

Genādijs MOSKVINS



**INTELEKTUĀLĀS
SISTĒMAS UN TEHNOLOĢIJAS**

Latvijas Lauksaimniecības universitāte
Tehniskā fakultāte
Lauksaimniecības Enerģētikas institūts



G.MOSKVINS

Intelektuālās sistēmas un tehnoloģijas

Mācību līdzeklis studiju priekšmetam
„Intelektuālās sistēmas un tehnoloģijas”

Jelgava 2008



Mācību līdzeklis sagatavots un izdots ESF projekta „Inženierzinātņu studiju satura modernizācija Latvijas Lauksaimniecības universitātē” ietvaros, projektu līdzfinansē Eiropas Savienība.

Moskvins G. Intelektuālās sistēmas un tehnoloģijas: Mācību līdzeklis. – Jelgava: LLU, 2008. – 136 lpp.

Mācību līdzeklī apkopoti intelektuālo sistēmu un tehnoloģiju (IST) pamatu uzbūves modeļi, principi, aktualitātes un problēmās, izskatīta mūsdienu IST struktūra un aktuālie uzdevumi, doti IST elementu un funkciju pamatjēdzieni un definējumi, veikta mūsdienu IST iespēju analīze un dots to attīstības tendenču novērtējums. Pamatojoties uz mūsdienu IST sistēmisko analīzi un nākotnes perspektīvām pētniecībā, lauksaimniecībā, ekonomikā, rūpniecībā un bionikā izveidotas IST vadības paradigmas un atbilstoša zināšanu bāze. Grāmatā tiek apspriesti daži teorētiskie, metodoloģiskie un praktiskie IST aspekti, kas kopumā papildina zināšanu bāzi par IST iespējām, uzbūves principiem un attīstības perspektīvām. Mācību līdzeklis var būt noderīgs izstrādājot intelektuālās sistēmas un tehnoloģijas rūpniecībā, ekonomikā un lauksaimniecībā.

ISBN 978-9984-784-62-5

© Genādijs Moskvins
© LLU Tehniskā fakultāte

Saturs

Ievads.....	5
1. Intelektuālo sistēmu un tehnoloģiju pamatjēdzieni un terminoloģija.....	12
2. Domas abstrakciju evolūcijas glosārijs.....	19
3. IST teorijas pamatu apgūšanas metodoloģija	22
3.1. Mehānikas metodoloģija	25
3.2. Elektronikas metodoloģija.....	26
3.3. Datormodelēšanas metodoloģija	28
3.4. Neiromodeļu metodoloģija.....	29
4. Kā domā cilvēks?.....	32
5. IST aktuālie uzdevumi	33
6. Saprātīgas domas akts.....	36
7. Par prātīguma principiem	39
8. Intelektā definīcijas.....	41
9. Domājošo mašīnu intelektuālie modeļi	45
10. Ārprāta brāļi	51
11. Mākslīgā intelekta filozofiski-metodoloģiskās problēmas	52
12. Mākslīgais intelekts un psiholoģijas problēmas	58
13. Mākslīgā intelekta sistēmu uzbūves principi.....	61
14. Domājošu mašīnu veidi	69
15. Domas vēsture.....	70
16. Mākslīgo neironu tīklu gūstā	74
17. Nanoroboti	75
18. IST, MI ... Kas tālāk?	77
19. IST attīstība	82
20. „Cilvēks” - tas skan sarežģīti!	85
21. Mašīnas intelekts.....	89
22. Skaitļošanas intelekts un modelēšana	93
23. Tūringa domājošās mašīnas	95
24. Bionisko intelektuālo sistēmu telpa laika korekcija	98
25. Mākslīgie neironu tīkli.....	104

26. Mākslīgie neironu tīkli un ekspertu sistēmas.....	106
27. Neironu tīklu nākotnes perspektīvas.....	107
28. Adaptācijas labirintos.....	108
29. Risināšanas veidi.....	109
30. MNT īpašības un attīstības virzieni	111
31. Mākslīgā intelekta aktualitātes.....	117
32. Virtuālā realitāte.....	122
33. Ģeniju inkubators.....	124
34. Izdzīvošana un mākslīgais intelekts.....	125
35. Gaidot apokalipsi	129
Literatūra	130

