

Umjetna inteligencija danas

Kovač, Lorena

Master's thesis / Diplomski rad

2015

Degree Grantor / Ustanova koja je dodijelila akademski / stručni stupanj: **University of Rijeka, Faculty of Humanities and Social Sciences / Sveučilište u Rijeci, Filozofski fakultet u Rijeci**

Permanent link / Trajna poveznica: <https://um.nsk.hr/um:nbn:hr:186:606497>

Rights / Prava: [In copyright](#) / [Zaštićeno autorskim pravom.](#)

Download date / Datum preuzimanja: **2024-04-25**



Repository / Repozitorij:

[Repository of the University of Rijeka, Faculty of Humanities and Social Sciences - FHSSRI Repository](#)



SVEUČILIŠTE U RIJECI
FILOZOFSKI FAKULTET U RIJECI
DIPLOMSKI STUDIJ POLITEHNIKE I INFORMATIKE

STUDENT: Lorena Kovač

NASLOV RADA:

UMJETNA INTELIGENCIJA DANAS

MENTOR:
Prof.dr.sc. Mile Pavlić

U Rijeci,

Sadržaj

1. UVOD	1
2. UMJETNA INTELIGENCIJA.....	3
2.1. Smjerovi u istraživanju umjetne inteligencije.....	5
2.1.1. Kognitivistička škola	6
2.1.2. Logička škola.....	7
2.1.3. Bihevioristička škola.....	8
2.2. Umjetna inteligencija i filozofija.....	10
2.3. Mogu li strojevi misliti?.....	11
3. KLASIČNE METODE U UMJETNOJ INTELIGENCIJI	15
3.1. Ekspertni sustavi (sustavi temeljeni na znanju)	16
3.1.1. Prednosti ekspertnih sustava	18
3.1.2. Nedostaci ekspertnih sustava.....	19
3.2. Strojno učenje	20
3.3. Rudarenje podataka.....	22
4. MODERNE METODE U UMJETNOJ INTELIGENCIJI.....	25
4.1. Umjetne neuronske mreže	26
4.1.1. Primjer modela umjetnog neurona.....	29
4.1.2. Duboko učenje	32
4.2. Evolucijsko računalstvo	33
4.3. Inteligentni agenti.....	36
4.3.1. Osnovna građa inteligentnih agenata	38
4.3.2. Primjena inteligentnih agenata.....	39
4.3.3. Softverski agent – softbot - bot	40
4.3.4. Višeagentski sustavi i raspodijeljena inteligencija.....	41
4.4. Big data – veliki podaci.....	43
5. KOMERCIJALNA PRIMJENA UMJETNE INTELIGENCIJE	45
5.1. Primjena inteligentnih sustava.....	46
5.2. Google i razvoj umjetne inteligencije.....	51
6. PERSPEKTIVE U RAZVOJU I PRIMJENI UMJETNE INTELIGENCIJE	54
6.1. Pravci novih istraživanja u umjetnoj inteligenciji	54
6.2. Istraživanja inspirirana radom mozga – neuromorfna arhitektura	55
6.2.1. Europski neuromorfni projekti	56
6.2.2. Neuromorfni projekt u SAD-u.....	58
6.3. Nova promišljanja umjetne inteligencije – projekt na MIT-u.....	60
6.4. Predviđanja o primjeni umjetne inteligencije.....	61
7. ZAKLJUČAK.....	64
8. METODIČKA OBRADA NASTAVNE JEDINICE	65
LITERATURA	73
POPIS SLIKA.....	77

Sažetak

U radu je obrađena umjetna inteligencija i njezina primjena danas. U svijetu velikih podataka i poplave informacija sustavi koji primjenjuju umjetnu inteligenciju postaju sve značajniji. Krajem dvadesetog stoljeća, usvajanjem bottom-up pristupa u razvoju umjetne inteligencije započinju novi pravci istraživanja. Obilježeni su značajnim porastom ulaganja u razvoj inteligentnih sustava, koje traje i danas. U radu su opisane različite metode kojima se ostvaruje umjetna inteligencija. One uključuju ekspertne sustave, strojno učenje, umjetne neuronske mreže, evolucijsko računalstvo te koncept inteligentnog agenta. Prikazani su primjeri uporabe sustava temeljenih na umjetnoj inteligenciji poput virtualnog asistenta, sustava za preporuke, autonomnih vozila. Posebno je istaknut Google, najveći komercijalni sustav umjetne inteligencije ikada izgrađen. Opisana su nova istraživanja umjetne inteligencije primjenom računalne arhitekture inspirirane radom mozga. Spomenuta su očekivanja primjene umjetne inteligencije u personaliziranim pametnim domovima povezivanjem predmeta u “*internet stvari*”. Na kraju se navodi kako konačni cilj, stvaranje strojeva koji misle, još nije postignut. Nastojanja da se taj cilj ostvari ne izazivaju samo podsmijeh, već zabrinutost javnosti za budućnost čovječanstva. Poticaj za brigu je mogućnost razvoja generalne umjetne inteligencije u desetljećima koja slijede.

Ključne riječi: umjetna inteligencija, umjetne neuronske mreže, inteligentni agent, Google, neuromorfna znanost, internet stvari, generalna umjetna inteligencija.

Abstract

This paper deals with artificial intelligence and its application today. In the world of big data and the flood of information systems using artificial intelligence are becoming increasingly important. At the end of the twentieth century a bottom-up approach to the development of artificial intelligence started new directions in research. They were marked by a significant increase in investment in the development of intelligent systems, which continues today. The paper describes various methods achieved in field of artificial intelligence. Those methods include expert systems, machine learning, artificial neural networks, evolutionary computing, and the concept of an intelligent agent. Given are examples of systems based on artificial intelligence such as a virtual assistant, customers recommendation system, self-driving car. Particularly emphasized is Google, the largest commercial system of artificial intelligence ever built. Described are current research projects in artificial intelligence which are using computer architecture inspired by the brain. There is also expectation of wide usage of artificial intelligence in the personalized smart homes by linking objects in the "internet of things". In the end it is stated that the final goal, the creation of machines that thin, has not yet been reached. Yet today, these efforts do not cause laughs, but the concern about aftermath of humanity if we reach that goal. Boost for concern is the possibility of the development of general artificial intelligence in the decades that follow.

Keywords: artificial intelligence, artificial neural networks, intelligent agent, Google, neuromorphic science, internet of things, general artificial intelligence.

1. UVOD

Nastojanja da se izgradi umjetna inteligencija doživljavaju posljednjih godina strelovit uspon obilježen brojnim investicijama u primjeni i razvoju inteligentnih sustava na područjima poput komunikacija, trgovine, zdravstvenih usluga, pretraživanja interneta, sustava za preporuke i sl. Dok se brojna znanstveno fantastična djela zabavljaju kreiranjem nekog oblika svjesne umjetne inteligencije, čini se da su prvi znaci inteligentnih strojeva već su prisutni u računalima koja uče, prepoznaju objekte i sve bolje koriste govorni jezik u komunikaciji s ljudima.

U radu se obrađuje umjetna inteligencija danas, odnosno ostvarena sposobnost računala da se ponašaju na inteligentan način. U drugom poglavlju je definiran pojam umjetne inteligencije, zatim su navedeni različiti smjerovi istraživanja proizašli iz nastojanja da se ostvare strojevi koji misle. Također je opisan filozofski pogled na mogućnost stvaranja jake umjetne inteligencije te je opisan Turingov test kojim se nastoji dokazati je li ispitano računalo sposobno za komunikaciju s ljudima, odnosno razmišljati. Prikazana je i obrnuta inačica testa kojom se danas koristimo pri automatskom raspoznavanju računala i ljudi. Treće poglavlje obrađuje klasične metode korištenja umjetne inteligencije poput ekspertnih sustava, strojnog učenja i tehnike rudarenja podataka. U četvrtom su poglavlju opisane moderne metode umjetne inteligencije temeljene na bottom-up pristupu, gdje se metodama poput umjetnih neuronskih mreža nastoji od osnovnih elemenata izgraditi kompleksno, inteligentno ponašanje. Takav pristup temelji se na dinamičkoj interakciji softverskih ili robotskih agenata smještenih u stvarnom okruženju. U osnovnim sastavnicama opisana je metoda korištenja genetskih algoritama kao načina kreativnog traženja rješenja kojeg izvode algoritmi. Peto se poglavlje bavi današnjom primjenom inteligentnih sustava, prikazom brojnih praktičnih primjera, a posebno je izdvojen Google kao najveći komercijalni ulagač i predhodnik u razvoju umjetne inteligencije.

Posljednje, šesto poglavlje, posvećeno je novim pravcima u istraživanju i primjeni umjetne inteligencije. Uskoro se primjerice očekuju veoma dobri rezultati u primjeni neuromorfne znanosti koja povezuje računalnu znanost i rad biološkog mozga stvaranjem nove generacije računalnih čipova i nove računalne arhitekture s naglaskom na paralelizmu. Obrazložen je projekt novog promišljanja umjetne inteligencije pokrenut na MIT-u, te najavljeni dolazak tehnologija poput interneta stvari kojom će se povezati predmeti koji se nose i kućanski aparati.

Glavni cilj rada je opisati tehnologiju umjetne inteligencije koja je danas prisutna te navesti primjere u kojima se ona koristi. Razlog za odabir teme je činjenica kako su posljednjih godina u naglom porastu razvoj strojnog učenja i primjena inteligentnih sustava, a ipak šira javnost o tome zna vrlo malo. Danas možemo pouzdano reći da će umjetna inteligencija obilježiti sljedeća desetljeća, pokrenuti novu tehnološku revoluciju i izmjeniti naše živote na način koji u ovom trenutku ne možemo sasvim sagledati.

2. UMJETNA INTELIGENCIJA

Umjetna inteligencija (UI), jest sposobnost digitalnog računala ili računalno-kontroliranog robota da izvodi zadaće obično povezane uz inteligentna bića. (Copeland, 2014)

Polje umjetne inteligencije, skraćeno UI, pokušava razumjeti **inteligentne entitete**. Tako je jedan od razloga za proučavanje ovog područja i to da bolje razumijemo sebe. Međutim, za razliku od filozofije ili psihologije, koje također proučavaju inteligenciju, polje umjetne inteligencije teži **izgradnji umjetnih entiteta**, kao i njihovom razumijevanju.

“Umjetna inteligencija propituje jednu od konačnih zagonetki. Kako je moguće da spor, maleni mozak, biološki ili elektronički, može percipirati, razumjeti i predviđati svijet, te manipulirati svijetom mnogo većim i mnogo kompleksnijim nego što je on? Kako da izgradimo nešto s takvim svojstvima? Ta su pitanja teška, ali, za razliku od putovanja brzinom većom od brzine svjetlosti ili antigravitacijskog uređaja, istraživač na području umjetne inteligencije ima čvrste dokaze da je zadatak moguće ostvariti. Sve što trebamo učiniti je pogledati u zrcalo da bismo vidjeli primjer inteligentnog sustava” (Norvig i Russell, 1995, str 3).

Dva su osnovna cilja kojima teži većina istraživanja umjetne inteligencije. Prvi, najvažniji cilj, je izgradnja inteligentnih strojeva. Drugi cilj je shvaćanje prirode inteligencije ili onoga što psiholozi nazivaju g-faktor u testovima inteligencije (pokušaj da se izmjeri generalna inteligencija koja se prostire kroz različite domene ljudskog djelovanja). Oba cilja imaju u svojoj biti potrebu za definiranjem pojma inteligencije.

Generalna umjetna inteligencija GUI (engl. AGI, Artificial General Intelligence) je pojam kojim se opisuju istraživanja usmjerena prema kreiranju strojeva sposobnih za generalno inteligentne radnje. Ključno je da generalna umjetna inteligencija mora biti u stanju ovladati različitim područjima i naučiti nova područja s kojima se prije nije suočila. Nije pri tom potrebno posjedovati jednake sposobnosti u svim domenama. Kada u tom kontekstu kažemo kako ljudski um ili GUI imaju generalnu svrhu, ne mislimo da se može riješiti bilo koji problem u svekolikim domenama, već da postoji potencijal rješavanja bilo kojeg problema u bilo kojoj domeni dobije li se pravilno iskustvo. “Dokle god se ne dokaže da GUI nije moguć, ostaje legitimna tema

istraživanja. A uzevši u obzir kompleksnost problema, nema razloga očekivati da GUI postigne cilj u kratkom periodu, pa će sve popularne teorijske kontroverze vjerojatno postojati i nakon što se ostvari GUI, kako je planirano. Činjenica da u istraživanjima ima malo konsenzusa treba nas učiniti opreznijima kada prosuđujemo novu ideju kao potpuno pogrešnu. Kako se već više nego jednom u povijesti događalo, pravi prodor mogao bi proizaći iz nečega što se protivi intuiciji” (Goertzel i Wang, 2007, str. 15).

Znanstvenici koji istražuju UI često rado govore o inteligentnim strojevima, ali među njima je veoma malo suglasnosti o osnovnom sastavu inteligencije. To ima za posljedicu da je među istraživačima veoma malo suglasnosti o tome što UI stvarno jest i što bi trebala biti. Složni su u tome da strojevima žele podariti attribute koje nisu u stanju definirati. Stoga umjetna inteligencija pati od nedostatka definicije svog područja djelovanja, a istraživanja umjetne inteligencije fokusirana su uglavnom na komponente inteligencije kao što su učenje, zaključivanje, rješavanje problema, percepciju te uporabu jezika (Schank, 1993, str 4). Roger C. Schank u članku objavljenom u knjizi *Osnove umjetene inteligencije* navodi neke značajke inteligencije. One ne moraju sveukupno biti prisutne, ali svaka za sebe integralni je dio inteligentnog entiteta. “Te su značajke: komunikacija, unutarnje znanje, znanje o svijetu, ciljevi i planovi te kreativnost. **Komunikacija** je moguća samo između inteligentnih subjekata, bez obzira na različite “vrste” inteligencije (komunikacija čovjeka i psa npr.) **Unutarnje znanje** odnosi se na očekivanje da inteligentni entitet ima neko znanje o sebi. Trebao bi znati kada mu je nešto potrebno, o nečemu misliti i znati da o tome misli. **Znanje o svijetu** pretpostavlja da inteligencija uključuje svjesnost o vanjskom svijetu i sposobnost pronalaženja i korištenja informacija koje netko ima o vanjskom svijetu. To također implicira posjedovanje memorije u kojoj su prošla iskustva kodirana i koja se mogu koristiti kao vodič za procesiranje novih iskustava. **Ciljevi i planovi** odnose se na ponašanje orijentirano ka cilju, u smislu poznavanja onoga što se želi postići (cilj) i poznavanje načina (plan) za postizanje toga cilja. Pravi kriterij kada se govori o planiranju je međuodnos planova i njihova pohrana na dovoljno apstraktan način kako bi se omogućilo da plan konstruiran za situaciju A bude usvojen i primijenjen na situaciju B. **Kreativnost** je pretpostavljena u određenim stupnjevima kod svakog inteligentnog entiteta. Definicija kreativnosti veoma je slaba (uključuje sposobnost pronalaženja novog smjera do izvora hrane, ukoliko je stari izvor

blokiran). Isto tako može značiti drugačiji pogled na nešto, što u značajnoj mjeri mijenja promatrani svijet. “Kreativnost prije svega znači sposobnost prilagodbe na promjene u okolini te sposobnost učenja iz iskustva. Možemo reći kako entitet koji ne uči vjerojatno nije inteligentan” (Schank, 1993, str 6).

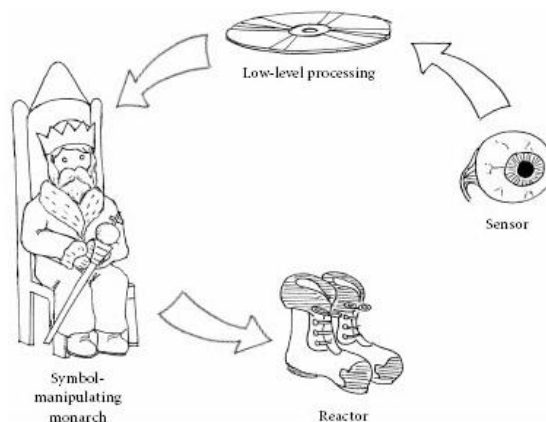
2.1. Smjerovi u istraživanju umjetne inteligencije

Tijekom razvoja umjetne inteligencije od 1950-ih do danas formirane su mnoge akademske škole, a svaku karakterizira vlastita metodologija istraživanja, akademski pogledi i fokusi istraživanja. Među njima se ističu kognitivistička škola, logička škola i bihevioristička škola. Ono što su istraživači 60-ih i 70-ih godina pokušavali napraviti je neka vrsta monolitnog entiteta koji posjeduje vlastite module za rasuđivanje određenih ulaznih podataka i razvoja hipoteza koje su zahtijevale sve više i više apstrakcije, pa su time postajale sve udaljenije od temeljnog ulaza. Takav je pristup doveo do spektakularnog neuspjeha i dijelom postao uzrokom **zime umjetne inteligencije**. Taj period, a riječ je 80-im godinama dvadesetog stoljeća, označava oskudno financiranje istraživanja i maleni napredak u umjetnoj inteligenciji. Nakon toga slijedi uspon **ekspertnih sustava** – odnosno tehnologija umjetne inteligencije koje su obilježene specifičnim ciljevima realiziranih sintezom velikih količina podataka iz određenog područja ljudske djelatnosti, kao što je računalni šahovski šampion Deep Blue ili MYCIN program za medicinsku dijagnostiku. Nakon ekspertnih sustava dolazi period u kojem mi sada djelujemo, a koje u jedan sustav povezuje pojmove poput strojnog učenja, vizualnog raspoznavanja znakova i s mnogo širim područjima primjene, ali izgrađenih od raspodijeljenih komponenti. Danas su znanstvenici oprezniji u predviđanjima, pa se o GUI-ju govori kao teoretskoj mogućnosti, dok se istovremeno mnogo radi na razvoju inteligentnih sustava koji upravljaju svakodnevnim ljudskim aktivnostima, poput pametnih agenata (robota ili softvera). Zato možemo govoriti o sve većoj prisutnosti tehnologija razvijenih u istraživanju umjetne inteligencije u današnjem društvu.

2.1.1. Kognitivistička škola

Kognitivističkoj školi pripadaju istraživači Herbert Simon, Marvin Minsky i Alan Newell. Njihova su istraživanja bila fokusirana na funkcionalnu simulaciju, sa računalima baziranim na ljudskoj noetici. Noetika ili teorija spoznaje je filozofska disciplina koja ispituje mogućnosti istinite spoznaje i raspravlja o spoznajnim izvorima, njihovom obujmu, pretpostavkama, granicama, kriteriju i objektivnoj vrijednosti spoznaje. Istraživači Newell i Simon zagovarali su 1950-ih “heuristički program” i zajedno napravili “Logic Theorist”. Taj računalni program je služio za simulaciju procesa ramišljanja prilikom dokazivanja matematičkih teorema. Šezdesetih su razvili program nazvan GPS (engl. General Program Solver) koji je simulirao općenita načela ljudskog načina rješavanja problema. Osamdesetih se Newell fokusirao na SOAR sustav, simboličku arhitekturu za generalno rješavanje problema, bazirano na mehanizmu učenja putem “komadanja” (koncept koji se smatrao bitnim u mnogim procesima percepcije, učenja i spoznaje kod ljudi i životinja), te memoriji baziranoj na pravilima za reprezentaciju operatora, kontrolu pretraživanja i sl.

Znanstveni cilj je bio izraditi senzorne module, primjerice vizualne sustave koji bi prevodili znanje iz svijeta u apstraktne simbole, nakon čega bi takve simbole predavali na obradu inteligentnoj jezgri, nekoj vrsti elektroničkog monarha. Prema ovom modelu, prikazanom na slici 2.1., monarh prima podatke od senzora, rasuđuje o njima, zatim emitira odluke koje ostvaruju reaktori.



Slika 2.1. Model monarha

Monarh u modelu je trebao manipulirati simbolima, zatim davati naredbe uređajima za pokretanje (najčešće kotačima) da se pokrenu. Takav je process idealizirao hijerarhiju sa mozgom na vrhu, očima i udovima na dnu (Shasha i Lazere, 2010 str. 19). Minsky je pak zauzeo fizikalni pogled, smatrajući da u svakodnevnim aktivnostima ljudi primjenjuju mnogo znanja stečenih i sakupljenih iz prijašnjih iskustava. Takva znanja, smatrao je, pohranjuju se u mozgu u strukturama poput okvira, pa je predložio strukturu reprezentacije znanja u okvirima. Minsky je vjerovao kako ne postoji unificirana teorija ljudske inteligencije. U svojoj knjizi *Društvo uma* je istaknuo kako je društvo uma opsežno društvo, individualno jednostavnih agenata sa određenim sposobnostima mišljenja (Shi, 2011, str 13).

2.1.2. Logička škola

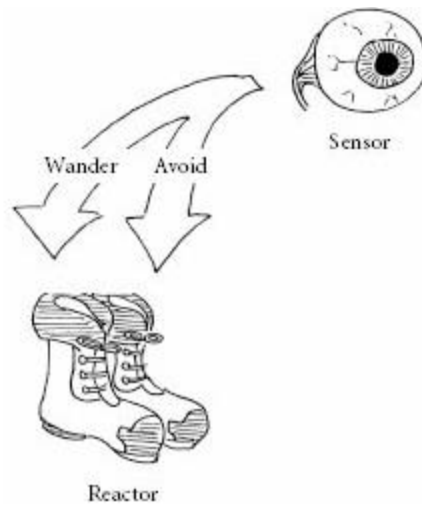
Logičku školu predstavljaju John McCarthy i Nils Nilsson. Logička perspektiva u istraživanju umjetne inteligencije opisuje objektivni svijet kroz formalizaciju. Ova se škola temelji na pretpostavkama da će inteligentni strojevi imati znanje o svojoj okolini i da će najsvestraniji inteligentni strojevi reprezentirati veći dio svoga znanja o okolini deklarativno. Deklarativno znanje je znanje o odnosima i načinima dekompozicije objekata, pojmova, događaja ili procesa. Ono služi za opis informacija neophodnih za rješavanje problema, odnosno opis objekata i činjenica, odnosa između objekata i/ili činjenica te opis drugih “pravila svijeta” upotrebljivih za neki sustav. Takvo se znanje naziva i statičko. Za razliku od toga, proceduralno znanje je ono o procedurama koje opisuju slijed akcija, načine promjene stanja radne okoline i znanje o postupcima za uporabu znanja te predstavlja znanje o tome kako se problemi rješavaju: opis metoda i postupaka za rješavanje problema, znanje o slijedu stanja i radnji, te znanje o znanju (meta-znanje). Takvo je znanje dinamičko. U konvencionalnim obradama podataka proceduralno znanje je u programima, a deklarativno u podacima smještenim u datotekama ili bazama podataka. Logička škola bila je fokusirana na konceptualnu reprezentaciju znanja, teorijsku semantiku i deduktivno zaključivanje, a McCarthy je tvrdio kako je sve moguće prikazati pomoću unificiranih logičkih okvira.

