *A Cultural Perspective in Social Interface*, artykuł dostępny na stronie: www.mic.atr.co.jp/yugo/publishment/generalpaper.ps.gz

⁸ N. Negroponte, dz. cyt. s. 125-131.

⁹ J. J. Morgan, *Submission for Interacting with Computers, Antropomorphism on Trial*, artykuł dostępny pod adresem: www.ccc.hw.ac.uk/~jeff/files/anthro.ps.gz

¹⁰ P. Wegner, *Why interaction is more powerful than algorithms*, Comms. Of ACM 40(5), 1997.