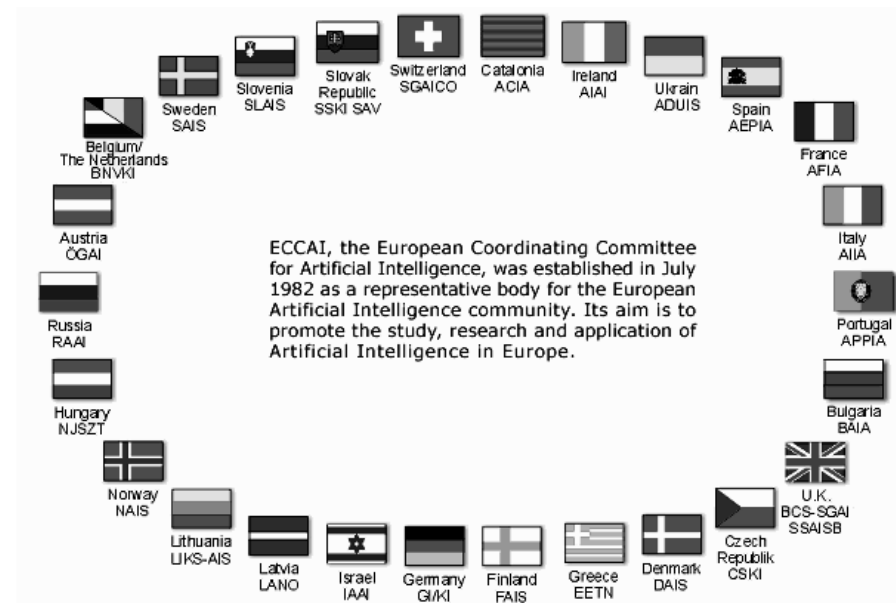
Ievads

Gēniji ir nemirstīgi. Toties muļķi ilgāk dzīvo (M.Geny)

Intelektuālās sistēmas un tehnoloģijas, ekspertu sistēmas un mākslīgais intelekts ir 21. gadsimta vadošais un perspektīvākais attīstības virziens, kas paver jaunas iespējas ķīmisko, bioloģisko, ekoloģisko un citu procesu, optimālās vadības, atbilstības kontroles, testēšanas un citu intelektuālo zināšanu iegūšanas procesu optimizācijā, izmantojot informācijas tehnoloģiju sasniegumus. Intelektuālās tehnoloģijas aizvien plašāk ieviešas tehniskajās un bioloģiskajās sistēmās. Adaptācijas princips nodrošina sarežģītu tehnoloģisko procesu nepārtrauktu optimālu vadību objektos ar nestacionāriem un nedeterminētiem parametriem un stohastiskām perturbācijām. Šādas īpašības ir raksturīgas lielākajai daļai lauksaimniecības objektu, procesu un sistēmu, jo tur dominē augi un dzīvnieki, pārstrādes un selekcijas produkti. Intelektuālo modeļu uzvedības nenoteiktības pakāpe ir jo lielāka, jo augstāks ir modelējamais intelekts. Jebkuras jaunas zināšanas ir savdabīgs pievilkšanas pols, domu barojošā vide. Intelektuālo tehnoloģiju attīstības pamatā ir bioloģisko objektu un subjektu analogijas noteikšana, modelēšana, maņu orgānu elektroniskā, bioķīmiskā, nanomolekulārā un cita fenomenoloģiskā imitācija, strukturālā modelēšana un pielāgošanas iespēju izpēte konkrētam automatizācijas objektam vai tehnoloģiskam procesam. MI uzņemas savdabīga intelektuālā interfeisa lomu, kas interpretē un lietotājam piedāvā nevis miglainas informācijas kopas, bet lietošanai gatavas, mērķtiecīgas un tādēļ lietderīgas zināšanas un apdomātus lēmumus, ieteikumus un rekomendācijas.

21. gadsimta pasaules materiālā un humanitārā iekārta balstīsies uz informācijas plūsmām un intelektuālām tehnoloģijām. Zinātne kļūs par galveno sabiedrības radošo spēku, bet zināšanu iegūšanas un izmantošanas tehnoloģijas kļūs par pamatprocesiem, kas 21. gadsimtā noteiks sociālo, ekonomisko un kultūrālo realitāti. Diemžēl var konstatēt, ka pārtikas un citu cilvēces izdzīvošanai nepieciešamo resursu daudzums uz mūsu planētas strauji samazinās. Pēc ANO datiem 1991. gadā pārtika kopumā pietika 100 dienām, bet 2007. gadā – jau tikai 23 dienām. ANO oficiāli atzīst, ka 1/3 cilvēces uz mūsu planētas šodien nav nodrošināta pat ar to minimālo resursu daudzumu, kas garantētu tai normālu fizioloģisko izdzīvošanu. Ekoloģiskās un sociālās problēmas,

pārtikas bāzes globālā pasliktināšanās, arī cilvēces kopskaita pieaugums liek meklēt gan tradicionālus papildus resursus, gan principiāli jaunās tehnoloģiskās iespējas.



1. att. Latvijas Nacionālā Automātikas Asociācija (LANO) – ECCAI dalībniece

Rezerves šajā jomā meklējamas, galvenokārt, izmantojot intelektuālo tehnoloģiju iespējas un zinātnes potenciālu. Pirmkārt tas attiecas uz bioinženieriju, nanotehnoloģijām (NT), informācijas tehnoloģijām (IT) un mākslīgo intelektu (MI). Latvija (LANO) ir Eiropas Mākslīgā intelekta koordinācijas komitejas (ECCAI) dalībvalsts (1. att.).

Mūsdienās intelektuālās tehnoloģijas aizvien plašāk ieviešas tehniskajās un biosistēmās. Adaptācijas princips nodrošina sarežģītu tehnoloģisko procesu nepārtrauktu optimālu vadību objektos ar nestacionāriem, nedeterminētiem parametriem un stohastiskām perturbācijām. Šādas īpašības ir raksturīgas lielākajai daļai lauksaimniecības objektu, procesu un sistēmu, jo tur dominē augi un dzīvnieki, to pārstrādes un selekcijas produkti [1, 2, 14-28].