2.1.3. Bihevioristička škola

Većina istraživanja u umjetnoj inteligenciji imala je polazište u previše apstraktnim i jednostavnim modelima stvarnog svijeta. Rodney Brooks je smatrao kako se nužno treba ići iznad tih apstraktnih modela, te umjesto toga uzeti kompleksan stvarni svijet kao zaleđe istraživanja. Smatrao je kako inteligentni sustav mora imati svoje reprezentacije u fizičkom svijetu i inspiraciju potražiti u biološkim modelima.

Brooks je kreirao robote koji posjeduju skup neovisno dizajniranih vještina (robot koristi senzore i pogonski mehanizam za različite potrebne vještine, međusobno neovisne). Smatrao je da bi se teorije i tehnologije umjetne inteligencije mogle testirati na zadacima rješavanja problema u stvarnom svijetu i u tim bi se testovima one poboljšavale. Godine 1991. Brooks je iznio teoriju inteligencije bez reprezentacije i inteligenciju bez razumijevanja u kojoj je tvrdio da je inteligencija određena dinamikom interakcije sa svijetom (okolinom). Taj je rad nazvao jednostavno “robot” odnosno “roboti bazirani na ponašanju”, a novi smjer istraživanja umjetne inteligencije je dobio naziv ”**nouvelle AI**”. Ključni aspekti koji karakteriziraju ovaj pristup prikazan na slici 2.2. su:

- **Smještajnost:** Roboti su smješteni u svijetu i svijet direktno utječe na ponašanje robota.
- **Utjelovljenost:** Roboti imaju tijela i direktno iskušavaju svijet.
- **Inteligencija:** Izvor inteligencije nije ograničen samo na računalne strojeve. Dolazi isto tako iz situacija u svijetu.
- **Pojavnost:** Inteligencija sustava javlja se iz njegove interakcije sa svijetom i ponekad indirektnim interakcijama između komponenata sustava.



Slika 2.2. Nouvelle AI pristup

Brooksov je robot trebao koristiti senzore i pratiti jednostavna pravila kao što su: “*izbjegavaj sudar*”, “*lutaj*”. Prema navedenim idejama, Brooks je programirao autonomne mobilne robote temeljene na raspodijeljenim mrežama i staničnim automatima i tako dobio relativno neovisne jedinice za, naprimjer, funkcioniranje napredovanja, ravnoteže robota itd. Njegov je robot upješno hodao i 1990-ih pokrenuo sasvim novi pristup u istraživanjima robotike i umjetne inteligencije.

2.2. Umjetna inteligencija i filozofija

Zagovornici **jake umjetne inteligencije** vjeruju kako je računalo koje se ponaša na inteligentan način sposobno posjedovati mentalna stanja, stoga može biti zaista svjesno, na isti način kao ljudi, a da bi se to ostvarilo, potreban je samo dovoljno dobar softver. Smatraju da bi računala, ukoliko imaju dovoljno jake procesore za obradu podataka i dovoljno dobar softver, bila u stanju doslovno misliti i postići svjesnost koju posjeduju ljudska bića. “Mnogi filozofi i istraživači umjetne inteligencije smatraju ovaj stav krivim pa čak i komičnim. Mogućnost kreiranja emocionalnog i svjesnog robota često je predmet istraživanja znanstvene fantastike, ali se danas rijetko smatra ciljem umjetne inteligencije” (Coppin, 2004, str 32). Ipak, to je ostao ultimativni cilj istraživačima umjetne inteligencije, a futuristima i umjetnicima inspiracija za stvaranje klasičnih književnih i filmskih ostvarenja.

Slaba umjetna inteligencija je manje kontroverzna. Zasniva se na ideji da se računala mogu programirati na inteligentne načine u svrhu rješavanja specifičnih problema bez da te probleme razumiju. Pojednostavljeno, smatra se da se **inteligentno ponašanje** može modelirati i tada koristiti primjenjeno na računalima kako bi se rješavali kompleksni problemi.

Takav stav pretpostavlja da inteligentno ponašanje računala nije dokaz da je ono zaista inteligentno, na način koji ljudi jesu. Jedan argument u korist te tvrdnje dao je John Searle misaonim eksperimentom koji se naziva “**kineska soba**”. Sastoji se od sobe u kojoj se nalazi osoba koja ne razumije kineski jezik. Toj se osobi prosljeđuju poruke napisane na kineskom jeziku. Osoba u sobi posjeduje upute za prevođenje te nakon primitka poruke i prevođenja šalje prevedenu poruku izvan sobe. Druga osoba, (ispitivač) koja razumije kineski i nalazi se izvan sobe, prima prevedene poruke. Ispitivač bi na temelju primljenih (prevedenih) poruka mogao zaključiti da se u sobi nalazi osoba koja razumije kineski. Searle zaključuje kako bi osoba koja se nalazi u sobi prošla Turingov test iako ne razumije kineski jezik i time dokazuje kako su računala sposobna izvesti složene radnje nad nizom njima besmislenih znakova, ali će im uvijek nedostajati njihovo razumjevanje.

2.3. Mogu li strojevi misliti?

Alan Turing utemeljio je teoriju klasičnog računalstva 1936. godine. U radu iz 1950. nazvanim *Computing Machinery and Intelligence* postavio je pitanje: mogu li strojevi misliti? Ne samo da je branio tvrdnju kako je to moguće, već je predložio test kojim bi se utvrdilo je li program to i postigao. Test je danas poznat pod nazivom **Turingov test**. Zasniva se na tome da prikladan sudac (čovjek) nije u stanju prepoznati razgovara li sa drugim čovjekom ili sa računalnim programom, odnosno da li je program čovjek ili nije. Turing je u radu predložio i protokole za provođenje svog testa; poput toga da program i čovjek sa sućem interagiraju odvojeno i samo u tekstualnom obliku, preko nekog medija poput teleprintera. Tako bi se u testu provjeravala samo sposobnost razmišljanja kandidata ne i njihova pojavnost. Turingov test i argumenti koje je naveo nagnali su mnoge znanstvenike na razmišljanje, ne samo je li on u pravu, već i o tome kako proći test. Znanstvenici su započeli pisati programe s namjerom istraživanja načina za prolazak testa. Joseph Weizenbaum je 1964. napisao računalni program pod nazivom Eliza koji je bio osmišljen tako da oponaša psihoterapeuta. Procijenio je kako je psihoterapeut posebno jednostavan tip osobe za oponašanje iz razloga što može davati neodređene odgovore o sebi i može postavljati pitanja temeljena na pitanjima i odgovorima korisnika. Tipično se takvi programi temelje na dvije strategije. Najprije se skeniraju ulazne riječi koje unosi korisnik, zatim se među njima pretražuju određene ključne riječi i gramatički oblici. Ukoliko je taj postupak bio uspješan, sljedi odgovor; prema predlošku kojim se popunjavaju praznine, korištenjem opet riječi iz ulaza. Na primjer, ako se kao ulaz unese izraz “Ja mrzim svoj posao”, program može prepoznati gramatiku rečenice i postojanje posvojne zamjenice *svoj* i isto tako može prepoznati ključnu riječ *mrzim* iz ugrađene liste poput *volim/mrzim/želim* i u tom slučaju može odabrati odgovarajući predložak i odgovoriti: “Što najviše mrzite u svom poslu?” Ukoliko ne može analizirati unos do te mjere, onda može postaviti vlastito pitanje nasumičnim odabirom iz pohranjenih uzoraka, koji mogu i ne moraju biti ovisni o ulaznoj rečenici. Na primjer, ukoliko je postavljeno pitanje “*Kako radi frižider?*” program može odgovoriti: “Što je tako zanimljivo u tome ‘*kako radi frižider?*’” ili može samo pitati “Zašto vas to zanima?”

Druga strategija, koja se koristi u novijoj verziji Elize, je izgradnja baze podataka predhodnih razgovora. Tako se programu omogućuje jednostavno ponavljanje fraza koje su prethodno utipkali korisnici pa iste fraze program ponovo koristi ovisno o ključnim rječima koje je u datom trenutku unio korisnik. “Weizenbaum je bio šokiran činjenicom da su mnogi ljudi bili prevareni koristeći Elizu. Tako da je svojoj najnaivnijoj verziji prošla Turingov test. Štoviše, čak nakon što im je rečeno da Eliza nije autentična umjetna inteligencija, ponekad su korisnici nastavljali duge razgovore s programom o svojim osobnim problemima, upravo kao da su vjerovali kako ih razumije” (Deutsch, 2011, str. 349).

Danas je pak u svakodnevnoj uporabi ono što se naziva obrnuti Turingov test ili **CAPTCHA** (engl. Completely Automated Public Turing test to tell Computers and Humans Apart). Koristi se u svrhu detekcije softvera, na način da generiranjem nasumičnih slika slova ili brojeva sprečavaju zlonamjernim računalnim programima automatski pristup mrežnoj stranici. Od korisnika se zahtjeva da odgovori na test ponovnim unosom prikazanih znakova. Test, dakle, utvrđuje je li korisnik čovjek ili računalo, jer upravo ono što je za ljude lagano računalima još uvijek predstavlja problem.



Slika 2.3. Primjer CAPTCHA testa

U tvrtci Google su ustavonovili kako su današnje tehnike umjetne inteligencije takve da računala uspješno u prepoznaju 98 % iskrivljenog teksta (Shet, 2014.), a CAPTCHA test oduzima korisnicima previše vremena. Zato su 2014. godine objavili test prikazan na slici 2.4. koji su nazvali “No CAPTCHA reCAPTCHA”.



Slika 2.4. Primjer No CAPTCHA reCAPTCHA testa

Ipak, u situacijama koje zahtijevaju veću razinu sigurnosti oba se testa koriste zajedno, kao što je prikazano na slici 2.5.



Slika 2.5. Primjer istodobnog CAPTCHA i No CAPTCHA reCAPTCHA testa

Međutim, dosadašnji pokušaji da se izrade strojevi koji misle i koji mogu razgovarati s ljudima nisu bili uspješni na način kako je to Turing zamišljao. “Toliko je uspjeha dala potraga za *”strojevima koji razmišljaju”* u 65 godina, koliko je prošlo od Turingovog članka: nimalo. U svakom drugom pogledu računalna znanost i tehnologija doživjeli su nevjerojatan napredak” (Deutsch, 2011, str. 356).

Jedni tvrde da je spomenuta kritika nepravedna jer moderna istraživanja umjetne inteligencije nisu fokusirana na prolazak Turingovog testa, a veliki je napredak postignut u onome što se naziva umjetna inteligencija u mnogim specijaliziranim aplikacijama. Međutim, nijedna od aplikacija ne izgleda poput “strojeva koji razmišljaju”.

Drugi smatraju da je kritika preuranjena iz razloga što su većim dijelom povijesti istraživanja umjetne inteligencije računala imala apsurdno malu brzinu procesora i kapacitet memorije u odnosu na današnja. Tako oni očekuju prodor u sljedećih nekoliko godina. Deutsch smatra kako ni to nije temeljni razlog ne postojanju umjetne inteligencije u ovom trenutku. “Ni ovo neće proći. Ne radi se o tome da je netko napisao chatbot koji bi mogao proći Turingov test, ali je potrebno još godinu dana da se izračunaju svi odgovori. Ljudi bi rado čekali. U svakom slučaju, ne radi se o tome, kada bi netko znao kako napisati takav program, ne bi bilo potrebe za čekanjem. U svom je članku Turing je procijenio kako će za prolazak testa UI program zajedno sa svim podacima trebati 100 MB memorije, računala neće trebati biti brža od onih u njegovo doba (oko 10 000 operacija u sekundi) i da će do godine 2000. netko moći govoriti o strojnom razmišljanju bez da mu se suprotstavljaju. Ja sam, poput Turinga, uvjeren da bi ga se *moglo* programirati da misli; i to bi zaista moglo uključivati onoliko malo resursa koliko je Turing predvidio, iako su veći redovi veličine danas dostupni. Ali s kojim programom? I zašto nema

naznaka takvog programa?” (Deutsch, 2011, str. 357). Inteligencija u generalnom smislu, kako ju je Turing zamišljao, jedna je od postavki ljudskog uma koje su zagonetne filozofima već tisućljećima; ostale uključuju svjesnost, slobodnu volju i smisao. Turing je izumio svoj test u nadi da će zaobići sve te filozofske probleme. Drugim riječima, nadao se kako je funkcionalnost moguće postići prije nego li je objašnjena. Nažalost, jako rijetko je praktična rješenja temeljnih problema moguće ostvariti bez ikakvog objašnjenja o tome zašto i kako djeluju. “U ovom trenutku za trenutačno stanje discipline pravilo jest: ako ga nije moguće programirati, onda nema veze s inteligencijom u Turingovom smislu” (Deutsch, 2011, str. 362-63). S time u vezi, Deutsch je predložio jednostavan test za prosudbu bilo kojih tvdrnji koje bi objasnile prirodu svijesti ili bilo kojeg računalnog zadatka: “Ako to ne možeš programirati, onda to nisi razumio”.

3. KLASIČNE METODE U UMJETNOJ INTELIGENCIJI

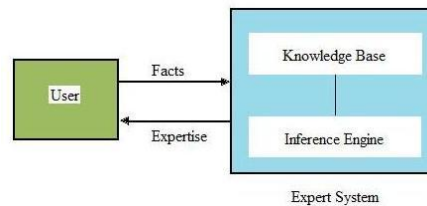
Klasična umjetna inteligencija zasniva se na korištenju **simboličkog pristupa** i metoda. Polazi od pretpostavke da se inteligentno ponašanje postiže razvojem sustava kroz logičko zaključivanje na temelju definiranog skupa pravila i činjenica – **simbola**. Primjenjuje se takozvani **TOP-DOWN pristup**, a temeljena je na znanju sa simboličkim opisom svijeta kroz definirani skup pravila. Pretpostavka je da se inteligentno ponašanje postiže spoznajnom analizom (koja je neovisna o biološkoj strukturi mozga) putem obrade simbola. Top-Down ovdje znači da se problem podijeli na podprobleme u rekurzivnom smislu, odnosno jedan komplicirani problem rješava se tako da se podijeli na više manjih, jednostavnijih problema.

Metode korištene u klasičnoj umjetnoj inteligenciji poput modeliranja odlučivanja kod ekspertnih sustava ili stabla odlučivanja postizale su dobre rezultate, ali su isto tako pokazivale ozbiljne nedostatke kod primjene na stvarnim robotima i u nesigurnim uvjetima. Posljedica toga je kombinatorna eksplozija, izazvana brojnim pravilima koje se pojavljuju u kompleksnom stvarnom svijetu kao i činjenica da je nemoguće predvidjeti sve situacije koje bi autonomni entitet mogao susresti. Stoga su kasnije razvijene novije metode, potaknute nastojanjem da se sustav prilagodi situacijama iz stvarnog svijeta.

Klasični pristup umjetnoj inteligenciji je dakle takav da se pogleda rad mozga izvana i kao rezultat promatranja pokuša replicirati njegova izvedba. Međutim, klasične metode nisu toliko dobre kada se nastoji “osvijestiti” situacija i pokuša napraviti gruba usporedba sa prije naučenim iskustvima – nešto što je izuzetno značajan aspekt svakodnevne inteligencije. Klasičan je pristup jako uspješan u radu s dobro definiranim zadacima na koje se primjenjuju skupovi jasnih pravila. To je posebno prikladno za slučajeve kada je mnoštvo takvih pravila potrebno obraditi, a zatim odgovore primjeniti u kratkom vremenu. Prednost računala u brzini memorijskog odgovora u tome igra važnu ulogu.

3.1. Ekspertni sustavi (sustavi temeljeni na znanju)

Koncept ekspertnog sustava zasniva se na tome da stroj (računalo) može prosuđivati o činjenicama iz određene domene i da može, u grubom smislu, raditi na isti način kako bi radio mozak stručnjaka iz tog područja. Da bi to postiglo, računalu je potrebno znanje iz te domene (**baza znanja**), neka **pravila** koja pripremaju stručnjaci i koja se moraju slijediti kada se pojavi nova informacija te neki **način komunikacije** koji se treba ostvariti sa korisnikom sustava. Takvi se sustavi nazivaju sustavi temeljeni na pravilima, sustavi temeljeni na znanju ili općenito ekspertni sustavi. Često se koriste, jer se većina ljudskog znanja može izraziti produkcijskim pravilima (ako-onda). Temeljne sastavnice ekspertnih sustava, prikazane na slici 3.1., jesu baza znanja i mehanizam zaključivanja.



Slika 3.1. Osnovne komponente ekspertnih sustava

Jedan od prvih uspješnih ekspertnih sustava bio je medicinski sustav nazvan MYCIN, a koristio se u svhu dijagnostike krvnih infekcija. Sadržavao je od prilike 450 pravila i tvrdilo se da je bolji u postavljanju dijagnoze od većine “mlađih” liječnika, a dobar poput nekih eksperata. Pravila u sustavu nisu bila teoretski generirana, već su bila sastavljena nakon razgovora provedenih s mnogim stručnjacima koji su davali izvještaje, iz vlastitog neposrednog iskustva. Na taj način pravila su, barem djelomično, odražavala nesigurnosti uočene u medicinskim stanjima. Općenito, struktura MYCIN-a bila je kao i kod svih ekspertnih sustava takva da produkcijska pravila imaju temeljni oblik:

AKO (uvjet) ONDA (zaključak)

Na primjer, pravilo bi moglo glasniti: AKO (kihanje) ONDA (gripa).

Međutim, moguće je da nekoliko uvjeta treba biti istovremeno ispunjeno kako bi se ostvarilo pravilo (zaključak), odnosno da pravilo iz uvjeta bude istinito. Isto tako, može samo jedan uvjet iz niza biti ispunjen da bi se izveo zaključak. Tada pravilo može imati oblik:

AKO (uvjet1 **i** uvjet2 **ili** uvjet3) **ONDA** (zaključak)

Primjenjeno na dijagnostiku gripe, pravilo bi onda moglo glasiti:

AKO (kihanje **i** kašljanje **ili** glavobolja) **ONDA** (gripa)

Postoje situacije kada se iz istog skupa činjenica može izvesti nekoliko mogućih zaključaka. To predstavlja problem za stručnjaka i za ekspertni sustav. Da bi takve situacije bile rješive uvode se samo za takve slučajeve dodatna pravila kako bi se odlučilo koja akcija treba biti primijenjena i to se naziva **razrješavanje konflikta**.

Kada nekoliko pravila ima ispunjene uvjete, odabire se jedno ovisno o primijenjenim kriterijima, koji mogu biti:

1. Pravilo najvišeg prioriteta - svakom se pravilu dodjeli prioritet pa se odabire onaj s najvećim, kod razrješavanja konflikta.
2. Uvjeti najvišeg prioriteta – svakom uvjetu se dodjeli prioritet. Da bi pravilo bilo odabrano, moraju sadržavati uvjete s najvišim prioritetima.
3. Najnoviji uvjeti – pravilo čiji su uvjeti ostvareni najrecentnije.
4. Najspecifičniji – odabire se pravilo čijih je uvjeta najviše ispunjeno (još se naziva i najduže poklapanje).
5. Limitirani kontekstom – pravila se podijele u grupe od kojih su samo neke aktivne u određenom trenutku. Da bi bilo odabrano, pravilo mora pripadati aktivnoj grupi. Tako se ekspertni sustav može prilagoditi različitim situacijama tijekom vremena.

Koja će se metoda za razrješavanje konflikta primijeniti ovisi isključivo o aplikaciji pa je vjerovatno da će se za jednostavne sustave primijeniti i jednostavne metode (Warwick, 2012, str.34).

Primjenjuju se kod ekspertnih sustava i dva načina napredovanja prema zaključcima:

a) ulančavanje pravila prema naprijed

Ono započinje poznatim podacima i odlikuje ga napredovanje prema zaključku (engl. forward chaining (forchaining), data driven processing, event driven, bottom-up, antecedent, pattern directed processing reasoning). Kod normalnog rada, ekspertnih sustava, skup činjenica postat će očit u određenom vremenu i one će ispaliti (pokrenuti) skup pravila, odpuštajući dalje ostala pravila i tako dalje sve dok se ne dostigne konačni zaključak. Način rada ku kojemu proces zaključivanja polazi od ulaznih podataka do konačnog cilja naziva se ulančavanje pravila prema naprijed. "Cilj takvog načina rada je otkrivanje svega što se može deducirati iz danog skupa činjenica" (Warwick, 2012, str. 35).

b) ulančavanje pravila unatrag

Postupak se primjenjuje kroz odabir mogućeg zaključka (neke hipoteze), nakon čega slijedi pokušaj dokazivanja valjanosti hipoteze traženjem valjanih dokaza. Ekspertni sustavi mogu se koristiti, dakle, i obratno. Kada je postignut cilj, provjerava se koje su se činjenice (podaci) pojavili da bi sustav donio zaključak koji jest. Isto tako, može se promatrati unazad kroz sustav da bi se procijenilo koje činjenice moramo unijeti u sustav da bi se specifični cilj ostvario.

3.1.1. Prednosti ekspertnih sustava

Ekspertni sustavi imaju brojne prednosti pred ostalim UI metodama. Prva je da su relativno laki za programiranje u računalima (uniformne linije koda u IF-THEN strukturi). Osim toga, svako je pravilo zasebni entitet s vlastitim skupom pravila za paljenje i vlastitim zaključcima. Ako se novo pravilo pokaže nužnim, može se dodati ukupnom sustavu. Ekspertni sustav je idealan za informacije iz prirodnih situacija u stvarnom svijetu jer ipak su to iste informacije sa kojima se suočavaju stručnjaci kada kažu: "u ovoj situaciji, postupam tako", a to se onda lako unosi u ekspertni sustav (Warwick, 2012, str.36). Struktura sustava je odvojena od podataka u sustavu, tako da se problemsko područje, u smislu strukture može primijeniti na različite domene. Ono što je različito je skup pravila i način na koji se pravila kombiniraju. Tako se ista struktura ekspertnog sustava može koristiti u medicinskoj dijagnostici i upravljanju strojevima, pod

uvjetom da se primjene drugačija pravila i drugi podaci. Veliku prednost ekspertni sustavi imaju u brzini odgovora, što konačno dovodi do značajnih ušteda u novcu i većoj sigurnosti u primjeni. Ekspertni sustavi koji se koriste u radu sa strojevima ili pokreću alarme za nadzor ili pak u financijskim transakcijama su izvrstan primjer.

3.1.2. Nedostaci ekspertnih sustava

Nekoliko je problema koji se susreću u izradi ekspertnih sustava. Jedan je način sakupljanja pravila, koji može biti nezgodan u smislu da je stručnjacima iz nekog područja ponekad teško jednostavnim terminima objasniti što rade u svakodnevnim situacijama. Na to se nadovezuje činjenica da različiti stručnjaci, kada ih se pita, različito razmišljaju o problemu pa je tada teško standardizirati pravila. Isto tako, treba napomenuti da stručnjaci znaju biti skupi. Najveći problem ekspertnih sustava je ono što se naziva kombinatorna eksplozija – sustav postaje prevelik. Ona je posljedica težnji ovih sustava da izvode zaključke kako bi dali odgovor na svaku moguću situaciju.

Bitno je uočiti da ekspertni sustavi nisu samo programirani prema mehanizmu za donošenje odluka i ne ponašaju se uvijek prema očekivanjima, odnosno može ih programirati da uče dok donose zaključke (Warwick, 2012, str. 38).

3.2. Strojno učenje

Strojno učenje je grana umjetne inteligencije. Značajni aspekt umjetne inteligencije je upravo sposobnost računala da uče. Općenito, za računalni program kažemo da uči iz iskustva E s obzirom na neku klasu zadaća T i mjeru izvedbe P, ako je njegova izvedba zadaće T, mjerena sa P poboljšana sa iskustvom E. Na temelju toga izvodi se skup objekata koji onda definiraju strojno učenje:

- zadaća (T), jedna ili više njih
- iskustvo (E)
- izvedba (P)

Dok računalo izvodi neki skup zadaća, iskustvo bi trebalo voditi ka povećanju izvedbe. Strojno učenje se uglavnom može svrstati u dvije kategorije:

- nadzirano – učenje s učiteljem
- nenadzirano – učenje bez učitelja

Učenje s učiteljem pretpostavlja rad sa skupom prethodno označenih podataka na kojima se uči. Za svaki primjer podatka kojim se radi potrebno je odrediti parove ulaznih i izlaznih objekata. Na primjer, klasifikacija podataka preuzetih s Twittera:

```
Baš volim zadnji Waitsov album!  
#moda prodajem tenisice Adidas! Tko je zainteresiran? #adidas  
Moj Hadoop cluster radi na velikoj količini podataka. #data
```

Kako bi klasifikator mogao znati izlazni rezultat svakog tweeta, potrebno je ručno unijeti odgovore. Izlazni objekti u ovom su primjeru: **muzika**, **obuća** i **bigdata**, prikazani na kraju reda.

Baš volim zadnji Waitsov album!	muzika
#moda prodajem tenisice Adidas! Tko je zainteresiran? #adidas	obuća
Moj Hadoop cluster radi na velikoj količini podataka. #data	bigdata

Da bi klasifikator, kada se uspješno pokrene, mogao pronaći smisao u podacima, potrebno je ručno upisati puno ulaznih podataka. Ono što imamo je skup za učenje koji se koristi za kasniju klasifikaciju podataka. Kada je riječ o **učenju bez učitelja**, računalu se prepušta nalaženje skrivenih uzoraka u gomili podataka. U tom slučaju nema unaprijed određenog odgovora; pokrene se algoritam koji izvodi strojno učenje i onda se pogleda ishod, zato je učenje bez učitelja možda više rudarenje podataka, nego učenje (Bell, 2014, str 9).

Mnogo je praktičnih primjera u kojima se primjenjuje strojno učenje, poput **softvera** koji nakon prvog korištenja, u radu s korisnicima, uči o njihovim navikama i s vremenom može predvidjeti želje korisnika. Koristi se i u **detekciji spamova**, kao filter za klasifikaciju neželjene e-pošte. Kada algoritam uoči poruku koju smatra spamom, upita korisnika za potvrdu. Ako korisnik potvrdi da je poruka spam, na temelju toga i na temelju iskustva, algoritam buduće slične poruke tretira na isti način.

Tehnike strojnog učenja mogu se primjeniti i na ekspertne sustave. U primjeni na ekspertne sustave, računalima je potreban ljudski input (učenje s učiteljem). Ekspertni sustavi bazirani na pravilima, po definiciji započinju ekstrakcijom niza pravila koju zadaju ljudski stručnjaci, zajedno sa ostalim informacijama iz domene problema koji se obrađuje. Rezultat je nastala baza pravila, koja mogu rezultirati novim pravilima, kada se aktiviraju. Tako ovisno o određenim ulaznim podacima i prema najboljoj seriji pravila (može ih biti nekoliko) pokrenu se ostala pravila u nizu dok se ne dostigne konačni zaključak ili konačno pravilo. Da bi se to ostvarilo, potrebno je prethodno aktivirati sva pravila u seriji. Pri tome se dobra rješenja nagrađuju, a loša kažnjavaju. Tako bi se sljedeći put kod postignutih istih uvjeta još vjerojatnije pokrenuo isti niz pravila. Loša rješenja se kažnjavaju tako da vjerojatnost sljedećeg izvođenja postane manja ili otklanjaju ručnim ispravkama zadanih postavki.

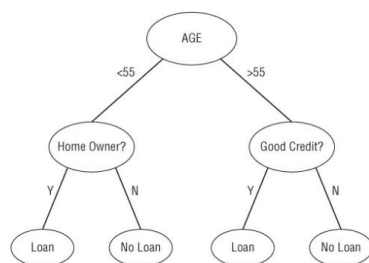
3.3. Rudarenje podataka

Danas opseg dostupnih informacija, za donošenje odluka, prelazi ljudske mogućnosti – imamo informacijsko preopterećenje. Stoga se brojne tvrtke bave upravo analizom i predstavljanjem najboljih opcija, pa nas onda savjetuju u donošenju najbolje odluke, uz određenu cijenu za svoje usluge. Bez obzira da li se oslanjamo na ljude ili strojeve, u izvlačenju znanja (ili uzoraka) iz kompleksnih dostupnih podataka govorimo o rudarenju podataka. Sustavi umjetne inteligencije posebno su pogodni za takve zadaće, zato što imaju sposobnost pohrane velikih količina podataka iz kojih mogu izvući svakakve vrste odnosa kako bi razvile uzorke, veze i smislene poveznice. Upravo tu sposobnost koriste današnji algoritmi pri kognitivnom računalstvu i onome što se naziva duboko učenje.

U mnogim situacijama nalazimo mnogo različitih dijelova podataka, a mi trebamo otkriti sličnosti, poveznice i odnose među dijelovima ili pak želimo odrediti najbitnije dijelove. Možda nas zanima predvidjeti ishod temeljen na podacima koji su nam trenutačno poznati, stoga trebamo razlučiti koji su dijelovi podataka relevantni, a koji nisu. Kupovina u supermarketu je dobar primjer i to zato što mnogi ljudi redovito obavljaju baš takvu kupovinu. Postoji oko 100 različitih tipova proizvoda (uglavnom hrane) u takvim trgovinama i svaki put kad osoba koristi supermarket, dobivaju se podaci o tome što je osoba kupila. U nekom periodu, napravi se statistička veza o tome što je osoba kupila i koliko često je taj proizvod kupila. Slično se napravi i za ostale proizvode i ostale osobe. Očiti je cilj da se na kraju može reći kako će “sljedeće srijede ta osoba doći i kupiti taj proizvod, a ako bude dostupan možda kupi i ovaj drugi, na temelju naših predviđanja”. To je način kojim se može zaraditi na rudarenju podataka. Temeljna statistička metoda koja se ovdje koristi je **korelacija** – iz koje je vidljivo kako se jedan dio podataka odnosi prema drugom. Hoće li se povećanjem jednoga i drugi povećati ili će se smanjiti (Warwick, 2012. str. 56).

Da bi se reducirala kompleksnost problema, može se koristiti i **tehnika stabla odlučivanja**. U kojoj se radi o tome da je čitava baza podataka raskomadana na manje porcije kojima se lakše

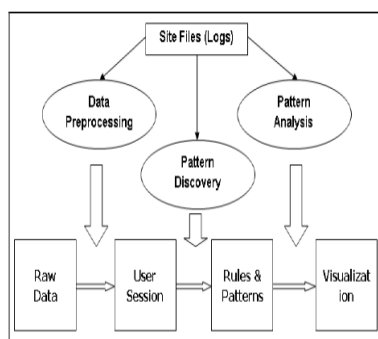
upravlja i to na način koji ovisi o zahtjevima korisnika. Na slici 3.2. prikazana je metoda stabla odlučivanja koja se može koristiti prilikom odobrenja kreditiranja klijenata ovisno o njihovoj dobi, prijašnjim zaduživanjima ili posjedovanju nekretnine kao zalogu otplate.



3.2. Primjer stabla odlučivanja

Na primjeru kupovine u supermarketu mogli bi odlučiti da želimo proučiti samo ženske kupce i tako stvaramo korisnički - specifičnu granu kojoj su relevantni izlazni podatci vezani samo uz ženske kupce. Tada bi UI sustav kompletno ignorirao podatke vezane uz muške kupce. Ukoliko želimo unijeti još neke kriterije (poput toga da se prate kupci koji troše samo više od 50 kn po posjeti ili redovno kupuju svježije povrće) tada se vrijeme analize dramatično smanjuje, a točnost predviđanja poboljšava.

Današnje primjene metoda poput tehnike rudarenja podataka veoma su korisne u marketingu jer omogućuju nabavu proizvoda i analizu uzoraka pretraživanja i ponašanja pa se tada može ciljano izraditi ponuda za bilo koju grupu ljudi. Na primjer, mnogo se koristi i rudarenje korištenja podataka s mreže. Pri tome se podacima prikupljenim rudarenjem s interneta nastoji otkriti skriveno znanje o korisnicima i njihovom ponašanju na mreži. Dobivene informacije donose prednost i povećanje zarade davateljima usluga koji onda mogu modificirati sadržaj stranica, poboljšati rad sustava ili personalizirati prikazani sadržaj. Proces rudarenja podataka o korištenju interneta prikazan je na slici 3.3.



Slika3.3. Proces rudarenja podataka korištenja Interneta

Proces rudarenja podataka sastoji se od tri osnovna koraka. U prvom se obrađuju sakupljeni podatci, što zahtjeva najviše vremena (tipično 60 do 90 % ukupnog vremena za izvođenje projekta), zatim se među podacima otkrivaju uzorci (pokretanjem algoritma za rudarenje podataka; asocijacijskih pravila, rudarenje sekvencijskih uzoraka, grupiranje uzoraka) i konačno obavljaju se analiza i prikaz dobivenih rezultata. U posljednjem koraku dobiveni se rezultati analiziraju kako bi se nevažni i trivijalni rezultati izdvojili od korisnih. Rudarenje podataka izvodi se danas u analizi poslovnih kretanja i financijskih transakcija na burzama. Moguće je na taj način predviđati trendove pa se potencijalni ishod može procijeniti ako su ostvareni određeni uvjeti.

“ Novo područje primjene rudarenja podataka je otkrivanje kriminalnih aktivnosti. Promatra se tipično ponašanje grupe ljudi ili pojedinca (može biti i stalno praćeno) pa se bilo koja devijacija brzo uočava. Na taj način mogu se identificirati zločini poput krađe kreditne kartice“(Warwick, 2012, str. 58).

4. MODERNE METODE U UMJETNOJ INTELIGENCIJI

Posljednjih je godina moderan pristup umjetnoj inteligenciji više usredotočen na bottom-up tehnike, odnosno na to da se uzmu osnovni gradivni blokovi inteligencije, koji se zatim postave zajedno u određene situacije te ih se ostavi da uče i da se razvijaju određeni vremenski period pa se pogledaju dobiveni rezultati. **Situacijskim pristupom** nastoji se ostvariti umjetna inteligencija koja je utjelovljena i smještena u stvarnom svijetu. Takav je pristup u istraživanju umjetne inteligencije nastao zadnjih dvadestak godina, a temelji se na izgradnji inteligentnih agenata koji se u svome okruženju ponašaju uspješno. Temelj ove metode je **BOTTOM-UP dizajn**. Bottom-up odnosi se na elementarna ponašanja koja se kombiniraju kako bi se ostvarila kompleksnija ponašanja.

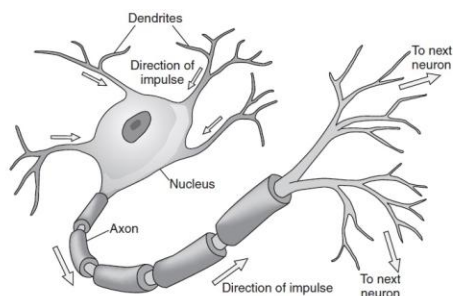
Novi pristup zagovara ideju da se inteligencija kod strojeva može ostvariti, ali kroz dovoljne motorne vještine i senzornu interakciju sa okolinom odnosno **smještenost**. Izraz smješten (engl. situated) nastao je u robotici i odnosio se na smještaj robota u okolini, međutim može se smatrati da se smještaj odnosi i na softverske agente pod uvjetom da se oni nalaze u dinamičkom okruženju (koje se brzo mijenja), da svojim ponašanjem mogu manipulirati i mijenjati okruženje koje su u stanju osjećati i percipirati. Naglasak je, dakle, na ponašanju i ne oslanja se na simbolički opis svijeta, već na kreiranju modela interakcija entiteta i njihove okoline. Smještajni pristup ulaže mnogo manje prioriteta apstraktnom zaključivanju ili vještinama koje zahtjevaju rješavanje problema.

Inteligentno ponašanje entiteta postiže se u njegovoj interakciji s okolinom i to kroz povezivanje jednostavnih procesnih elemenata koji rade paralelno (poput neurona u mozgu). To je temeljna ideja prema kojoj funkcioniraju i umjetne neuronske mreže, pa se prema njihovoj implementaciji ovaj pristup naziva još i **konektivistički**.

Prvi temeljni koncept moderne umjetne inteligencije jest razmotriti način rada biološkog mozga u smislu; osnovnih funkcija, razvoja i prilagodbe tijekom vremena. Drugi se temelji na potrebi za dobivanjem relativno jednostavnih modela temeljnih elemenata - gradivnih elemenata - mozga. Treće, te gradivne elemente potrebno je oponašati tehnološkim dizajnom – možda elektroničkim krugom, možda računalnim programom s ciljem da simuliraju gradivne blokove mozga. Umjetni gradivni blokovi tada se mogu zajedno priključiti na različite načine kako bi djelovali slično mozgu. “Može se činiti kako je cilj kopirati originalni mozak na neki način, međutim samo je u pitanju traženje inspiracije iz biološkog načina djelovanja kako bi ga se koristilo u tehnološkom dizajnu. Tada se umjetna verzija koristi prednostima biološkog mozga, poput sposobnosti generalizacije ili jednostavno kategoriziranje događaja u jednu ili drugu kategoriju” (Warwick, 2012, str. 88).

4.1. Umjetne neuronske mreže

Točan način na koji mozak omogućuje misao jedna je od najvećih misterija znanosti. Tisućama godina je poznato kako snažan udarac u glavu može dovesti do nesvjestice, privremenog gubitka pamćenja ili trajnog gubitka mentalnih sposobnosti. To je dalo naslutiti kako je mozak na neki način povezan sa razmišljanjem (Norvig i Russell, 1995, str. 564). Poznato je da je **neuron** ili živčana stanica (prikazana na slici 4.1.) osnovni funkcionalni dio tkiva živčanog sustava, uključujući mozak. Svaki neuron sastoji se od tijela stanice, koji se još naziva **soma** i u kojem je smještena stanična jezgra. Iz tijela stanice dalje se granaju brojna vlakna koja se nazivaju **dendriti** i jedno najduže vlakno (najduži dendrit) koji se naziva **akson**. Dendriti se granaju u grmolike mrežaste strukture oko stanice dok se akson produljuje na veću udaljenost - obično do jednog centimetra (to je 100 puta veće od promjera tijela stanice), a ponekad u ekstremnim slučajevima i do jednog metra. Naposljetku, akson se također grana u strukture i podstrukture koje su povezane sa dendritima i staničnim tijelima drugih neurona. Povezane spojnice neurona nazivaju se **sinapse**.



Slika 4.1. Građa neurona

Svaki neuron formira sinapse s ostalim neuronima, a može ih biti nekoliko desetaka do nekoliko stotina tisuća. Obično se neuron nalazi u stanju mirovanja, a signal koji prima je u obliku elektrokemijskog pulsa koji do njega dolazi putem dendrita od ostalih neurona. Svaki puls mijenja električni potencijal tijela stanice – neki od dendrita doprinose signalnom potencijalu (nazivaju se uzbuđujući), dok ga neki umanjuju (nazivaju se inhibitorni). Ukoliko ukupni signal na dendritima u bilo kojem trenutku dostigne vrijednost praga, tada će stanica ispaliti elektrokemijski puls koji se još naziva **akcijski potencijal** u svoj akson odnosno u ostale neurone kako bi se i oni nakon toga pobudili. Kratko nakon toga neuron se vraća u stanje mirovanja i čeka da se ponovo izgradi puls na njegovim dendritima. Obratno, ako nije postignuta vrijednost praga, tada neuron neće okinuti. To je **sve ili ništa** process u kojem neuron ili okida ili ne (Warwick, 2012, str. 90).

Taakva struktura nastaje djelomično iz genetskih razloga, ovisno o građi mozga predaka, a djelomično razvojem individue kroz životna iskustva. Kako osoba uči, veze akson – dendrit jačaju u njezinom mozgu (pozitivno) ili slabe (negativno) i tako čine da se osoba manje ili više ponaša na određeni način. Tako je mozak izuzetno plastičan na način da se prilagođava i funkcionira drugačije ovisno o uzorcima signala koje prima i nagradama ili kaznama koje su s time povezani. Kada činimo nešto ispravno, kao odgovor na određeni događaj, znači da će neuronski putevi uključeni u odlučivanje jačati, tako da idući put kada se pojavi isti događaj mozak vjerovatnije izvodi sličan odabir. Istodobno, ako se kao odgovor na određeni događaj nešto čini neispravnim, tada neuronski putevi uključeni u taj odabir slabe. Na taj način će mozak činiti manje pogrešaka (Warwick, 2012, str.91). To je osnova na kojoj se temelje biološki rast i razvoj mozga i koja omogućuje mozgu izvođenje operacija.

Ideje preuzete iz strukture biološke neuonske mreže i njezine metode učenja temeljne su sastavnice koje se primjenjuju u **umjetnim neuronskim mrežama (UNM)**, kod kojih je cilj upotrijebiti tehnološka sredstva za ostvarenje nekih svojstava originalne biološke verzije.

Opća definicija mogla bi glasiti: **Neuronska mreža** jest skup međusobno povezanih jednostavnih procesnih elemenata, jedinica ili čvorova, čija se funkcionalnost temelji na biološkom neuronu.

“Neurofiziološka istraživanja, koja su nam omogućila bolje razumijevanje strukture mozga, daju naslutiti da je modelu mozga najsličniji model u kojem brojni procesni elementi podatke obrađuju paralelno. Područje računarstva koje se bavi tim aspektom obrade informacija zovemo neuro-računarstvo, a paradigmu obrade podatka umjetnom neuronskom mrežom (engl. Artificial Neural Network)” (Bašić, 2008, str. 4).

Kada uspoređujemo mozgove sa digitalnim računalima, prikazano na slici 4.2., vidimo kako oni odrađuju veoma različite zadatke i imaju veoma različita svojstva.

atribut	mozak	računalo
tip elementa za procesiranje	neuron (100 različitih vrsta)	bistabil
brzina prijenosa	2 ms ciklus	ns ciklus
broj procesora	oko 10^{11}	10 ili manje
broj veza među procesorima	10^3 - 10^4	10 ili manje
način rada	serijski, paralelno	serijski
signali	analogni	digitalni
informacije	ispravne i neispravne	ispravne
pogreške	nefatalne	fatalne
redundancija	stotine novih stanica	eventualno rezervni sustav

Slika 4.2. Razlike između digitalnih računala i mozga

Posebna istraživanja rade se na području arhitekture računala koja bi na pogodniji način od konvencionalne von Neumannove arhitekture omogućila učinkovitu primjenu umjetne neuronske mreže. “Konvencionalna arhitektura računala temelji se na sekvencijalnoj obradi podataka, koja nema mnogo zajedničkog sa strukturom i načinom funkcioniranja mozga” (Bašić, 2008, str. 4). Čipovi u računalima mogu izvoditi naredbe u desecima nanosekundi, dok neuroni za paljenje trebaju mikrosekunde. Mozgovi to nadoknađuju više nego dobro, time što su svi neuroni i sinapse simultano aktivni, dok današnja računala imaju jednu ili svega nekoliko CPU-a. Neuronska mreža koja radi na serijskom računalu zahtjeva stotine ciklusa kako bi odlučila hoće

li neuronu slična jedinica ispaliti signal ili neće dok u pravom mozgu svi neuroni to čine u jednom koraku. Tako, iako su računala milijun puta brža u sirovoj brzini okidanja, mozak je na kraju milijun puta brži u onome što radi.

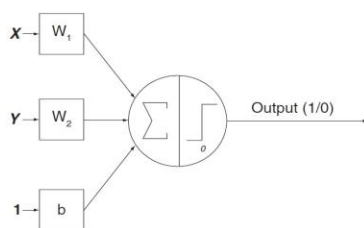
Jedna je od privlačnosti metode umjetnih neuronskih mreža u nadi da se mogu konstruirati uređaji koji će kombinirati usporednost mozga sa brzinom računala. Mozak može obavljati kompleksne zadatke – poput prepoznavanja lica – u vremenu manjem od sekunde, što bi značilo da je moguće u tom periodu izvesti samo sto ciklusa. Serijsko računalo zahtijeva milijarde ciklusa da obavi isti zadatak manje uspješno. Mozgovi su tolerantniji na pogreške od računala. Greška na hardveru koja izmjeni jedan bit može upropastiti čitavo računanje, ali moždane stanice umiru svakog trena bez ikakvog učinka na cjelokupnu moždanu funkciju. Dodatno, mozgovi su konstantno suočeni s novim informacijama pa ipak uspijevaju s njima nešto učiniti. Još jedna povoljna osobina neuronskih mreža je u njihovoj sklonosti prema postepnom padu izvedbe kada se uvjeti pogoršaju.

4.1.1. Primjer modela umjetnog neurona

McCulloch-Pitts model umjetnog neurona prikazan na slici 4.3. naziva se još i Threshold Logic Unit (TLU). Model oponaša funkcionalnost biološkog neurona, pri čemu koristi zadanu vrijednost praga (engl. threshold). Mada je TLU model veoma sličan stvarnom neuronu valja istaknuti kako je to samo jedna od mogućnosti koja se koristi u modeliranju umjetnih neuronskih mreža. U primjeni modela koristi se slijedeća analogija: signali su opisani numeričkim iznosom, a na ulazu u neuron množe se težinskim faktorom koji predstavlja jakost sinapse; signali pomnoženi težinskim faktorima zatim se sumiraju analogno sumiranju potencijala u tijelu stanice; ako je dobiveni iznos iznad definiranog praga, neuron daje izlazni signal, odnosno neuron “pali”.

Ulazi x i y pomnože se s težinskim faktorima (težinama) w_1 i w_2 koji su im dodijeljeni, a zatim se zbroje. Ukupna vrijednost se onda upoređuje sa zadanom vrijednosti praga b . Prag je efektivno negativna vrijednost koju suma težine i ulaza mora prijeći. Ukoliko je suma težine i ulaza jednaka ili veća od praga b , neuron “pali”, i daje **izlaz 1**. Ukoliko je ta suma manja od b ,

neuron ne pali i daje **izlaz 0**. Tako dobiveni izlaz iz neurona može se pomnožiti dalje sa svojim težinskim faktorom i postati ulaz za sljedeći neuron.



Slika 4.3. McCulloch-Pitts model umjetnog neurona

Kao praktičan primjer uporabe umjetnih neuronskih mreža može se uzeti jednostavan test provjere prikladnosti kandidata za dobivanje pozajmice. Ulaz x je 0 ako kandidati nikada prije nisu platili pozajmicu, a 1 ukoliko jesu; ulaz y je 0 ako imaju uštedevinu ispod nekog minimuma. Pretpostavimo da ako kandidati zadovoljavaju kriterij da su x i y 1, tada će dobiti pozajmicu, inače ne. Jednostavnije neuronske mreže moguće je dakle konstruirati tako da obavljaju određeni zadatak. To je moguće za mreže koje se sastoje od TLU perceptrona i koje obavljaju unaprijed zadanu logičku funkciju, jer u tom slučaju možemo pratiti što i kako točno mreža radi. U slučaju kada se koriste složenije prijenosne funkcije ili se radi s realnim brojevima, tipično se gubi nadzor nad načinom kako mreža obrađuje podatke. U tom je slučaju uobičajeno da se definira arhitektura mreže, a prije postupka obrade podatka obavi postupak učenja mreže ili treniranja.

Za razliku od konvencionalnih tehnika obrade podataka gdje je postupak obrade potrebno analitički razložiti na određeni broj algoritamskih koraka, kod ovog tipa neuronskih mreža takav algoritam ne postoji. Znanje o obradi podataka tj. znanje o izlazu kao funkciji ulaza, pohranjeno je implicitno u težinama veza između neurona. Te se težine postupno prilagođavaju kroz postupak učenja neuronske mreže sve do trenutka kada je izlaz iz mreže, provjeren na skupu podataka za testiranje, zadovoljavajući. Pod postupkom učenja kod neuronskih mreža podrazumijevamo iterativan postupak predočavanja ulaznih primjera (uzoraka, iskustva) i eventualno očekivana izlaza. Ovisno o tome da li nam je u postupku učenja *á priori* znan izlaz iz mreže, pa ga pri učenju mreže koristimo uz svaki ulazni primjer, ili nam je točan izlaz nepoznat, razlikujemo dva načina učenja: učenje s učiteljem – učenje mreže provodi se primjerima u obliku para (ulaz, izlaz), učenje bez učitelja – mreža uči bez poznavanja izlaza (Bašić, 2008, str. 12-13).

Mreža se nakon inicijalizacije može modificirati da poboljša izvedbu na temelju parova ulaz/izlaz. To se čini do te mjere da se algoritme za učenje može napraviti generalnima i efikasnim, što povećava vrijednost neuronskih mreža kao psiholoških modela i čini ih korisnim alatima za stvaranje širokog spektra aplikacija s visokim performansama. Pri tome je moć obrade mreže, pohranjena u snazi veza između pojedinih neurona tj. **težinama** do kojih se dolazi postupkom prilagodbe odnosno učenjem iz skupa podataka za učenje.

Neuronska mreža **obrađuje podatke distribuiranim paralelnim radom** svojih čvorova.

Neke osobitosti neuronskih mreža naspram konvencionalnih (simboličkih) načina obrade podataka su sljedeće:

- Vrlo su dobre u procjeni nelinearnih odnosa uzoraka.
- Mogu raditi s nejasnim ili manjkavim podacima tipičnim za podatke iz različitih senzora, poput kamera i mikrofona, te u njima raspoznavati uzorke.
- Robusne su s obzirom na pogreške u podacima, za razliku od konvencionalnih metoda koje pretpostavljaju normalnu raspodjelu obilježja u ulaznim podacima.
- Stvaraju vlastite odnose između podataka koji nisu zadani na eksplicitan simbolički način.
- Mogu raditi s velikim brojem varijabli ili parametara.
- Prilagodljive su okolini.
- Sposobne su formirati znanje učeći iz iskustva.

Neuronske mreže odlično rješavaju probleme klasifikacije i predviđanja, odnosno općenito sve probleme kod kojih postoji odnos između prediktorskih (ulaznih) i zavisnih (izlaznih) varijabli, bez obzira na visoku složenost te veze (nelinearnost). Danas se neuronske mreže primjenjuju u mnogim segmentima života poput medicine, bankarstva, strojarstva, geologije, fizike itd., najčešće za sljedeće zadatke:

- raspoznavanje uzoraka,
- obrada slike,
- obrada govora,
- problemi optimizacije,

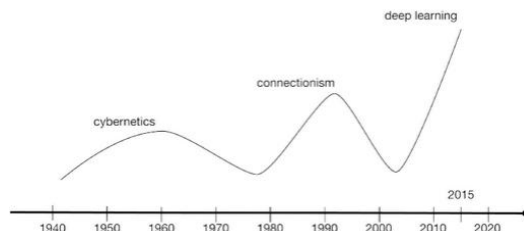
- nelinearno upravljanje,
- obrada nepreciznih i nekompletnih podataka,
- simulacije (Bašić, 2008, str. 9).

4.1.2. Duboko učenje

Tema o kojoj se sve više govori je i **duboko učenje** (engl. deep learning). Ono predstavlja podkategoriju strojnog učenja, a realizira se umjetnim neuronskim mrežama koje pomažu u prepoznavanju govora, računalnog vida i obrade prirodnog jezika.

Posljednjih godina duboko učenje doživljava nagli rast, kao što je vidljivo na slici 4.4. Razvoj dubokog učenja pomogao je kovanju napredka u područjima poput percepcije objekata, strojnog prevođenja i prepoznavanja glasa – svaki je od njih predmet kojega istraživači umjetne inteligencije dugo nastoje probiti. Ključni faktori poboljšanja izvedbe algoritama dubokog učenja su dostupnost velikih skupova podataka za treniranje, ostvareni masovnim povezivanjem računala te povećanje kapaciteta memorije i brzine računala. Veće umjetne neuronske mreže mogu danas izvoditi kompleksnije zadatke s većom točností.

Pri klasičnom strojnom učenju računalno bi kreiralo znanje kroz nadzirano iskustvo, što znači da je ljudski operater pomagao stroju u učenju dajući mu stotine ili tisuće praktičnih primjera za učenje, a greške su se ispravljale ručno. “Razlika današnjih sustava sa dubokim učenjem je u tome što sada istraživači nastoje konstruirati sustav koji sam kreira svoje osobine koliko god je to izvedivo. Dakle, duboko učenje je uglavnom bez nadzora. Uključuje, na primjer, neuronske mreže velikih razmjera koje omogućuju računalu učenje i samostalno “razmišljanje” bez potrebe za izravnom ljudskom intervencijom” (Dormehl L, 2014).



Slika 4.4. Povijest razvoja umjetnih neuronskih mreža

4.2. Evolucijsko računalstvo

Tehnike pretraživanja, u smislu traženja rješenja, zadnjih su godina pronašle inspiraciju u proučavanju biološke evolucije. Evolucijsko računalstvo čini skup algoritama za pretraživanje, proizašlih iz tehnika umjetne inteligencije koji se temelje na teoriji evolucije. Ona pruža alternativnu, veoma moćnu strategiju kada se zahtjeva pretraživanje, po mogućnosti najboljeg, rješenja problema, **odabirom** iz niza mogućih rješenja. Moguće je, ako je potrebno, ostvariti nova rješenja koja nisu prije razmatrana, odnosno ostvariti kreativnost.

Evolucijsko računalstvo sastoji se od ponavljanja sljedeće procedure:

1. Započni populacijom mogućih dizajna (kandidata)
2. Ocijeni svakog od njih pridavanjem rezultata iz "funkcije dobrote"
3. Zapamti dizajn koji dobije najbolji rezultat
4. Kreiraj novu populaciju odabirom dizajna najpodobnijih kandidata pa ih malo izmjeni nasumično ili ih izmjeni drastično zajedničkom kombinacijom različitih dizajna (Shasha i Lazere, 2010, str. 15).

U biološkom procesu evolucije, u nekom trenutku postoji populacija jedinki neke vrste koja čini generaciju. Te jedinke miješaju se međusobno (obično parenjem) da bi proizvele novu generaciju i tako, tijekom vremena vrsta preživljava i u najboljem slučaju napreduje. Kako se okoliš mijenja, da bi neka vrsta opstala ona se mora prilagoditi tim promjenama u okolišu. Međutim, ukupno gledano taj process je veoma spor, možda su potrebni milijuni godina. Oponašanjem (modeliranjem) ovog prirodnog procesa u računalima, moguće je postići tehniku koja prilagođava rješenja problema iz populacije potencijalnih rješenja. Pri tome se može postići najbolje moguće rješenje ili barem rješenje koje je funkcionalno. Različita rješenja u jednoj generaciji pomiješana su genetički parenjem, da bi proizveli novu, poboljšanu generaciju. Tako dobivena rješenja dalje se miješaju, na različite načine, da bi se ostvarila iduća generacija i tako dalje, sve dok se mnogo – ako je moguće tisuće – generacija kasnije ne dostigne mnogo bolje rješenje originalnog problema. "Softverske generacije moguće je dobiti u mnogo kraćem

periodu, nego prirodne. Radi se o milisekundama tako se ne mora čekati milijun godina da se pojavi rješenje” (Warwick, 2012, str. 102).

Najpoznatiji pristup evolucijskom računalstvu je **metoda genetskih algoritama** (GA). U toj je tehnici svaki član populacije definiran svojim genetskim kodom (računalnim kromosomima) koji ga jednoznačno opisuju. Taj opis može biti zapisan u binarnom obliku. Da bi se od jedne generacije dobila iduća, kromosomi jednog člana mješaju se (sparuju) s kromosomima drugog člana, a na kromosome se primjenjuju križanje i mutacije, što je također inspirirano biološkim procesima. (Fox i Fan, 2015.)

Kod binarnih kromosoma direktno je povezan s karakteristikama svakog člana (jedinke). Nastale razlike između pojedinih kromosoma odnose se na stvarne razlike u njihovim svojstvima. Kao jednostavan primjer možemo uzeti jednog člana **A** opisanog nizom **0101**, dok je drugi član **B** na primjer opisan nizom **1100**. Proces križanja sastoji se od toga da se uzme dio koda od A i da ga se pomiješa sa dijelom koda od B i tako se dobije novi član, pripadnik nove generacije. Na primjer, prvi dio (prve dvije znamenke) od A pomiješaju se sa drugim dijelom (zadnje dvije znamenke) od B daju **0100** - novi kod. Za duži kod odnosno više znamenki (što obično i jest slučaj u praksi) postupak je isti. **Proces mutacije**, koji se primjenjuje manje učestalo, uključuje odabir jedne znamenke (po mogućnosti nasumično) koja se zatim promijeni. Tako, ako uzmemo A kao 0101 i mutiramo ga mijenjajući treću znamenku, dobiti ćemo 0111 u sljedećoj generaciji.

Kod rada sa GA, prvi zadatak je konstrukcija kromosoma sa utvrđenom duljinom, koji predstavljaju populaciju, jednostrano određujući svaku pojedinu jedinku. Također je potrebno odrediti veličinu populacije i odlučiti hoće li se dozvoliti njezin rast. Ako je određeno da nema rasta populacije tada će neke jedinke u generaciji morati biti ubijene, odnosno neće dalje napredovati. To će biti najslabije jedinke ili pak (za potrebe raznolikosti) jedinke koje su veoma slične ostalima, ali ne toliko dobre pa mogu biti ubijene. Isto tako, ubijeni mogu biti i klonovi, zato što populacija identičnih jedinki nije poželjna. Zatim je potrebno odlučiti koliko će se križanja i mutacija pojaviti, što se uglavnom određuje empirijski, za potrebe specifičnog problema koji se rješava (Warwick, 2012, str. 104). Iz iteracije u iteraciju jedinke u populaciji poprimaju sve poželjnija svojstva. Algoritam obično završava dosezanjem zadanog broja iteracija. Kada je uvjet završetka ispunjen iz dobivene populacije odabire se najbolja jedinka i ona predstavlja rješenje problema.

Vjerojatno najbitniji aspekt u radu s genetskim algoritmima je određivanje mjerljivosti jedinki – što je dobro, a što je loše u zadanim osobinama. Za tu svrhu izrađuje se sveobuhvatna funkcija koja se naziva **funkcija dobrote** (engl. fitness function) kojom se određuje pogodnost jedinke. Funkcija dobrote dodjeljuje vrijednosti dobrote svakom kromosomu u trenutnoj populaciji. Dobrota (engl. fitness) kromosoma ovisi o tome koliko dobro taj kromosom rješava problem na koji se GA aplicira (Mitchell, 1999, str 4.). Funkcija dobrote može uključivati vrijednosti poput brzine, cijene, snage, duljine ili bilo čega što je značajno za određeni problem. Odabir prikladne funkcije dobrote obično je ključan faktor u primjeni genetskog algoritma.

Za pokretanje algoritma potrebna je inicijalna populacija koja se može dobiti nasumično ili putem određenog broja grubih procjena rješenja. Dobrota svakog pojedinog člana prve generacije određuje se funkcijom dobrote. Nakon toga odabire se par kromosoma za parenje - oni koji postignu bolji rezultat u dobroti imaju veće šanse za parenje. Primjene se križanje i mutacija, za svaki pojedini slučaj i kao rezultat dobiju se jedan ili više potomaka. Tada se proces može ponoviti na ostalim parovima kromosoma. Dobiveni rezultat je nova populacija sastavljena od originalnih kromosoma. Svaki od dobivenih kromosoma tada se testira primjenom funkcije dobrote te se iz populacije eliminiraju kromosomi koji dalje ne sudjeluju u procesu. Čitav process ponavlja se generaciju za generacijom – po mogućnosti više tisuća generacija – dok se ne uoči da nastupaju samo manje promjene u funkciji dobrote primjenjene na računanju najboljih kromosoma. U tom trenutku može se zaključiti da je postignuto rješenje. Ponekad su te promjene male u nekoliko generacija, pa vrijednost može čak i opadati pa se zatim ponovo poboljšavati što zapravo ovisi o kompleksnosti problema. Ponekad se odredi da se process zaustavlja nakon što je u nekoliko generacija uočena određena, zadana vrijednost funkcije dobrote i može se smatrati da je rješenje “*dovoljno dobro*”.

Jedna razlika između GA i prirodnih procesa je u tome što kod GA populacija ostaje stalna. Glavni razlog za to jest upravljanje algoritmom jer bi rast populacije zahtijevao da se za svaku novu jedinku ponovo izračunava funkcija dobrote, a to zahtijeva određene vremenske resurse u radu računala pa bi se velikim rastom populacije značajno povećalo vrijeme izračuna. Međutim, takav pristup smanjuje raznolikost u populaciji. Ponekad je potrebno da GA izvodi operacije prilagodbom na nastale promjene u uvjetima, kao na primjer robot koji se treba nositi sa

različitim uvjetima u okolišu. U takvim je situacijama raznolikost u populaciji jako bitna kako bi se relativno brzo pojavile velike promjene u individui ako i kada je to potrebno.

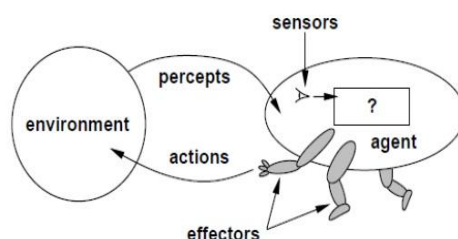
4.3. Inteligentni agenti

Agent je nešto što može percipirati svoju okolinu putem **senzora** (osjetila) i djelovati na tu okolinu putem **efektora**. **Ljudski agent** ima oči, uši i ostale organe za senzore te ruke, noge, usta i ostale dijelove tijela za efektore. **Robotski agent** zamjenjuje kamere i infracrvene lokatore za senzore, a različite motore za efektore. **Softverski agent** ima kodirane nizove bitova za perceptore i akcije koje izvodi (Norvig i Russell, 1995, str 31).

Temeljeno na konceptu agenta, predložena je **nova definicija umjetne inteligencije** :

- Umjetna inteligencija je grana računalstva. Njezin je cilj konstruiranje agenata koji pokazuju određena inteligentna ponašanja.

Tako je istraživanje inteligentnih agenata u srži problema umjetne inteligencije. Intelligentno računalo je prvi cilj, a ultimativni cilj je umjetna inteligencija (Shi, 2011, str. 502). Osnovna karakteristika **racionalnog agenta** bila bi da čini ispravnu stvar. Može se reći da su ispravne akcije one u kojima će agent biti najučinkovitiji. Dakle, to zahtjeva odluku o tome kako i na koji način provesti evaluaciju agentovog uspjeha.



Slika 4.5. Interakcija agenta s okolinom putem senzora i efektora

U svrhu evaluacije uspjeha koristi se izraz **mjera izvedbe** (engl. performance measure) kod određivanja kriterija uspješnosti agenta. Pri tome se inzistira na objektivnom mjerenju, nametnutom od strane nekog autoriteta. Drugim rječima, mi kao vanjski promatrači postavljamo standard o tome što znači biti uspješan u okolini i koristimo ga za mjerenje izvedbe agenata. Kao

primjer, zamislimo slučaj agenta koji bi trebao očistiti prljavi pod. Prihvatljiva mjera izvedbe bila bi količina prljavštine koju je agent očistio u osam sati. Sofisticiranije mjerenje uključivalo bi i faktor potrošnje električne energije kao i faktor stvaranja količine buke u radu agenta. Također je važan i vremenski faktor, odnosno ukoliko želimo mjeriti količinu prljavštine koju je agent sakupio u prvom satu, nagradili bi agenta koji je u početku najučinkovitiji (čak kada bi kasnije radio malo ili nimalo), a kažnjavali bi one koji rade konzistentno. Ovdje treba biti oprezan u smislu da se razlikuje racionalni agent od sveznajućeg. Sveznajući agent zna stvarni ishod svojih akcija i prema njemu se može ponašati, ali to je u stvarnosti nemoguće. “Zapravo, ovdje treba istaći da se racionalnost bavi očekivanim uspjesima s obzirom na to što je percipirano. Bitno je da, ako označimo kako bi inteligentni agent trebao uvijek činiti ispravne stvari, postaje nemoguće konstruirati agenta koji udovoljava ovim specifikacijama – osim ako ne poboljšamo izvedbe kristalnih kugli” (Norvig i Russell, 1995, str 33).

Racionalnim se u bilo koje vrijeme može smatrati nešto što udovoljava navedenim kriterijima:

1. Mjeri izvedbe koje definira stupanj uspjeha.
2. Sve što je dosad agent percipirao (sveukupna povijest percipiranja naziva se slijed percepta).
3. Što agent zna o okruženju.
4. Akcije koje agent može izvesti.

Iz ovoga se izvodi **definicija idealnog racionalnog agenta**:

Za svaki mogući slijed percepta, idealni racionalni agent trebao bi poduzimati bilo koje akcije kojima očekuje da će maksimalno povećati svoju mjeru izvedbe, temeljeno na dokazima dobivenim preko sljeda percepta i bilo kojeg ugrađenog znanja koje posjeduje.

Ponašanje agenta ovisi isključivo o slijedu percepta do određenog trenutka. Tada je moguće opisati bilo kojeg agenta izradom tablica akcija koje se poduzimaju kao odgovor na svaki mogući slijed percepta. Takva se lista naziva **mapa**; od sljeda percepta do akcija. Idealno mapiranje tada opisuje idealnog agenta. Određenjem akcija koje agent treba poduzeti kao odgovor na bilo koji slijed percepta omogućuje se kreiranje idealnog agenta. Ponašanje agenta može biti zasnovano i na njegovom vlastitom iskustvu i na ugrađenom znanju prilikom izgradnje agenta, za određeno okruženje u kojemu agent djeluje.

4.3.1. Osnovna građa inteligentnih agenata

Zadaća tehnika umjetne inteligencije jest kreiranje **agent programa**: funkcije koja implementira mapiranje agenta od percepta do akcija. Pretpostaka je da će se program izvoditi na nekoj vrsti računalne naprave koju nazivamo **arhitektura**.

Program se odabire tako da ga arhitektura prihvaća i izvodi. Arhitektura može biti jednostavno računalo ili može uključivati hardver za specijalne namjere poput obrade slika iz kamere ili filtracije audio unosa. Može uključivati softver koji omogućava izolaciju između temeljnog računala i agent programa, kako bi se moglo programirati na višoj razini.

Općenito; arhitektura omogućava perceptore iz senzora dostupnima program, izvodi program, i dobavlja programske akcije efektorima kako su one generirane. Odnos između agenata, arhitekture i programa može se sažeti ovako:

$$\text{AGENT} = \text{ARHITEKTURA} + \text{PROGRAM}$$

Pojavnost sveukupnog kompleksnog inteligentnog ponašanja postiže se kroz skup interakcija jednostavnih entiteta koji su sami poluautonomni agenti. Ti agenti (kao kod UNM) mogu biti u formi neurona koji su jednostavno ulančani zajedno, a inteligentno ponašanje ostvaruju svojom brojnošću. Kod GA moguće je da se populacija gena, kao agenata, poboljšava putem evolucijskog procesa uz vanjsku procjenu – funkciju dobrote. U svakom je slučaju vidljivo da pojedini agenti imaju malo ili nimalo znanja o tome što ostali agenti čine. Oni su relativno neovisni, ali ipak pod utjecajem ostalih agenata u smislu postizanja ciljeva u okolini. Konačni rezultat može ostvariti samo jedan agent (u slučaju GA) ili skup zajednice agenata (u slučaju UNM).

Autonomnost sustava određena je mjerom kojom je njegovo ponašanje određeno vlastitim iskustvom. Zaista autonoman inteligentni agent trebao bi samostalno djelovati uspješno u različitim vrstama okruženja, ukoliko ima dovoljno vremena za prilagodbu.

Autonomni agent je sustav koji:

- je smješten u okolinu,

- je dio te okoline,
- koju osjeća,
- na koju djeluje,
- u vremenu,
- s ciljem ostvarenja vlastitih planova (nijedan čovjek ne navodi njegove odabire u akcijama),
- kako bi te akcije mogle utjecati na njegova buduća opažanja (on je strukturno udružen sa svojom okolinom)

Posljednja osobina koja je navedena, razlikuje autonomne agente od ostalog softvera. Pretpostavka je da bi neki generalno inteligentni sustav trebao biti autonomni agent iz razloga što je za generalizaciju znanja, između različitih i vjerojatno novih područja, potrebno učenje. Učenje zahtjeva osjetila, a često i izvođenje akcije. “Autonomni je agent pogodan za učenje, pogotovo učenje nalik ljudskom. Da bi sve to izvodio, agent mora imati ugrađene senzore za osjetila, efektore za izvršenje akcije i mora imati primitivne motivatore, koji motiviraju njegove akcije. Senzori, efektori i motivatori jesu primitivi koji moraju biti ugrađeni u agenta ili se moraju razviti u bilo kojem agentu” (Goertzel i Wang, 2007, str 38).

4.3.2. Primjena inteligentnih agenata

Jedan pristup UI je da se posebno fokusira na ideju agenata, i njihove individualne identitete kako bi se proizvelo ukupno nastalo ponašanje. Svaki element može se smatrati dijelom društva koje obično percipira ograničene aspekte svoje okoline, a na koju može djelovati ili samostalno ili u suradnji s ostalim agentima. Na taj način pojedinačni agenti usklađuju svoje ponašanje s ostalima kako bi obavili specifičnu zadaću. Ključna razlika između tog pristupa i klasične umjetne inteligencije je u tome što je sveukupna inteligencija raspodijeljena između agenata umjesto da je smještena u nekom centralnom središtu.

U korištenju agenata za rješavanje problema može se kompleksni problem razlomiti na manje probleme od kojih je svaki mnogo lakši za rješavanje. Agenti se tada mogu koristiti za rješavanje tih manjih problema – udružujući se u ostvarivanju konačnog rješenja. Jedna prednost ovoga je u tome da svaki agent sadrži informaciju o svom “manjem” problemu – ne mora znati ništa o

generalnom problemu. Međutim, mnogo je načina na koje se ovakva primjena može ostvariti, pa se tako mogu pronaći različite definicije toga što je agent i što može raditi.

“Neki agenti imaju zadane akcije, dok su neki fleksibilni i prilagodljivi. Neki su autonomni, neki su kompletno ovisni o odlukama ostalih. Većina je responzivna na okolinu u kojoj postoje iako to može biti u smislu okoline vanjskog svijeta ili akcija agenata koji ga okružuju. Zamislite jedan neuron u sredini svog mozga, na primjer, on je pod utjecajem ostalih neurona, ne izravno pod bilo kakvim vanjskim utjecajem” (Warwick, 2012, str. 109).

Moguće je kreirati agentski sustav tako da svi agenti imaju istu snagu i sposobnosti; međutim može se odrediti da neki agenti nadvladaju odluke ostalih agenata – tada govorimo o **slojevitoj arhitekturi** u kojoj akcija ili odluka jednog agenta, nižeg prioriteta, biva prevladana odlukom agenta višeg prioriteta.

4.3.3. Softverski agent – softbot - bot

Računalni agenti tranutačno su veoma aktualna tema u umjetnoj inteligenciji kao i prijelomna točka u razvoju softvera nove generacije. Postoje široki spektar mogućih softverskih agenata. Na primjer, takvi se agenti danas koriste kod promatranja financijskih tržišta. Oni u realnom vremenu provjeravaju kretanja dionica i njihove cijene na tržištima. Agent može biti i često jest odgovoran za kupovinu i prodaju roba pa u tom slučaju mora znati ako (čovjek) dobavljač ima u nekom trenutku smjernice koje dionice su zanimljive, a koje treba izbjegavati – u tom slučaju softverski agent može trebati “razumijevanje” instrukcija govornog jezika. Agenti su idealni za ovakve transakcije zato što jednostavno miruju i promatraju aktivnosti, a rednju izvode samo kada se pojave “pravi uvjeti”. Ne samo što je ovo teško izvodivo za ljude, već jednom kada je odluka potrebna, agent je može izvesti gotovo smjesta. U vremenu koje bi ljudskom brokeru bilo potrebno da napravi istu odluku (nekoliko sekundi ili minuta) prodaja može već biti gotova. Rezultat toga je da veoma veliki omjer dnevnih financijskih transakcija širom svijeta izvode, ne ljudi, već softverski agenti. Uredi financijskih kuća u gradskim središtima u Londonu i New Yorku zagušeni su računalima. Brokeri koji su nekada izvršavali transakcije sada su uključeni tako što promatraju UI agenata, dobavljaju za njih informacije i povremene naredbe – nakon toga ih puštaju u rad. Istovremeno su drugi uključeni u izradu novih UI agenata. “Više nije slučaj da

tvrtka koja izvede najbolju pogodbu zaradi najviše novca, sada su to tvrtke koje ostvare najbolje UI agente” (Warwick, 2012, str. 110). Takav je agent u stanju nadzirati mnoštvo faktora istovremeno, čuvajući povijesne podatke o cijeni dionica kroz neko razdoblje; trendovima u istraživanjima; uspoređujući te podatke s ostalim dionicama; povezujući ih s omjerima u financijskim transakcijama i ostalim vanjskim informacijama prevedenim iz svakodnevno novih podataka. Kako je udruženo mnogo faktora, moguće je potrebno ugraditi u agenta neke od tehnika rudarenja podataka ili agent treba imati pristup podatcima kako i kada treba je potrebno.

Osnovna akcija koju agent izvodi jest da uzima informaciju s jednog od brojnih ulaza, obradi tu informaciju, poveže je sa povijesnim podacima te donese odluku za djelovanje bilo fizičku ili u obliku softverskog izlaza. Ovo se može postići dizajnom agenta temeljenom na pravilima ili tablicama koje sadrže podatke. Ukoliko se povijesni podaci ignoriraju tada možemo govoriti o **refleksnom agentu**. Kada agent sadrži elemente koje zahtijevaju planiranje kako bi postigli interni cilj ili se usmjeravaju prema vanjskom cilju tada govorimo o **agentu orijentiranom ka cilju**. Istodobno, ako se elementi planiranja prilagođavaju kako bi na odgovarajući način reagirali na vanjske promjene u okolini, po mogućnosti kao posljedica vlastitih akcija agenta, onda govorimo u **agentima koji uče**. Konačno, agenti se mogu temeljiti na modelima dobivenim iz stvarnog svijeta kojega onda nastoji oponašati u izvedbi i tada se naziva **agent temeljen na modelu**.

4.3.4. Višeagentski sustavi i raspodijeljena inteligencija

U 1990-ima multiagentski sustavi MAS (engl. multi-agent system) postali su veoma zastupljeni u istraživanjima distribuirane inteligencije. Višeagentski sustavi pretežno istražuju inteligentno ponašanje prilikom koordinacije više agenata. Za postizanje globalnog, zajedničkog cilja ili njihovih pojedinačnih ciljeva, agenti dijele znanje o problemu i pristupima rješavanju kroz međusobnu suradnju. Raspodijeljena inteligencija najviše istražuje rješavanje problema u fizički ili logički raspodijeljenim sustavima (Shi, 2011, str. 499).

Raspodijeljenu inteligenciju možemo definirati kao:

- područje umjetne inteligencije koje proučava principe i oblikovanje višeagentskih sustava.

Od 1970-ih raspodijeljena inteligencija postaje sve zanimljiviji predmet istraživanja postepenim razvojem računalnih mreža, komunikacija i tehnologijom paralelnog dizajna računalnih programa. Naglim razvojem interneta u 1990-ima pružaju se nove, izvrsne mogućnosti za primjenu informacijskih sustava i DDS (eng. Decision Support System) sustava. Opseg, domet i kompleksnost primjene novih informacijskih sustava povećani su danas na visoki stupanj, a raspodijeljena inteligencija u njima igra ključnu ulogu. Osobitosti takvih sustava jesu:

- Podatci, znanje i kontrola raspodijeljeni su ne samo logički, već fizički.
- Svaki sudionik u rješavanju problema povezan je računalnom mrežom.
- Svaka komponenta surađuje s ostalima kako bi se riješio problem koji jedna komponenta samostalno ne može.

Implementacija raspodijeljene umjetne inteligencije može nadoknaditi nedostatke tradicionalnih ekspertnih sustava i veoma poboljšati izvedbu sustava znanja. Prednosti sustava raspodijeljene inteligencije uključuju:

- **Poboljšanje kapaciteta rješavanja problema** - zbog raspodijeljenosti karakteristika inteligentnog sustava, prije svega njegov se kapacitet rješavanja problema značajno povećava i cijeli sustav reducira vrijeme odgovora i točnost rješenja, također je otporniji na greške. Zatim, sustav je lakše proširivati, a zahvaljujući obilježjima pojedinih modula moguće je cijeli sustav napraviti veoma fleksibilnim.
- **Poboljšanje učinkovitosti rješavanja problema** - zbog toga što čvorovi raspodijeljenog inteligentnog sustava mogu paralelno rješavati problem i tako povećati učinkovitost.
- **Proširenje dosega primjene** - u raspodijeljenim inteligentnim sustavima, različita područja i različiti stručnjaci iz istog područja mogu surađivati prilikom rješavanja problema kada jedan određeni stručnjak to ne može sam. U isto vrijeme moguće je uključiti više nestručnjaka kod rješenja problema i dobiti jednaki učinak ili takav da nadilazi eksperta.

- **Reduciranje kompleksnosti softvera** - raspodijeljeni sustavi rastavljaju zadaće na podzadaje pa tako reduciraju kompleksne probleme.

Tamo gdje se primjenjuju višeagentski sustavi, oni mogu djelovati u suradničkom odnosu tako da svaki od agenata dobavi dio odgovora na problem, a sveukupno rješenje dobije se skupnom kohezijom izlaza sa brojnih agenata. Također je moguće postaviti agente da djeluju natjecateljski, pojedinačno ili u grupi pa samo mala grupa aktivnih agenata dobavlja sveobuhvatno konačno rješenje.

Raspodijeljene višeagentske sustave (u vidu senzora) primjenio je Alasdair Allan na projektu mreže raspodijeljenih peer-to-peer teleskopa, koji su djelujući autonomno, rasporedili opažanja tako da reagiraju na vremenski kritične događaje i u to doba (2009. godine) polučili su golemi uspjeh u vidu otkrića tada najudaljenijih objekata ikada pronađenih. On smatra da će se geografski raspodijeljeni agenti koji preko senzora prikupljaju podatke iz svoje okoline, primjenjivati sve više na svakodnevne objekte. Neizbježno, količina računala koji će umjesto nas (u drugom planu) izvršavati zadaće prikupljanja i obrade podataka u realnom vremenu povećavat će se, a za desetak godina svaki dio odjeće, nakita i svega što nosimo uz sebe, mjeriti će, vagati i računati, a svijet će biti pun senzora. ”Prikupljeni podaci biti će u oblaku, u halou uređaja čiji će zadatak biti da nam pružaju usluge opažanja i konstantnog računanja, dok šetamo oni će izračunavati, međusobno se konzultirati, predviđati i iščekivati naše potrebe. Bit ćemo okruženi mrežom distribuiranih senzora i računala” (Allan A., 2013).

4.4. Big data – veliki podaci

U istraživanjima na području astronomije i genetike javlja se tijekom 2000-ih eksplozija velike količine sakupljenih podataka. Iz genomike nastao koncept **velikih podataka** (engl. **big data**) proširio se od tada na sva područja ljudske djelatnosti (Mayer-Schönberger, 2013, str. 16).

U srži pojma big data nalazi se predviđanje. Iako se opisuju kao grana računalne znanosti i umjetne inteligencije, odnosno područja strojnog učenja, ova osobina je zavaravajuća. “S big data ne pokušava se naučiti računalo da “misli” poput ljudi. Ovdje se radi o primjeni matematike na goleme količine podataka da bi se dobio zaključak temeljen na vjerojatnosti: koja je

vjerojatnost da je pristigla email poruka spam; da je ispis “hte” zapravo “the”, da su putanja i brzina pješaka koji pretrčava cestu dovoljni da on pređe na dugu stranu tako da vozilo bez vozača treba samo malo usporiti. Ključno je da ovakvi sustavi imaju dobru izvedbu zbog toga što su hranjeni s mnogo podataka na kojima temelje svoja predviđanja. Štoviše, načinjeni su tako da sami sebe poboljšavaju s vremenom, na način da bilježe najbolje signale i uzorke za promatranje dok se hrane novim podacima.

U budućnosti – i prije nego što mislimo – mnoge pojavnosti našeg svijeta koje su danas isključivo u domeni ljudske prosudbe biti će proširene ili zamjenjene računalnim sustavima. Kao što je internet radikalno promijenio svijet dodajući komunikaciju računalima, tako će koncept big data promijeniti temeljne aspekte života dodajući im kvantitativnu dimenziju koju prije nisu imali (Mayer-Schönberger, 2013, str. 29). “Digitalno doba možda je olakšalo i ubrzalo obradu podataka, omogućilo računanje milijuna brojeva u kratkom trenu. Ali kada kažemo da podaci “govore” - mislimo nešto više i drugačije.

Big data se temelji na tri bitna pomaka u razmišljanju koji su povezani i međusobno se pojačavaju. Prvi je sposobnost analize velike količine podataka iz određenog područja, umjesto da se promatra samo manji skup podataka. Drugi je nastojanje da se usvoje “zbrkani podaci” iz stvarnog svijeta, bez inzistiranja na privilegiji točnosti. Treći je povećano uvažavanje korelacija umjesto nekada nedostižne uzročnosti” (Mayer-Schönberger, 2013, str. 44).

Korelacije su korisne u svijetu malih podataka, ali u kontekstu velikih podataka one dolaze do sjaja. Pomoću njih lakše, brže i jasnije nego prije odabiramo uvide. U svojoj srži korelacije kvantificiraju statističke odnose između dvaju vrijednosti podataka. Jaka korelacija znači da kada se promijeni vrijednost nekog podatka, vrijednost drugog ima veliku vjerojatnost da se također promijeni. Kod korelacija radi se o vjerojatnosti, ne o sigurnost, međutim kada je korelacija jaka, vjerojatnost povezanosti je visoka.

“Mnogi korisnici Amazona mogu se u to uvjeriti odabirom police sa preporučenim knjigama. Inovativna preporuka od strane Amazon sustava rezultate dobiva stvaranjem korisnih korelacija među njima, bez poznavanja temeljnih uzroka. Znati *što*, ne *zašto*, je dovoljno dobro” (Mayer-Schönberger, 2013, str. 115).

5. KOMERCIJALNA PRIMJENA UMJETNE INTELIGENCIJE

“S masivnom količinom računalne snage, strojevi sada mogu prepoznavati objekte i prevoditi govor u realnom vremenu. Umjetna inteligencija napokon postaje pametnija” (Hof, 2013).

Svjedoci smo eksplozije primjene umjetne inteligencije. Ona podrazumjeva slabu UI, odnosno primjenu pametnih algoritama u aplikacijama za rješavanje kompleksnih problema. Taj uzlet rezultat je kombinacije nekoliko faktora poput; jeftinog paralelnog računanja u oblaku, velikih podataka i sve boljih algoritama za strojno učenje. Danas možemo govoriti o tome da: “umjetna inteligencija već mijenja svijet.” Naslovi poput ovoga u svim svjetskim medijima postali su uobičajeni, a uzlet umjetne inteligencije se uspoređuje s industrijskom revolucijom. Najavljuju se duboke promjene u društvu i na tržištu rada temeljene na predviđanjima da će *strojevi* u potpunost preuzeti obavljanje mnogih rutinskih poslova koje sada obavljaju ljudi. Razlog je sve šira primjena metoda razvijenih u umjetnoj inteligenciji na različita područja ljudske djelatnosti, od fundamentalnih znanosti poput fizike i biologije do upravljanja kompleksnim sustavima.

Rast umjetne inteligencije koji sada vidimo povezan je sa rastom velikih podataka. “Jednostavno rečeno, eksplozivan rast podataka je osnova za stvaranje slabe umjetne inteligencije. Ako tome dodate podatke koje generiraju korisnici – gotovo 300 000 tweetova, 220 000 fotografija na Instagramu, 72 sata YouTube video sadržaja, 2.5 miliona komadića sadržaja koji dijele korisnici Facebooka, u *svakoj minuti*, koje poslovne tvrtke moraju nadgledati i prihvatiti – jasno je kako nijedna organizacija niti (ili čovjek) nije u stanju nositi se s time bez pomoći umjetne inteligencije” (Senior, 2015).

Poslovne i osobne pogodnosti koje proizlaze iz UI nadilaze sposobnosti kretanja kroz goleme količine informacija i automatizacije poslova rutinskog znanja. Neki UI obrađuju dobivene podatke tako da sakupe i istaknu ono što je značajno i vrijedno pojedinim korisnicima i njihovom stanju “*potrebe*”. “Umjetna inteligencija je svuda oko nas - mi više nemamo ruku na zatvaraču. Jednostavan čin povezivanja s nekim putem tekstualnih poruka, e-maila ili telefonskih poziva koristi inteligentne algoritme za usmjeravanje informacija. Gotovo svaki proizvod kojega dotaknemo dizajniran je u suradnji ljudi i umjetne inteligencije, a zatim izgrađen u

automatiziranim tvornicama. Kada bi svi UI sustavi odlučili sutra stupiti u štrajk, naša bi civilizacija bila obogaljena: ne bismo dobili novac od banke, uistinu, naš novac bi nestao; komunikacije, transport i proizvodnja bili bi zaustavljeni. Srećom, naši inteligentni strojevi nisu još dovoljno inteligentni da bi organizirali takvu zavjeru“ (Kurzweil, 2012, str. 325).

5.1. Primjena inteligentnih sustava

Postoje slušna pomagala sa algoritmima koji filtriraju pozadinsku buku, sustavi za pronalaženje rute prikazom mape i ponudom savjeta u navigaciji za vozače, sustavi za preporuku knjiga i muzike (obzirom na prijašnje kupovine korisnika) i sustavi za pomoć u medicinskim odlukama kod rane dijagnoze karcinoma. Postoje robotski kućni ljubimci, roboti koji čiste, kose travu, kiruški roboti, roboti spasioci i više od milijun industrijskih robota. Svjetska populacija robota danas premašuje 10 milijuna. Izrađuju se moderni programi za prepoznavanje govora u kojima sustavi koriste statističke tehnike strojnog učenja koje automatski grade modele iz promatranja uzoraka koji se koriste. Kod programa za prevođenje koriste se isti modeli pa tako programeri u izgradnji sustava ne moraju uopće govoriti jezike na kojima rade (Bostrom, 2014. str 31).

Bitna vrijednost algoritma je upravo u njihovoj brzini, koja je pak određena hardverom. Zato što algoritmi omogućuju rješavanje kompleksnih zadataka velikom brzinom postali su pokretačka snaga današnjeg svijeta, ostvareni u kodovima i telekomunikacijskim infrastrukturnama koje čine “mrežu”. Algoritam koji odlučuje koji film ćete gledati, poput onih koje primjenjuje Netflix, je relevantan zbog brzine kojom može provjeriti na tisuće ulaznih faktora korektno i rezultat vratiti korisniku gotovo trenutačno. Kada rezultat ne bi bio gotovo trenutačan, tada kao alat ne bi bio toliko učinkovit. ”Danas nam algoritmi dozvoljavaju da u svoje dane uvrstimo više sadržaja. Činimo više stvari u manje vremena, u svijetu algoritama i računala to je poznato kao uklanjanje latencije” (Steiner, 2012, str. 232).

Pokretačka snaga razvoja boljih i pametnijih algoritama je globalno tržište s uvijek većim zahtjevima na tehnologijama koje su jeftinije, koje štede vrijeme i energiju ili se pak koriste u predviđanju. Zadnjih godina nastale su brojne tvrtke koje specijalizirane u umjetnoj inteligenciji, a najviše se koriste primjenom algoritama strojnog učenja. Primjeri sustava u kojima se danas koriste umjetna inteligencija jesu: autonomna vozila, sustavi za prepoznavanje slika, sustavi za

prepoznavanje prirodnog govora, program za odgovore -Wolfram Alpha, sustavi za preporuke koji *uče* navike korisnika, odnosno ostvaruju personalizaciju, sustavi za automatsku trgovinu, sustavi za nadzor, obranu i sigurnost, sustavi za inteligentno raspoređivanje.

Novost u umjetnoj inteligenciji je zaista mnogo javno i komercijalno dostupnih primjera. **Autonomna vozila** kao što su Googleov automobil ili Daimlerov poluautonomni kamion Inspiration, tehnologije su za koje se smatra kako će dovesti do znatnog smanjenja nesreća na prometnicama, povećati kapacitet prometa, olakšati ljudima zahtjevnost izvođenja vožnje i mnoge druge pogodnosti.



Slika 5.1. Samoupravljaajući automobil

Vozilima bez vozača već je legalno dozvoljeno upravljanje na javnim cestama (u američkoj saveznoj državi Nevadi), uz određene restrikcije, iako se masovna primjena u javnosti širom svijeta očekuje tak krajem desetljeća. Kraljevina Nizozemska je 2014. najavila petogošnji plan pripreme za uvođenje autonomnih kamiona na svojim prometnicama. Tehnologija koja promatra cestu i upozorava vozača na nadolazeće opasnosti već je ugrađena u vozila. Jedna takva tehnologija zasniva se na uspješnom modelu vizualnog procesuiranja u mozgu koja je kreirana na MIT-u. Naziva se **MobilEye** i sposobna je upozoriti vozača na opasnosti poput nadolazećeg sudara ili djeteta koje pretrčava ispred automobila. Tu tehnologiju primjenjuju u automobilima proizvođači kao što su Volvo i BMW.

Prepoznavanje lica je dovoljno poboljšano zadnjih godina pa se sada koristi u automatskom prelazu granice između Europe i Australije. Ministarstvo obrane Sjedinjenih Američkih Država upravlja sa sustavom za prepoznavanje lica sa preko 75 milijuna fotografija od aplikacija za vizu. “Sustavi za nadzor koriste rastuću sofisticiranu UI tehniku rudarenja podataka za analizu

glasa, videa ili teksta u velikim količinama pobranog iz svjetskih medija za elektroničku komunikaciju te pohranjenog u ogromnim centrima za podatke” (Bostrom, 2014, str. 33).

Prepoznavanje govora, odnosno “Razumijevanje prirodnog jezika, pogotovo kao ekstenzija automatskog prepoznavanja govora, sada je ušao u mainstream. Siri, automatizirana osobna asistentica na iPhone 4S mobilnim uređajima uskomešala je svijet mobilnog računalstva.” Siri možete pitati da učini gotovo sve što bi svaki pametni telefon sa samopoštovanjem mogao učiniti (na primjer, “*Gdje mogu pronaći indijsku hranu u blizini?*” ili “*Pošalji mojoj ženi poruku da sam krenuo*” ili “*Što ljudi misle o novom filmu Breda Pitta?*”) i gotovo svaki put Siri će udovoljiti. Siri će i zabavljati uz malo neproduktivnog čavrljanja. Ako je zapitate što je smisao života, odgovorit će “42”, što će ljubitelji Vodiča kroz galaksiju za autostopere prepoznati kao - odgovor na ultimativno pitanje o životu, svemiru i svemu ostalom” (Kurzweil, 2012, str. 330).

Program Wolfram Alpha je program za odgovore (za razliku od programa za pretraživanje) kojega je razvio britanski matematičar Stephen Wolfram. Godine 2002, nakon 10 godina istraživanja, Wolfram je objavio knjigu *Nova vrsta znanosti* u kojoj objašnjava svoje kontroverzne poglede na neadekvatnost znanosti baziranoj na matematici da otključa tajne prirodnog svijeta. Predložio je gledište da bi se kompleksnost prirode mogla bolje shvatiti kroz proučavanje računalnih modela temeljenim na staničnim automatima – uljučujući primjene u svakojakim znanstvenim nastojanjima, poput vremenske prognoze, rast umjetnih organizama, objašnjenje ponašanja burze kapitala kao i razumijevanje porijekla svemira. Priroda, kako on tvrdi, funkcionira kao računalo. Godine 2009. Wolfram Research (tvrtka koju je osnovao Stephen Wolfram) predstavio je Wolfram Alpha, program za pretraživanje dizajniran da odgovara na osnovna pitanja, posebno ona koja se mogu izraziti u jednadžbama i pri tome koristi veliku bazu podataka umjesto da pretražuje internet. Na primjer, ukoliko pitate Wolfram Alpha: “*Koliko ima prim brojeva manjih od milijun?*” dobiti ćete odgovor 78 498. Program nije potražio odgovor, izračunao ga je, korisniku vratio odgovor i nakon toga je prikazao jednadžbe koje je koristio za dobivanje rezultata. “Ako u konvencionalnu tražilicu upišete isto pitanje, dati će vam poveznice na kojima možete naći traženi algoritam. Tada bi trebali te jednadžbe uključiti u sustav kao što je Matematica (također razvijen od strane dr. Wolframa), ali očito je da takav način traži puno više posla i razumijevanja, nego kada jednostavno upitate Alphu” (Kurzweil,

2012, str. 351). Alpha za pravo sadrži 15 milijuna linija koda iz Matematice. Ono što čini Alpha jest da doslovno računa odgovor iz podataka kojima raspolaže (oko 10 trilijuna bajtova). Možete postaviti široki raspon pitanja, na primjer: *“koja zemlja ima najveći BDP po broju stanovnika.”* Alpha se koristi i kao dio Appleovog Siri programa; ako Siri postavite činjenično pitanje, onda se ono prebacuje Alphi na obradu. Alpha se isto tako koristi kod pretraživanja predstavljenim u Microsoft Bing pretraživaču. “To je impresivan sustav, koji koristi izrađene metode i ručno obrađene podatke. To je testament razlozima zbog kojih smo u prvom redu i kreirali računala. Dok mi otkrivamo i sastavljamo znanstvene i matematičke metode, računala su mnogo bolja u njihovim implementacijama, nego je to ljudska inteligencija bez pomagala. Većina poznatih znanstvenih metoda kodirane su u Alphi, zajedno sa kontinuirano ažuriranim podacima iz različitih područja, od ekonomije do fizike“ (Kurzweil, 2012, str. 352). “Dokazivanje teorema i rješavanje jednadžbi sada su već toliko utvrđeni da ih se gotovo ne primjećuje kao umjetnu inteligenciju. Rješavatelji jednadžbi uključeni su u automatske dokazivače teorema kao što je Matematica, Formalne metode verifikacije uključujući automatske dokazivače teorema, rutinski se koriste kod proizvođača čipova za provjeru dizajna električnih krugova prije proizvodnje” (Bostrom, 2014, str 31).

Sustavi za preporuke koji u radu koriste strojno učenje potihno preuzimaju globalna tržišta. Danas su veoma rašireni u primjeni, a namjenjeni su korisnicima različitih usluga poput e-trgovine. Kada je pokrenut, 1995. godine, Amazon je bio knjižara koja je u sebi imala ugrađen personalizirani sustav za preporuke koji je funkcionirao tako da trenutačno daje rezultate (preporuke) korisniku. Google nam može isporučiti pametnije oglašavanje, Netflix, svojim korisnicima nudi preporuke na temelju njihovih preferenci odnosno onoga što su korisnici već gledali ili spremili za gledanje. “U Amazonu je pritisak za korisničkim podacima beskrajan: dok čitate na vašem Kindle uređaju, podaci o frazama koje ističete, koje stranice birate, da li čitate redom ili preskačete dijelove, se dostavljaju u Amazonov server i mogu se koristiti za indicaciju koja bi vam se knjiga sljedeća mogla svidjeti. Kada se nakon čitanja e-knjiga s Kindle uređaja na plaži ulogirate, Amazon je u mogućnosti suptilno prilagoditi svoju stranicu prikladno onome što ste čitali: Ako ste proveli puno vremena s posljednjim djelom Jamesa Pattersona, ali samo ste

letimično pogledali novi vodič za prehranu, mogli bi ugledati mnogo komercijalnih trilera i manje knjiga o zdravlju” (Pariser, 2011, str. 59).

Automatizirana trgovina dionicama je primjer primjene softvera na tržištima kapitala na kojem se mogu uočiti koristi od takvih sustava, ali i nepredvidljive posljedice uzrokovane *čudnim* ponašanjem softvera. Botovi su 2000. godine doprinosili manje od 10 % u svakom trgovanju dionicama u SAD-u. Veliki igrači na Wall Streetu znali su za njihovo postojanje, ali tada algoritmi nisu pomicali tržišta. Visokofrekventno trgovanje dionicama koje izvode računala sve to mijenja. Na Wall Streetu su algoritamski ratovi postali na neki način natjecanje u tome tko će istu stvar učiniti brže. “Kako se je Wall Street punio hakerskim talentima kroz 1990-e tako je sve više tvrtki nastojalo primijeniti iste strategije i slične algoritme. Da bi stvorili nekakvu prednost u odnosu na ostale, trgovci su sve više nastojali dobiti bolji hardver, bolja računala i bolje telekomunikacijske linije “(Steiner, 2012, str. 232). Do početka 2008. godine, automatski botovi na tržištima dionica činili su 60 % svog prometa dionicama, a financijska industrija provela je sedam godina isisavajući svakog sposobnog inženjera, fizičara i generalno renesansne ljude nakon završetka studija, nudeći im veoma visoke plaće i dodatke na zaradu. Wall Street je postao najveći unajmljivač matematičara, inženjera i znanstvenika i u tome pretekao industriju poluvodiča, farmaciju ili telekomunikacijske tvrtke. Tako da nije začudna činjenica da je jedan dio razloga za čudno divljanje burze iz 2010. godine koji se naziva **munjeviti slom** (engl. Flash Crash) bila i algoritamska trgovina koja je bila programirana da u obzir uzima volumen trgovanja, ne cijenu ili vrijeme; trgovina se izvodila nevjerojatno brzo: u 20 minuta, umjesto u nekoliko sati kako bi bilo tipično za takve naredbe” (Simpson, 2010).

Vojska SAD-a i tajne službe širom svijeta velikim djelom predvode primjenu umjetne inteligencije putem **robota za razminiravanje, sustava za nadzor, dronova za napade** i ostalih neimenovanih vozila. Oni za sada još veoma ovise o navođenju od strane ljudskih operatera, ali provode se istraživanja na proširenju njihovih autonomnih sposobnosti.

Inteligentno raspoređivanje je također veoma uspješno. Sustav rezervacija avionskih tvrtki koristi sofisticirane metode prilikom izrade rasporeda i kreiranja cijena. Poslovni subjekti koriste široki spektar UI tehnologija u kontroli nabave i skladištenja proizvoda. UI tehnologije podloga su mnogih usluga na inetrnetu. Softver upravlja svijetskim prometom e-pošte, sprečava poplavu

spamova. Softver koji koristi UI komponente odgovoran je za automatsko odobrenje ili odbijanje kredita i transakcija te kontinuirano prati aktivnosti korisničkih računa tražeći znakove neovlaštenog korištenja. Sustavi za povrat informacija u širokom spektru koriste strojno učenje, Googleova tražilica je u ovom trenutku najveći UI sustav ikada izgrađen (Bostrom, 2014, str. 32).

5.2. Google i razvoj umjetne inteligencije

Google se posebno ističe kao tvrtka koja od svojeg osnutka radi na realizaciji umjetne inteligencije. Možemo Google smatrati za jedan sustav ili projekt koji povezuje različite inovativne metode, ideje i inovativne ciljeve za razvoj skupa sustava umjetne inteligencije. Google sustav koristi tehnike strojnog učenja da bi automatski klasificirao i obradio veliku količinu informacija s mreže ili za razvoj Google Glass - računala koje se nosi. Google Glass sadrži optički zaslon koji se stavlja na glavu i omogućuje trenutnu i stalnu povezanost sa internetom.

Tvrtka Google ima pristup razvoju umjetne inteligencije koji se odražava u primjeni raspodijeljenih sustava i paralelnog rada računala. Ponekad je to motivirano potrebom da se podatci sakupe s geografski udaljenih mjesta (web stranica sa servera ili podataka o vremenu ili prometu prikupljenih iz senzora), a ponekad potrebom da se obrade ogromne količine podataka koju jednostavno jedna CPU ne može obaviti. Proučavanje raspodijeljenih sustava i paralelnog rada donosi neobične probleme poput kontrole istodobnosti, otpornosti na greške, efikasnosti algoritama i komunikacije. Neki od istraživača u Googleu nastoje odgovoriti na temeljna teoretska pitanja dok su drugi okupirani konstruiranjem sustava koji rade na najvećoj mogućoj skali. Google se od ranih dana morao nositi s tim problemima u primjeni prilikom rada tražilice. U vremenu velikih podataka, ne može se očekivati od nijednog pojedinačnog stroja da se nosi sa količinom informacija i zahtjevom za obradom kako bi ispunio Googleovu misiju da: “organizira svjetske informacije i učini ih univerzalno dostupnima i korisnima” (Google, 2014).

Google ima **graf znanja** (engl. **knowledge graf**), tehnologiju sposobnu analizirati prirodni jezik, tehnologiju neuronskih mreža i nevjerojatnu mogućnost dobivanja povratne informacije od korisnika. “Ukoliko kombiniramo sve to zajedno s ljudima, u petlju koja neprekidno daje povrat,

naš sustav postaje inteligentan” (Clark, 2013). Graf znanja funkcionira dinamički, odnosno prikazuje ono što Google smatra relevantnim na temelju unosa korisnika. Iako Google još nije postigao cilj - stvaranje umjetne inteligencije moguće je vidjeti prve naznake sustava koji ima izgled neovisne inteligencije u alatima za predviđanje koje tvrtka koristi. Primjer je Google Now sustav koji korisnicima dostavlja relevantne podatke u nekom vremenskom trenutku na osnovu postavki prikupljenih od individualnih korisnika.

Četvrtog prosinca.2009. Google je objavio kako prestaje koristiti sustav za indeksiranje stranica kod pretraživanja i uvodi “**personalizirano pretraživanje za svakoga**”. Od tada sustav koristi 57 *signala* prikupljenih od korisnika kako bi prilagodio pretraživanje svakom individualnom korisniku sustava. “Signale sustav koristi kako bi izradio pretpostavku o tome tko je korisnik i koje mrežne stranice bi se korisniku svidjele. Signale čini sve - od mjesta na kojem se korisnik ulogirao, koji browser koristi, što je korisnik prethodno tražio. Čak ako i niste ulogirani, on će prilagoditi svoje rezultate i prikazati vam stranice za koje je predvidio da je najveća vjerojatnost da ćete ih *kliknuti*” (Pariser, 2011, str. 15). Sustav filtrira rezultate pretraživanja u skladu s identitetom korisnika na temelju najnovijih metoda strojnog učenja.

Svaki korisnik Google sustava pomaže stvaranju umjetne inteligencije tako da odabire rezultate koji *najbolje* odgovaraju onome što je korisnik unio u tražilicu (odabirom fotografija objekta ili video sadržaja, na primjer), to je način na koji ovaj ogroman sustav uči i naziva se **pojačano učenje (engl. reinforced learning)**. Pojačano učenje je ono što daje možda najveće šanse Googleu u razvoju umjetne inteligencije. Zato što Google raspolaže najšire korištenom tražilicom, ima stotine milijuna Gmail-ova, YouTube i Android korisnike, kompanija ima kompletnu prednost u podešavanju svojih pristupa umjetnoj inteligenciji kao odgovor ljudima (korisnicima). To je kao da svaki korisnik Googleovih servisa sudjeluje u isprobavanju njihove UI tehnike, ispravljajući pri tome pogreške u traženju i izvođenje akcija kada je traženje uspješno. Njihov suparnik u istraživanju i primjeni novih tehnologija IBM, testira primjenu superračunala Watson na medicinsku dijagnostiku koristeći također tehniku pojačanog učenja u kojoj sustav ima sposobnost brze sinteze velike količine informacija i generiranje hipoteza kao odgovor na pitanja.

Google također radi na automatkoj izgradnji **trezora znanja** (engl. Knowledge Vault), masivne baze podataka koja će dati do sada neviđen pristup činjenicama u svijetu. “Google gradi najveće spremište znanja u ljudskoj povijesti – i to radi bez bilo kakve ljudske pomoći. Umjesto toga trezor znanja automatski sakuplja i spaja informacije širom mreže u jedinstvenu bazu činjenica o svijetu i ljudima i objektima u njemu” (Hudson H., 2014). Trezor znanja je tip **baze znanja** – sustav koji pohranjuje informacije tako da ih mogu istodobno razumjeti strojevi i ljudi. Tamo gdje baze podataka obrađuju brojeve, baza znanja obrađuje činjenice. Postojeća baza, koja se naziva i **graf znanja** (engl. **Knowledge Graph**), oslanja se na masovni izvor (engl. crowdsourcing) za dobavljanje podataka, ideja ili sadržaja od velike skupine ljudi, posebno zajednice na internetu, kako bi sakupila informacije. Kada su u Googlu uočili da se trend rasta grafa smanjio, jer ljudi ipak imaju određena ograničenja, odlučili su automatizirati proces izgradnje trezora korištenjem algoritma koji automatski izvlači informacije svuda iz mreže. U tom procesu se koristi strojno učenje kako bi se *sirove podatke* pretvorilo u primjenjive *djelice znanja*.

Tvrtka Google ulaže veliku količinu njima dostupnih resursa u razvoj umjetne inteligencije i nastoji okupiti i otkupiti manje tvrtke, znanstvenike i stručnjake iz polja umjetne inteligencije kao i vizionare koji bi radili na ostvarenju potpuno autonomnog inteligentnog (i možda svjesnog) sustava globalnog razmjera.

6. PERSPEKTIVE U RAZVOJU I PRIMJENI UMJETNE INTELIGENCIJE

Važno je spomenuti i kako je: “samo u zadnjem dijelu 2014. godine u industriju povezanu sa umjetnom inteligencijom uloženo više od pola milijarde dolara” (Waters, 2015). Golema novčana sredstva trebala bi potaknuti razvoj i primjenu umjetne inteligencije na način koji je sada teško predvidjeti. Velike tvrtke koje ulažu mnogo resursa u razvoj umjetne inteligencije nisu sasvim transparentne u svojim inovacijama, kako bi zadržale financijsku dobit i ostvarile prednost u odnosu na konkurenciju, to je dodatni razlog što je u ovom trenutku teško govoriti koliko daleko smo zapravo stigli u razvoju umjetne inteligencije.

Znanstveni skup održan na Sveučilištu Stanford u ožujku 2013. dobar je prikaz inovativnih projekata i procesa širom svijeta kad je riječ o akademskim istraživanjima umjetne inteligencije.

6.1. Pravci novih istraživanja u umjetnoj inteligenciji

Pravci novih istraživanja se mogu podijeliti u dvije glavne kategorije: oni koji su fokusirani na računalno - računalno interakcije te one koji su usmjereni na interakciju čovjek - računalno. Kada govorimo o interakcijama **računalno – računalno**, u istraživanjima prevladava:

- Izgradnja inteligentnih robota: reintegracija UI
- Cjeloživotno strojno učenje
- Povjerenje i autonomni sustavi
- Slabo nadzirano učenje iz multimedije

Kod sustava temeljenih na interakciji **čovjek - računalno** noviji smjerovi istraživanja su:

- Analiza mikroteksta
- Kreativnost i (rani) razvoj kognicije
- Dobrobit pomoću podataka: od samopraćenja do promjena u ponašanju
- Shikakeologija: dizajn okidača za promjene u ponašanju

Shikake je termin iz japanskog jezika koji predstavlja fizičke ili psihološke okidače za implicitne ili eksplicitne promjene u ponašanju sa ciljem rješavanja problema. “Shikakeologija je interesantna nova kategorija koja je kompletno u trendu sa rastućim pametnim stvarima; Internet stvari (engl. IoT) ili bolje internet svega, kvantificirano ja (pojam koji uključuje obradu podataka sakupljenih samopraćenjem), kreiranje navika i kontinuirano praćenje kretanja osobe. Primjer je kanta za smeće koja dopadljivom muzikom potiče ljude na bacanje otpadaka” (Swan, 2013).

6.2. Istraživanja inspirirana radom mozga – neuromorfna arhitektura

Posljednjih godina, tehnološke tvrtke i akademski istraživači nastoje izraditi takozvanu **neuromorfnu računalnu arhitekturu** – čine je čipovi koji oponašaju sposobnost ljudskog mozga da bude ujedno analitički i intuitivan kako bi omogućili stvaranje konteksta i značenja iz velike količine podataka. “Znanstvenici koji predvode ovakva istraživanja sebe nazivaju neuromorfni inženjeri. Umjesto da se o mozgu razmišlja kao o računalu, oni nastoje izraditi računala koja nalikuju mozgu. Na taj način čovječanstvo će dobiti ne samo bolje razumijevanje rada mozga, nego bolja, pametnija računala” (Hof, 2015).

Dok ljudski mozak ima 100 trilijuna sinapsi i troši svega 20W, današnja superračunala u nastojanju simulacije rada mozga troše snagu reda veličine MW. Vodeća nastojanja da se ostvari sustav čija su svojstva sličnija mozgu dostigla su novu prekretnicu, proizvodnjom tranzistorskog čipa koji sadrži više od 4000 neurosinaptičkih jezgri. Svaka se jezgra sastoji od računalnih komponenti koje odgovaraju njihovm biološkom dvojniku – jezgrene memorijske funkcije slične su sinapsama među neuronima, procesori predstavljaju jezgrine neurone, a komunikacija se ostvaruje vodičima sličnim neuronskim aksonima (Greenemeier, 2014).

Cilj neuromorfnih znanstvenika, je izgradnja računala koje ima neke ili sve značajke koje ima mozak, a današnja računala nemaju, to su **niska potrošnja energije pri radu; tolerancija na kvar** (kvar jednog tranzistora stvara ozbiljne probleme u mikroprocesoru, ali mozak stalno gubi pojedine neurone, a to ne uzrokuje poteškoće u radu živčanog sustava) i **nepotrebnost programiranja** (mozgovi uče i mijenjaju se spontano kroz svoju interakciju sa svijetom, umjesto da slijede zadane putove i grane predodređenog algoritma). Da ostvare cilj, znanstvenici bi morali poznavati rad mozga, međutim to je još uvijek velika nepoznanica. Ne postoji način na koji bi, za sada, mogli proučavati mozak na temeljnoj razini. S druge strane, prikladne računalne simulacije mogle bi odgovoriti na neka pitanja o temeljnim funkcijama mozga i obratno. Za pravo, postizanjem dobrog oponašanja rada mozga dogodila bi se prekretnica u računalstvu koja bi mogla konačno poslužiti boljem poznavanju moždanih funkcija, razvoju umjetne inteligencije i moguće svjesnih računalnih sustava. "Moglo bi se dogoditi da modeli budu prvi pa onda pomognu mapiranju mozga. Neuromorfno inženjerstvo moglo bi, drugim riječima, otkriti temeljne principe mišljenja prije nego to učini neuroznanost" (The Economist, 2013).

Veliki jaz u nerazumijevanju rada mozga nalazi se u srednjoj skali, odnosno u srednjem stupnju anatomije mozga. Znanosti je poznat rad pojedinačnih neurona. Također je relativno dobro poznato kako rade pojedinačne moždane polutke i gangliji (nakupine neurona koje izgrađuju periferni živčani sustav), gdje su u mozgu smješteni centri za govor ili vid. "Međutim, nejasno je kako se neuroni u ganglijima i moždanim polutkama organiziraju, a to je upravo razina organizacije na kojoj se ostvaruje razmišljanje – i pretpostavlja se, nastanjuje svjesnost" (The Economist, 2013).

6.2.1. Europski neuromorfni projekti

U EU se izvodi projekt pod nazivom **HBP** (engl. Human Brain Project) koji navodi cilj na vlastitim službenim stranicama: "U razdoblju od 10 godina, HBP istraživači simuliraju rad mozga, razviju računalne tehnologije inspirirane mozgom, mapiraju bolesti mozga, izvedu ciljano mapiranje mozgovog miša i čovjeka, razviju šest ICT (engl. Information and Communication Technology) platformi, provedu prijenos istraživanja u proizvode i usluge te primjene programe u edukaciji i upravljanju znanjem." Svi ti napori bit će napravljeni u

kontekstu strategije odgovornog istraživanja i inovacija (engl. Responsible Research and Innovation, RRI).

U EU su neuromorfna istraživanja na sveučilištima Manchester i Heildeberg također dala određene rezultate. Temelje se na dva različita neuromorfna pristupa. Prvi prototip čipa isporučen je 2008. godine i naziva se **SpiNNaker**. Čip je izgrađen na Sveučilištu u Manchesteru i osnova mu je digitalno računalo. "Upravo je u tijeku nastavak projekta koji se naziva BIMPA, on se fokusira na konstrukciji većih strojeva, uz SpiNaker zajedno sa softverskom infrastrukturom i aplikacijama. SpiNaker je danas dio HBP projekta. SpiNaker u svojoj jezgri sadrži mrežu mikroprocesora, a brzinu u radu postiže time da nije programirano te umjesto da prebacuje relativno velike blokove podataka pod kontrolom centralnog sata (kako to čine konvencionalna računala), njegovi procesori izbacuju mnogo sitnih šiljaka informacija kada i kako im odgovara. Taj je postupak sličan (i to je namjerno tako) radu pravih neurona. Signali prolaze kroz neurone u obliku električnih šiljaka nazvanih **akcijski potencijal**, koji u sebi sadrže malo informacija osim one da su se dogodili. Takvi asinhroni signali, bez sinhroniziranog sata, mnogo brže obrađuju podatke, također troše manje energije (čime je ispunjen prvi zahtjev), a ako procesor zakaže, sustav će ga zaobići (tako je ispunjen i zahtjev za tolerancijom na kvar). Upravo zbog toga što ih se ne može jednostavno programirati, većina programera izbjegava asinhronne signale, ali kao način za oponašanje rada mozga, oni su savršeni (The Economist, 2013).

Drugi čip nazvan **Spiky** razvijen je na Sveučilištu u Heildebergu i analogan je. Analognim radom računala nastojalo se smanjiti potrošnju snage. Analogni rad čipa zasniva se na reprezentaciji brojeva kao točaka na vremenski kontinuirano promjenjivom naponu određenog raspona. Tako, na primjer 0,5 V ima drugačije značenje nego 1 V i 1,5 V opet ima drugo značenje. Membrane na čipu kojih ima 384 upravljane su različitom provodljivosti kroz ionske kanale u nekom trenutku. Kada je vodljivost jako visoka kapacitivna membrana integrira naboj koji protječe kroz različite ionske kanale. Ako napon na membrani dostigne zadanu vrijednost praga, okida se proces generiranja šiljka. Neuronski parametri su također statistički određeni tako da u modelu, baš kao u prirodi, ne postoje dva jednaka neurona (University of Heidelberg, 2011).

6.2.2. Neuromorfni projekt u SAD-u

U SAD-u je također u tijeku neuromorfni projekt, pod nazivom **SyNAPSE** (engl. Systems of Neuromorphic Adaptive Plastic Scalable Electronics). Financiran je djelom od strane državne agencije za obranu (DARPA) u suradnji sa tvrtkom IBM i Sveučilištem Cornell. Konačni cilj projekta je izgraditi neurosinaptičko superračunalo veličine kutije za cipele s 10 milijardi neurona i 100 triljuna sinapsi koje bi trošilo 1KW snage.

“Organska tehnologija slična mozgu danas ne postoji, a njezin razvoj zahtijevao bi previše vremena i novca. U međuvremenu, pod pritiskom smo poplave podataka koja se ne može odgađati. Tako je naša ključna inovacija nova, **ne**-von Neumanova, paralelna, distribuirana, događajima pokretana, skalabilna arhitektura: ona se može sintetizirati iz današnje tehnologije i istovremeno služiti kao vodilja za buduće tehnologije. Za novu arhitekturu potreban je potpuno novi način razmišljanja, programiranja i učenja. To je **kognitivno računarstvo** – novi spoj silicija i softvera” (Modha, 2013).

Von Neumanova arhitektura odvaja memoriju i obradu podataka pa zato zahtjeva visoku propusnost (engl. bandwidth), što dovodi do povećanja potrošnje snage. Nadalje paralelni i događajima pokretani procesi ne mogu se dobro implementirati na model serijske obrade kod konvencionalnih računala. ”Mi primjenjujemo integraciju mrežaste memorije i neurona kako bi prijenos podataka ostao na lokalnoj razini i koristimo asinhroni događajima pokretani dizajn gdje se svaki krug evaluira paralelno i bez sata, tako se snaga rasipa samo kada je apsolutno nužno. Takav izbor arhitekture omogućuje gusto integrirane sinapse uz ultra nisku potrošnju snage i garantirano izvođenje u realnom vremenu” (Merrola, 2014).



Slika 6.1. IBM-ova ploča sa 16 neuromorfih čipova

IBM i Cornell su čip pod nazivom **TrueNorth** izgradili tako da se ponaša kao energetski efikasna šiljkasta neuronska mreža. Za razliku od normalnih neuronskih mreža koje signale obrađuju u pravilnim intervalima, ove mreže ispaljuju signale jedino kada električni signal dostigne specifičnu vrijednost. Ispaljivanjem signala utječu na nabijanje ostalih umjetnih neurona - slično radu mozga.

Istraživači su testirali čip na analizi video isječka, gdje je trebalo razabrati ljude, bicikliste i ostale objekte od pozadine te identificirati svaki objekt. U usporedbi sa simulatorom koji radi na istom tipu neuronske mreže (u kojoj se koristi moderan mikroprocesor za općenitu namjenu), čip nove konfiguracije trošio je 176 000 puta manje energije dok je slike obrađivao 100 puta brže.

IBM zamišlja neurosinaptički čip kao građevni blok za novu vrstu superračunala koji oponašaju ljudski mozak korištenjem milimetarskih mikročipova koji bi trošili malo energije i moglo bi ih se ugraditi u naočale, satove i ostale nošljive predmete. Kako su ovi čipovi jako dobri u sortiranju senzornih ulaznih signala, to ih čini sjajnim kandidatima u medicinskoj dijagnostici, a u IBM-u se nadaju budućoj integraciji SyNAPSE čipa i Watson superračunala. “Ova arhitektura u stanju je riješiti širok raspon problema, poput vida, sluha, u stanju je efikasno procesirati šumovite senzorne podatke u realnom vremenu, uz potrošnju snage nekoliko puta manje od konvencionalnih računalnih arhitektura” (Modha, 2013).

6.3. Nova promišljanja umjetne inteligencije – projekt na MIT-u

Mnogi znanstvenici smatraju kako se u polju istraživanja UI, utemeljenog prije 50 godina, puno vremena provelo lutajući u divljini, mijenjajući visoko ambiciozne ciljeve za relativno skromne skupine stvarnih dostignuća. Neki od pionira ovoga područja sada na MIT-u, u suradnji s novijim generacijama mislioca, udružuju snage za masivne preinake cijele ideje. Ovoga puta odlučni su napraviti to kako treba – i s prednošću sadašnjeg znanja, iskustva i brzog rasta novih tehnologija te uvida iz znanosti neuralnog računalstva, vjeruju da imaju za to dobre šanse. Projekt šireg razmjera pokrenut je 2009. godine. Cilj projekta nazvanog **MMP** (engl. Mind Machine Project) je ponovo osmisliti umjetnu inteligenciju, za novo doba. Namjera istraživača uključenih u projekt je da se vrate unazad i revidiraju temeljne pretpostavke o prirodi mozga, spoznaje, računanja i inteligencije te izrada inteligentnih strojeva. Neil Gershenfeld jedan od voditelja projekta, kaže: “projekt revidira temeljna načela u svim područjima obuhvaćenim umjetnom inteligencijom, uključujući prirodu uma i pamćenja te načina na koji se inteligencija manifestira u fizičkom obliku” (Chandler, 2009).

U realizaciji projekta nastoji se zahvatiti tri područja u kojima su istraživanja umjetne inteligencije “zapela”, a to su: um, memorija i tijelo.

- 1) **Um - izrada modela za misao.** Ljudski je um izuzetno kompleksan. Kroz evoluciju, um se razvio tako da je sastavljen od mnogo dijelova koji možda i nisu usuglašeni, slično kao i različita područja UI. Postoje dijelovi koji su dobri u rješavanju određenih problema, ali se onda nastoji njima riješiti sve. Zato je potrebno da se naprave sustavi izgrađeni od mnogo dijelova koji rade zajedno, poput različitih elemenata uma. Smatraju da nedostaje ekologija modela. Sustav koji može rješavati problem na mnogo načina, kako to um čini.
- 2) **Memorija – sakupljanje i korištenje iskustva.** Mnogo je pokušaja da se nametne umjetna konzistencija sustava i pravila za zbrkanu i kompleksnu prirodu ljudskih misli i pamćenja. Zato je, kaže Gershenfeld, sada moguće akumulirati cijelo životno iskustvo osobe i tada prosuđivati korištenjem skupa podataka koji su puni višeznačnosti i nepostojanosti. Tako mi funkcioniramo – ne prosuđujemo na temelju preciznih istina.

Računala moraju naučiti načine prosuđivanja koji rade *sa*, a ne izbjegavaju, višeznačnost i nepostojanost (Chandler, 2009).

- 3) **Tijelo – mjerljivi substrati (podloge) za utjelovljenje inteligencije.** Računala se programiraju pisanjem sekvenci linija koda, ali um tako ne radi. U umu se sve događa svuda, cijelo vrijeme. Novi pristup programiranju nazvan RALA (engl. Reconfigurable Asynchronous Logic Automata) nastoji ponovo primjeniti računalstvo na temelju fizikalnih načela, da sadržava fizikalne jedinice vremena i prostora pa bi se opis sustava poravnao sa sustavom kojeg predstavlja. “To bi moglo dovesti do računala koja pokreće fini paralelizam koji koristi mozak” (Chandler, 2009).

Grupa čak predlaže odbacivanje Turingovog testa, koji je već dugo priznati standard za određivanje umjetne inteligencije. Umjesto toga, istraživači MMP-a žele ispitati strojno shvaćanje dječje knjige - radije nego ljudsko shvaćanje drugog ljudskog bića – kako bi bolje razumjeli sposobnost umjetne inteligencije da obrađuje i vraća misli.

6.4. Predviđanja o primjeni umjetne inteligencije

Ono o čemu se sve više govori je **internet stvari** ili **IoT** (engl. Internet of Things). Ostvaruje se kada uzmemo objekte iz svakodnevnog života i dodamo im mogućnost da *sakupljaju podatke* (o našim korisničkim navikama), *komuniciraju međusobno* i s nama. Tim terminom nisu uključena osobna računala, pametni telefoni i tableti, govori se dakle samo o svakodnevnim stvarima poput kućanskih aparata. Predviđanja govore kako će rast IoT uređaja premašiti ostale uređaje za povezivanje. “Do 2020, broj pametnih telefona, tableta i osobnih računala će doseći 7,3 milijarde jedinica. Nasuprot tome, IoT će se širiti mnogo brže, što će rezultirati brojem od 26 do 50 milijardi jedinica, u istom tom trenutku” (Plummer, 2014). Povezanost stvari, poplava informacija sakupljenih od senzora, koji će biti sve jeftiniji i manji zahtijevati će inteligentne sustave koji će upravljati uređajima na korisnicima neprimjetan način, u nekim budućim pametnim gradovima. Kao rješenje za svijet preplavljen informacijama nudi se personalizacija. **Personalizacija** se provodi primjenom inteligentnih sustava na podatke sakupljene od korisnika. Podatci su sakupljeni preko interneta na društvenim mrežama ili putem onoga što korisnik unosi

u tražilicu ili iz različitih senzora, sustava za praćenje ili korištenjem kreditnih kartica. Rezultati pretraživanja prilagođeni najprije osobnim interesima korisnika, skrojeni prema navikama i preferencama korisnika, vijesti prilagođene interesima korisnika, filtriranje “bitnih” tema i prilagođavanje ponude oglašivača navikama potrošača. Sve to temelji se na mogućnosti pametnih sustava da obrade i interpretiraju sakupljene podatke. Peronalizacija oglasa koju provode oglašivači ponudu zaključuje na temelju pametnog softvera o našim namjerama i potrebama preko sakupljanja uzoraka sa stranica na kojima kupujemo ili na temelju naših objava na društvenim medijima.

Kako je sakupljanje i korištenje podataka automatizirano dobro je stoga postaviti pitanje o **automatiziranoj zaštiti privatnosti**, odnosno o kreiranju osobnog agenta, softvera koji će umjesto korisnika voditi računa o njegovim interesima, ali i zaštititi korisnika od neželjenih sadržaja ili zlorabe osobnih podataka, poput nekog posrednika između čovjeka i interneta. Otvaraju se tako brojne mogućnosti za komercijalnu primjenu softverskih agenata u zaštiti pojedinaca na internetu. “Kada se radi o zaštiti privatnosti, većini ljudi to se čini dosadnim, teškim i nepotrebnim, jer nadaju se da će u anonimnosti mnoštva, njihovim vlastitim dosadnim životima i nadom da se njima ništa neće dogoditi imati dovoljnu zaštitu. Ali oni se varaju. Na koji način možemo donijeti prividni red i kontrolu nazad pojedincu suočenom s rastućom kompleksnosti i transparentnosti u svijetu? Sljedeći korak u kontroli privatnosti mogao bi biti korisnički agent koji živi u oblaku i koji autonomno djeluje u cilju zaštite digitalnog života i vlasništva“ (Blonder, 2014). Pitanje koje bi se moglo postaviti u skoroj budućnosti je i zaštita **privatnosti neuronskih podataka**. To je zamisao koja se bavi pitanjima privatnosti i sigurnosti obzirom na tokove podataka dobivenih iz mozga osobe. Nekoliko razloga mogli bi biti značajni za brigu. Prvo, personalizirani zdravstveni podaci već su sada predmet rasprava kada se radi o pitanjima osobnih podataka i sve što se tiče misli i mentalnog djelovanja kao i potencijalne patologije ima veću osjetljivost i s time povezan tabu. Drugo, pitanja o zaštiti privatnosti neuronskih podataka može postati problem iz razloga što nije teško izmjeriti neku razinu električnih ili nekih drugih moždanih aktivnosti. Razvoj novih tehnologija može omogućiti hvatanje i obradu neuronskih aktivnosti velikog broja ljudi u realnom vremenu. Već postoje brojni primjeri uređaja koji mjere neuronske aktivnosti poput EEG uređaja, Google Glassa kao i brojnih aplikacija za analizu fizičkog i kognitivnog stanja. Treće, u nekom trenutku, strojevi koji

raspolazu velikim podacima i strojnim učenjem mogu biti u stanju uspostaviti valjanost i korisnost neuralnih podataka u odnosu na brojna stanja ljudske zdravstvene, fizičke i mentalne izvedbe. Četvrto, unatoč osjetljivosti toka neuronskih podataka, kao i svaki drugi oblik osobnih podataka (gdje dva elementa podatka počinju tvoriti identifikaciju), privatnost, sigurnost i anonimnost mogu biti praktično nemogući. “U najgorem slučaju, može doći do pojave zlonamjernog hakiranja, virusa i spamova koji ciljaju na tijek neuronskih podataka” (Swan, 2014).

Na kraju treba reći da iako danas prevladava primjena umjetne inteligencije na rješavanje određenih, vrlo specifičnih zadataka, postoje i realna očekivanja kako će generalna umjetna inteligencija zaživjeti negdje sredinom ovog stoljeća. Zato mnogi znanstvenici iskazuju zabrinutost kako će generalna umjetna inteligencija, koju bi mogli postići u skoroj budućnosti, utjecati na čovječanstvo. Možemo se zapitati: što bi stvaranje generalne umjetne inteligencije značilo za našu civilizaciju, društveni poredak i evoluciju ljudske vrste ili na cjelokupan život na Zemlji koji poznajemo? Također, što bi mogla donjeti integracija ljudi i strojeva u budućnosti? Nemamo ništa slično za usporedbu, tako da možemo samo nagađati o posljedicama stvaranja generalne umjetne inteligencije. *Generalna umjetna inteligencija je opasna!* Ovo je jedna od glavnih primjedbi koja je stara koliko i područje umjetne inteligencije. Poput svake druge znanosti i tehnologije, GUI posjeduje mogućnost da se zloupotrebi, ali to nije razlog da se zaustavi istraživanje GUI, kao i u bilo kojoj drugoj znanstvenoj disciplini. Mišljenje da je generalna umjetna inteligencija fundamentalno opasna zbog toga što će dovesti do katastrofe je obično zasnovano na različitim pretpostavkama da će inteligentni sustav u konačnici htjeti dominirati svemirom. Nick Bostrom, voditelj Instituta za Budućnost čovječanstva, kaže u knjizi *Superinteligencija*: “Ako jednoga dana izgradimo strojni mozak koji nadilazi ljudski mozak u generalnoj inteligenciji, ta nova superinteligencija mogla bi postati veoma moćna. I kao što sudbina gorila danas ovisi više o ljudima nego o njima samima, tako bi sudbina naše vrste mogla ovisiti o akcijama superinteligencije. Uopće nije upitno da se nalazimo na pragu velikog prodora u umjetnoj inteligenciji ili da možemo predvidjeti sa bilo kakvom preciznošću kada bi se takav razvoj mogao pojaviti. Čini se da bi to moglo biti u neko doba u ovom stoljeću, ali ne znamo zasigurno” (Bostrom, 2014, str. 11).

7. ZAKLJUČAK

Umjetna inteligencija je u slabom obliku danas veoma prisutna, u različitim tehnologijama i područjima ljudskih aktivnosti. Spoznajom da umjetna inteligencija ne mora biti zaista inteligentna u generalnom smislu, već dovoljno pametna za rješavanje problema, napravljen je veliki pomak od akademskih istraživanja do primjene tehnika razvijenih u umjetnoj inteligenciji pa se danas mogu naći zaista brojni primjeri uporabe inteligentnih sustava.

Može se pratiti komercijalna primjena umjetne inteligencije od ekspertnih sustava i strojnog učenja u kojima su napravljeni prvi koraci, do današnjeg razvoja umjetnih neuronskih mreža velikih razmjera koje omogućuju računalima da uče, prepoznaju objekte iz okoline, koriste prirodni jezik i sl. Novije metode umjetne inteligencije nastoje izraditi inteligentne agente koji su kreativni, koji uče u interakciji sa okolinom. Evolucijsko računalstvo koristi načela evolucije i prirodne selekcije pri traženju rješenja pomoću genetskih algoritama. Kreirani inteligentni agenti u interakciji s okolinom mijenjaju svoje ponašanje i prilagođavaju se promjenama bez ljudskog vođenja. Postoje zaista brojni primjeri u kojima se primjenjuju inteligentni sustavi u današnjem povezanom društvu, gdje količina podataka koju je potrebno klasificirati, obraditi i vratiti korisnicima postaje golema, a na predviđanjima i preporukama koje izvodi pametni softver zarađuju se ili gube značajne količine novaca. Za projekte istraživanja ljudskog mozga i neuromorfne računalne arhitekture očekuje se da će donijeti novi prodor u izgradnji umjetne inteligencije kroz idućih nekoliko godina. Umjetna inteligencija je ipak relativno mlada znanstvena disciplina čiji je razvoj posljednjih godina uzrokovan jednako porastom ulaganja znatnih novčanih sredstava u istraživanja kao i razvojem i širenjem Interneta. Taj nagli razvoj pokrenuo je također niz pitanja o posljedicama primjene umjetne inteligencije; na privatnost, sigurnost pojedinca i ekonomiju. Pogotovo se ističe strah od izgradnje jake, možda svjesne umjetne inteligencije kao nečega s čime se čovječanstvo još nije susrelo, a posljedice čega nismo u stanju predvidjeti.

8. METODIČKA OBRADA NASTAVNE JEDINICE

Priprema za izvođenje nastavnog sata informatike

Nastavni predmet: Informatika

Nastavna cjelina: Inteligentni sustavi

Nastavna tema: Inteligentni agenti

Nastavna jedinica: Inteligentni agenti

Tip nastavnog sata: obrada novog gradiva

Vrijeme izvođenja: 1 sat

Mjesto izvođenja: računalni kabinet

Cilj (svrha) sata: upoznati učenike s pojmom inteligentnog agenta te navesti primjere uporabe inteligentnih agenata

Ishodi učenja:

- definirati pojam agenta
- definirati pojam racionalnog agenta
- definirati pojam inteligentnog agenta
- navesti osnovnu strukturu agenta
- objasniti interakciju agenta s okolinom

- objasniti načelo uspješnosti agenta i mjeru izvedbe
- navesti primjere korištenja inteligentnih agenata u rješavanju problema
- definirati višeagentske sustave
- navesti primjer uporabe višeagentskih sustava
- objasniti ulogu višeagentskih sustava u povezanom svijetu

Nastavna načela:

- **načelo postupnosti** (gradivo se izlaže od jednostavnijeg ka složenom. Tijekom izvođenja sata učenici se upoznaju sa konceptom inteligentnog agenta, na kraju se ponavlja naučeno).
- **načelo zornosti** (uz pomoć projektora prikazati shemu interakcije agenta s okolinom).
- **načelo primjerenosti** (primjeri uvažavaju dob i predznanje učenika).
- **načelo individualizacije** (za vrijeme izlaganja prati se sudjelovanje učenika te se pruža dodatno objašnjenje, ako je potrebno).

Nastavne metode:

- metoda usmenog izlaganja,
- metoda demonstracije,
- metoda razgovora,
- metoda pisanja i rada s tekstom.

ORGANIZACIJA NASTAVNOG SATA

ETAPA	SADRŽAJ	OBLICI RADA	METODE RADA	VRIJEME
Uvodni dio	– priprema za rad – najava teme sata – motivacija učenika	Frontalni	Usmeno izlaganje, metoda razgovora	2
Središnji dio	– obrada gradiva – demonstracija primjera	Frontalni, individualni rad, frontalni rad s elementima individualizacije	Usmeno izlaganje, metoda demonstracije, razgovor	30
Završni dio	– ponavljanje i utvrđivanje gradiva	Frontalni rad	Metoda razgovora	5

TIJEK NASTAVNOG SATA

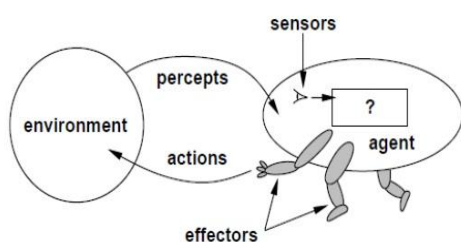
Uvodni dio

Najava teme koja se obrađuje na satu na način da razgovorom sa učenicima potaknemo njihovu motivaciju za što boljim razumjevanjem inteligentnih agenata.

Ono što su već naučili u sklopu prethodnih tema, na primjer korištenje umjetnih neuronskih mreža kao primjer bottom-up tehnike razvoja inteligentnih sustava.

Središnji dio

Definirati pojam agenta. Agent je nešto što može percipirati svoju okolinu putem **senzora** (osjetila) i djelovati na tu okolinu putem **efektora**. **Ljudski agent** ima oči, uši i ostale organe za senzore te ruke, noge, usta i ostale dijelove tijela za efektore. **Robotski agent** zamjenjuje kamere i infracrvene lokatore za senzore, a različite motore za efektore. **Softverski agent** ima kodirane nizove bitova za perceptore i akcije koje izvodi. Prikazati i objasniti interakciju agenta s okolinom.



Definirati pojam racionalnog agenta. Osnovna karakteristika **racionalnog agenta** bila bi da čini ispravnu stvar. Može se reći da su ispravne akcije one u kojima će agent biti najučinkovitiji. Dakle, to zahtjeva odluku o tome kako i na koji način provesti evaluaciju agentovog uspjeha.

U svrhu evaluacije uspjeha koristi se izraz mjera izvedbe (engl. performance measure) kod određivanja kriterija uspješnosti agenta. Pri tome se inzistira na objektivnom mjerenju, nametnutom od strane nekog autoriteta. Drugim rječima, mi kao vanjski promatrači postavljamo standard o tome što znači biti uspješan u okolini i koristimo ga za mjerenje izvedbe agenata.

Kao primjer, navodi se slučaj agenta koji bi trebao očistiti prljavi pod. Prihvatljiva mjera izvedbe bila bi količina prljavštine koju je agent očistio u osam sati. Pitamo učenike o tome koje bi faktore mogli uključiti u mjeru izvedbe rada agenta. Očekujemo kako će učenici navesti mjere poput potrošnje energije. Sofisticiranije mjerenje uključivalo bi dakle i faktor potrošnje električne energije kao i faktor stvaranja količine buke u radu agenta. Također je važan i vremenski faktor, odnosno ukoliko želimo mjeriti količinu prljavštine koju je agent sakupio u

prvom satu, nagradili bi agenta koji je u početku najučinkovitiji (čak kada bi kasnije radio malo ili nimalo), a kažnjavali bi one koji rade konzistentno.

Nakon toga objasnimo učenicima osnovnu funkciju agenata u inteligentnim sustavima. Objasnimo kako je zadaća tehnika umjetne inteligencije kreiranje **agent programa**: funkcije koja implementira mapiranje agenta od percepta do akcija.

Pretpostaka je da će se program izvoditi na nekoj vrsti računalne naprave koju nazivamo **arhitektura**. Program se odabire tako da ga arhitektura prihvata i izvodi. Zatim pitamo učenike što sve uključuje pojam arhitektura u informatici? Očekujemo da će odgovoriti kako je to hardver. Objasnimo učenicima da arhitektura može biti jednostavno računalo ili može uključivati hardver za specijalne namjere poput obrade slika iz kamere ili filtracije audio unosa. Može uključivati softver koji omogućava izolaciju između temeljnog računala i agent programa, kako bi se moglo programirati na višoj razini.

Općenito; arhitektura omogućava perceptore iz senzora dostupnima program, izvodi program, i dobavlja programske akcije efektorima kako su one generirane. Odnos između agenata, arhitekture i programa može se sažeti ovako:

AGENT = ARHITEKTURA + PROGRAM

Definiramo pojam **autonomnog inteligentnog agenta** te navedemo njegova svojstva. Objasnimo učenicima da je autonomnost sustava određena je mjerom kojom je njegovo ponašanje određeno vlastitim iskustvom. Zaista autonoman inteligentni agent trebao bi samostalno djelovati uspješno u različitim vrstama okruženja, ukoliko ima dovoljno vremena za prilagodbu. Autonomni agent je sustav koji:

- je smješten u okolinu,
- je dio te okoline,
- koju osjeća,
- na koju djeluje,
- u vremenu,
- s ciljem ostvarenja vlastitih planova (nijedan čovjek ne navodi njegove odabire u akcijama),

- kako bi te akcije mogle utjecati na njegova buduća opažanja (on je strukturno udružen sa svojom okolinom).

Posljednja osobina koja je navedena, razlikuje autonomne agente od ostalog softvera. Pretpostavka je da bi neki generalno inteligentni sustav trebao biti autonomni agent iz razloga što je za generalizaciju znanja, između različitih i vjerojatno novih područja, potrebno učenje. Učenje zahtjeva osjetila, a često i izvođenje akcije.

Autonomni je agent pogodan za učenje, pogotovo učenje nalik ljudskom. Da bi sve to izvodio, agent mora imati ugrađene senzore za osjetila, efektore za izvršenje akcije i mora imati primitivne motivatore, koji motiviraju njegove akcije. Senzori, efektori i motivatori jesu primitivi koji moraju biti ugrađeni u agenta ili se moraju razviti u bilo kojem agentu.

Nakon objašnjenja funkcija i strukture, navode se primjeri korištenja inteligentnih agenata. Računalni agenti trenutno su veoma aktualna tema i prijelomna točka u razvoju softvera nove generacije. Postoje široki spektar mogućih softverskih agenata koji se još nazivaju i **softbot**. Objasnimo učenicima kako se na primjer, softverski agenti danas koriste kod promatranja financijskih tržišta.

Oni u realnom vremenu provjeravaju kretanja dionica i njihove cijene na tržištima. Agent može biti i često jest odgovoran za kupovinu i prodaju roba pa u tom slučaju mora znati ako (čovjek) dobavljač ima u nekom trenutku smjernice koje dionice su zanimljive, a koje treba izbjegavati – u tom slučaju softverski agent može trebati “razumijevanje” instrukcija govornog jezika. Agenti su idealni za ovakve transakcije zato što jednostavno miruju i promatraju aktivnosti, a radnju izvode samo kada se pojave “pravi uvjeti”. Ne samo što je ovo teško izvodivo za ljude, već jednom kada je odluka potrebna, agent je može izvesti gotovo smjesta. U vremenu koje bi ljudskom brokeru bilo potrebno da napravi istu odluku (nekoliko sekundi ili minuta) prodaja može već biti gotova. Rezultat toga je da veoma veliki omjer dnevnih financijskih transakcija širom svijeta izvode, ne ljudi, već softverski agenti.

Jedan pristup inteligentnih sustava se posebno fokusira na ideju agenata, i njihove individualne identitete kako bi se proizvelo ukupno nastalo ponašanje. Na taj način se razvijaju i koriste **višeagentski sustavi**. Svaki element može se smatrati dijelom društva koje obično percipira

ograničene aspekte svoje okoline, a na koju može djelovati ili samostalno ili u suradnji s ostalim agentima. Tako pojedinačni agenti usklađuju svoje ponašanje s ostalima kako bi obavili specifičnu zadaću.

U korištenju agenata za rješavanje problema može se kompleksni problem razlomiti na manje probleme od kojih je svaki mnogo lakši za rješavanje. Agenti se tada mogu koristiti za rješavanje tih manjih problema – udružujući se u ostvarivanju konačnog rješenja. Jedna prednost ovoga je u tome da svaki agent sadrži informaciju o svom “manjem” problemu – ne mora znati ništa o generalnom problemu.

Opseg, domet i kompleksnost primjene novih informacijskih sustava povećani su danas na visoki stupanj, a raspodijeljeni inteligentni agenti u njima igraju ključnu ulogu. Osobitosti takvih sustava jesu:

- Podatci, znanje i kontrola raspodijeljeni su ne samo logički, već fizički.
- Svaki sudionik u rješavanju problema povezan je računalnom mrežom.
- Svaka komponenta surađuje s ostalima kako bi se riješio problem koji jedna komponenta samostalno ne može.

Tamo gdje se primjenjuju višeagentski sustavi, oni mogu djelovati u suradničkom odnosu tako da svaki od agenata dobavi dio odgovora na problem, a sveukupno rješenje dobije se skupnom kohezijom izlaza sa brojnih agenata. Također je moguće postaviti agente da djeluju natjecateljski, pojedinačno ili u grupi pa samo mala grupa aktivnih agenata dobavlja sveobuhvatno konačno rješenje.

Navedemo kao primjer korištenja raspodijeljenih višeagentskih sustava (u vidu senzora) projekt mreže raspodijeljenih peer-to-peer teleskopa, koji su djelujući autonomno, rasporedili opažanja tako da reagiraju na vremenski kritične događaje i u to doba (2009. godine) polučili golemi uspjeh u vidu otkrića tada najudaljenijih objekata ikada pronađenih.

Obajsniti učenicima mogućnosti primjene inteligentnih agenata u budućim sustavima te navesti kako se smatra se da će se geografski raspodijeljeni agenti koji preko senzora prikupljaju podatke iz svoje okoline, primjenjivati sve više na svakodnevne objekte. Neizbježno, količina računala koji će umjesto nas (u drugom planu) izvršavati zadaće prikupljanja i obrade podataka u realnom vremenu povećavat će se, a za desetak godina svaki dio odjeće, nakita i svega što nosimo uz

sebe, mjeriti će, vagati i računati, a svijet će biti pun senzora. Prikupljeni podaci biti će u oblaku, u halou uređaja čiji će zadatak biti da nam pružaju usluge opažanja i konstantnog računanja, dok šetamo oni će izračunavati, međusobno se konzultirati, predviđati i iščekivati naše potrebe. Bit ćemo okruženi mrežom distribuiranih senzora i računala.

Završni dio

Nakon izlaganja ukratko ponoviti današnje gradivo i ključne pojmove. Zahvaliti se učenicima na pažnji!

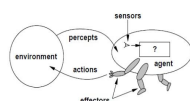
NASTAVNA SREDSTVA I POMAGALA

- literatura za pripremanje nastavnika: Russell S. i Norvig P. Artificial intelligence: a modern approach, Prentice-Hall Inc., New Jersey, 1995.
- obavezna literatura učenika: Janja Linardić, Darka Sudarević, Davor Šokac: www.informatika, udžbenik informatike s cd-om za gimnazije.
- školska ploča, marker
- projektor i računalo (koristi nastavnik kako bi prezentirao gradivo)
- računalo (koriste učenici kako bi pratili nastavu)

PLANIRANJE RADA NA ŠKOLSKOJ PLOČI

INTELIGENTNI AGENTI

Agent je nešto što može percipirati svoju okolinu putem **senzora** (osjetila) i djelovati na tu okolinu putem **efektora**.



Interakcija agenta s okolinom:

Odnos između agenata, arhitekture i programa: **AGENT = ARHITEKTURA + PROGRAM**

Autonoman inteligentni agent trebao bi samostalno djelovati uspješno u različitim vrstama okruženja.

Višeagentski sustavi pretežno istražuju inteligentno ponašanje prilikom koordinacije više agenata.

LITERATURA

1. Allan A., The inevitability of smart dust, Why general purpose computing will diffuse into our environment, dostupno 1.5.2015. na: <http://radar.oreilly.com/2013/01/the-inevitability-of-smart-dust.html#more-54191>
2. Bell J., Machine Learning: Hands-On for Developers and Technical Professionals, John Wiley & Sons Inc, 2014.
3. Blonder G., A Guardian "Agent" to Protect You from Digital Fraud (2014), dostupno, 24.02.2015.,na:<http://blogs.scientificamerican.com/guest-blog/2014/09/26/a-guardian-agent-to-protect-you-from-digital-fraud/>
4. Bostrom N., Superintelligence, Oxford University Press, 2014.
5. Chandler D.L., Rethinking artificial intelligence MIT News Office, 7.12.2009., dostupno 26.4.2012, na: <http://newsoffice.mit.edu/2009/ai-overview-1207>
6. Clark J., Google research chief: 'Emergent artificial intelligence? Hogwash!',2013, dostupno 26.4.2015. na: http://www.theregister.co.uk/2013/05/17/google_ai_hogwash/
7. Copeland B., Artificial intelligence (AI), dostupno 24.2.2015., na: <http://www.britannica.com/EBchecked/topic/37146/artificial-intelligence-AI>
8. Dalbelo Bašić B., Čupić M., Šnajder J., Umjetne neuronske mreže, Fakultet elektrotehnike i računarstva, Zagreb, svibanj 2008.
9. Deutsch D., The Beginning of Infinity, Penguin Books Ltd, 2011.

10. Dormehl L., Why Google Is Investing In Deep Learning, 2014, dostupno 26.4.2015. na:
<http://www.fastcolabs.com/3026423/why-google-is-investing-in-deep-learning>
11. Dunagan J., Mind in a Designed World, Institute for the Future, 2010., dostupno 24.02.2015.na:<http://www.iftf.org/uploads/media/SR-1345%20Mind%20in%20a%20Designed%20World%20%281%29.pdf>
12. Fox A.E., Fan W., The Effects of Fitness Functions on Genetic Programming - Based Ranking Discovery For Web Search, dostupno 23.3.2015. na:
<http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.3.1463&rep=rep1&type=pdf>
13. Goertzel B. i Wang P. (Eds.), Advances in Artificial General Intelligence: Concepts, Architectures and Algorithms, IOS Press, 2007.
14. Google, Distributed Systems and Parallel Computing, 2014, dostupno 26.4.2015. na:
<http://research.google.com/pubs/DistributedSystemsandParallelComputing.html>
15. Greenemeier L. Brain-Inspired Computing Reaches a New Milestone, 2014. dostupno 23.3.2015. na:<http://blogs.scientificamerican.com/observations/2014/08/07/brain-inspired-computing-reaches-a-new-milestone/>
16. Hof R., Deep learning, dostupno 9.3.1015.
na:<http://www.technologyreview.com/featuredstory/513696/deep-learning/>
17. Hof R. D., Neuromorphic Chips, dostupno 25.3.2015. na:
<http://www.technologyreview.com/featuredstory/526506/neuromorphic-chips/>
18. Hudson H., Google's fact-checking bots build vast knowledge bank, 2014, dostupno 9.5.2015. na: <http://www.newscientist.com/article/mg22329832.700-googles-factchecking-bots-build-vast-knowledge-bank.html#.VU35S44vnLU>
19. Kurzweil R., How to Create a Mind: The Secret of Human Thought Revealed, Viking Penguin, 2012.
20. Mayer-Schönberger V., Cukier K., Big data: A Revolution That Will Transform How We Live, Work and Think, Houghton Mifflin Harcourt Publishing Company, 2013.

21. Merrola P., Artur J., A Digital Neurosynaptic Core Using Embedded Crossbar Memory with 45pJ per Spike in 45nm, 2014, dostupno 7.5.2015. na: <http://www.modha.org/papers/012.CICC1.pdf>
22. Mitchell M., An introduction to genetic algorithms, A Bradford Book, The MIT Press, 1999.
23. Modha D., Introducing a Brain-inspired Computer, 2014, dostupno 5.4.2015. na: <http://www.research.ibm.com/articles/brain-chip.shtml>
24. Modha D., Systems That Perceive, Think, and Act, 2013, dostupno 5.4.1015. na: <http://www.theatlantic.com/sponsored/ibmcognitivecomputing/archive/2013/06/systems-that-perceive-think-and-act/276708/>
25. Pariser E., The Filter Bubble, The Penguin Press, 2011.
26. Plummer C., Top 10 Strategic Predictions for 2015 and Beyond: Digital Business Is Driving 'Big Change' (2014), dostupno na: <https://www.gartner.com/doc/2864817?srcId=1-3132930041#a-152809420>
27. Russell S. i Norvig P. Artificial intelligence: a modern approach, Prentice-Hall Inc., New Jersey, 1995.
28. S.Jagan, Dr.S.P.Rajagopalan, A Survey on Web Personalization of Web Usage Mining, International Research Journal of Engineering and Technology (IRJET), 2015.
29. Schank R., Partridge D., Wilks Y., (Eds.), The foundations of artificial intelligence, Cambridge University Press, 1993.
30. Senior D., Narrow AI: Automating The Future Of Information Retrieval,2015, dostupno 5.5.1015. na: <http://techcrunch.com/2015/01/31/narrow-ai-cant-do-that-or-can-it/>
31. Shasha D., Lazere C., Natural computing: DNA, quantum bits, and the future of smart machines, W. W. Norton & Company, Inc. 2010.

32. Shet V., Are you a robot? Introducing No CAPTCHA reCAPTCHA, dostupno 3.12.2014. na: <http://googleonlinesecurity.blogspot.com.es/2014/12/are-you-robot-introducing-no-captcha.html>
33. Shi Z., Advanced Artificial Intelligence, World Scientific Publishing Co. 2011.
34. Simpson S., What Caused The Flash Crash?, 2010, dostupno 26.4.2015. na: <http://www.forbes.com/2010/06/30/what-caused-flash-crash-personal-finance-panic.html>
35. Steiner C., Automate this: how algorithms came to rule our world, Penguin Group (USA) Inc., 2012.
36. Swan M., What's new in AI? Trust, Creativity, and Shikake, dostupno 4.3.2015. na: <http://futurememes.blogspot.it/2013/03/the-aaai-spring-symposia-held-at.html>
37. Swan M., Neural Data Privacy Rights, dostupno 26.4.2015. na: <http://futurememes.blogspot.com/2014/06/neural-data-privacy-rights.html>
38. The Economist, Neuromorphic computing, The machine of a new soul, 2013, dostupno 5.4.2015. na: <http://www.economist.com/news/science-and-technology/21582495-computers-will-help-people-understand-brains-better-and-understanding-brains>
39. University of Heidelberg , 2008, dostupno 7.5.2015. na: http://www.kip.uni-heidelberg.de/cms/vision/projects/facets/neuromorphic_hardware/single_chip_system/the_spikey_chip/
40. Warwick K., Artificial intelligence: the basics, Routledge, 2012.
41. Waters R., Investor rush to artificial intelligence is real deal, 2015, dostupno 9.5.2015. na: <http://www.ft.com/intl/cms/s/2/019b3702-92a2-11e4-a1fd-00144feabdc0.html#axzz3Ny5kj89q>
42. Wikipedia, Teorija spoznaje, dostupno 26.4.2015 na: http://hr.wikipedia.org/wiki/Teorija_spoznaje

POPIS SLIKA

Slika 2.1. Model monarha, (Shasha, Lazare, 2010, str. 28)

Slika 2.2. Nouvelle AI pristup (Shasha, Lazare, 2010, str. 32)

Slika 2.3. Primjer CAPTCHA testa, preuzeto sa: <http://captcha.com/>

Slika 2.4. Primjer No CAPTCHA reCAPTCHA testa, preuzeto sa:

<https://www.google.com/recaptcha/intro/index.html>

Slika 2.5. Primjer istodobnog CAPTCHA i No CAPTCHA reCAPTCHA testa, preuzeto sa:

<https://www.google.com/recaptcha/intro/index.html>

Slika 3.1. Osnovne komponente ekspertnih sustava, preuzeto sa :

<http://intelligence.worldofcomputing.net/ai-branches/expert-systems.html#.VVyNUY4vnLV>

Slika 3.2. Primjer stabla odlučivanja (Bell J., 2014, str. 99)

Slika 3.3. Process rudarenja podataka korištenja Interneta, (S.Jagan, i Dr.S.P.Rajagopalan, 2015, str. 8)

Slika 4.1. Građa neurona (Warwick, 2012, str. 91)

Slika 4.2. Razlike između digitalnih računala i mozga (Bašić, Ćupić, 2008, str. 5)

Slika 4.3. McCulloch-Pitts model umjetnog neurona (Warwick, 2012, str. 93)

Slika 4.4. Povijest razvoja umjetnih neuronskih mreža, preuzeto sa:

<http://www.iro.umontreal.ca/~bengioy/dlbook/intro.html>

Slika 4.5. Interakcija agenta s okolinom putem senzora i efektor (Norvig i Russell, 1995, str. 32)

Slika 5.1. Samoupravljaajući automobil, preuzeto sa:

<https://plus.google.com/+GoogleSelfDrivingCars/photos>

Slika 6.1. IBM-ova ploča sa 16 Čipova, preuzeto sa:

<http://www.research.ibm.com/articles/brain-chip.shtml>