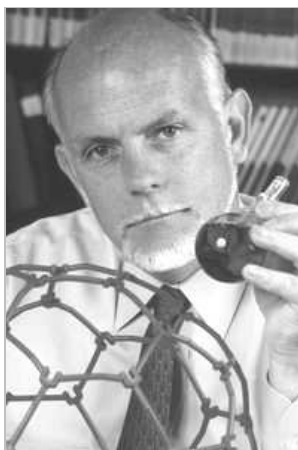
Jaunās intelektuālās tehnoloģijas drīz vadīs ne tikai mūsu telpisko ķermeni, bet, pats svarīgākais, indivīda dabu un viņa fizisko ķermeni. 21. gadsimta tehnoloģiju – robottehnikas, gēnu inženierijas, mākslīgā intelekta, informācijas tehnoloģiju un

nanotehnoloģiju attīstība notiek tieši tādā virzienā. Iedomājieties, ka jūs skatāties jaudīgā mikroskopā un atklājat neticami mazus mehānismus, kuri sadala viņus aptverošo vielu atsevišķās molekulās, bet pēc tam savāc molekulas kopā tā, lai veidotos šo mehānismu precīzas kopijas. Kopijas, protams, darīs to pašu. Pēc 20 paaudzēm katrs mehānisms pārveidosies miljonos mehānismu. Vai izdosies apturēt šo procesu, vai tie savaldzinās visu pasauli? Vai tā nav kaut kāda fantastikas pasaka par tehnoloģiju, kura vairs nav kontrolējama? Tā ir pasaule, kurā mēs dzīvojam, kur tādi mehānismi eksistē praktiski visur.

Tādu mehānismu neskaitāmi miljoni dzīvo katra cilvēka zarnās. Mēs tos saucām par baktērijām, un viņi savaldzināja pasauli jau miljardus gadu atpakaļ, ilgi līdz tam brīdim, kad parādījās cilvēki. Mēs apejamies ar tiem kā nākas, pretējā gadījumā tie mūs iznīcinātu. Evolucionisti līdz galam tā arī neizskaidroja, ko uzskatīt par baktēriju priekštečiem, bet atkārtot dabīgo eksperimentu mēs nevaram. Taču dabai bija daudz laika uz šo lietu – miljardiem gadu, bet mums, mirstīgajiem, ir jānodemonstrē veiksmīgu rezultātu līdz tam momentam, kad beigsies grants uz zinātniski-pētniecisko darbu veikšanu. Jebkurā gadījumā, pat vienkāršākajām baktērijām ir apbrīnojami sarežģīta uzbūve ar DNS virknēm, kuras satur pilnas instrukcijas, kas attiecas uz vielu maiņu un vairošanos. Tomēr daudzi biologi tic, ka viņi atrodas uz mikroba sintēzes sliekšņa laboratorijas apstākļos. Kļuvis iespējams mākslīgi radīt garas DNS virknes, tagad zinātnieki nosaka, kādi gēni ir vissvarīgākie. Ja laboratorijā radīs jaunu mikrobu, tas nozīmē, ka zinātnieki seko dabas receptēm. Bet ja būs radīta pavisam jauna dzīvības forma „no nulles”, nesejojot dabai? Kas būs, ja šī dzīvības forma būs vairāk līdzīga mehānisko ierīču veidam, kurus jau prot konstruēt cilvēks, bet tikai tā būs ļoti maza izmēra? Cik maza? 1959. gadā amerikāņu fizikas sabiedrības sanāksmē Ričards Feinmans, laikam pats slavenākais mūslaiku fiziķis, uzstājās ar referātu ar nosaukumu „*Tur, lejā ir daudz brīvas vietas*”. Samazinātus līdz atoma mērogam, kā atzīmēja Feinmans, visus 24 „*Britānikas*” enciklopēdijas sējumus var izvietot uz kniepadatas galviņas. Viņš aicināja savus kolēģus izstrādāt iedarbības iespējas uz vielu un to vadīšanu atomu līmenī.

Pēc trīsdesmit gadiem zinātnieki IBM laboratorijā Kalifornijā izvietoja 35 ksenona atomus uz niķeļa kristāla virsmas tādā veidā, lai veidotos uzraksts „IBM” ar

lieliem burtiem. Tas bija īstenots, izmantojot skenējošo, tuneļa tipa mikroskopu (STM), ierīci, kuru izstrādāja *IBM laboratorijā Cīrihē*, kura izgudrotāji ieguva Nobela prēmiju fizikā. To pasludināja par nanorevolūcijas sākumu – tehnoloģiju rašanās moments mērogiem, kas tūkstoš reizes mazāki par mikroelektronikas pasaulē sastopamajiem. Pašlaik tomēr to saukt par „revolūciju” nevar. IBM zinātnieki pierādīja, ka var operēt ar atsevišķiem atomiem, bet šī iespēja līdz šim vēl nedeva kaut kādu praktisku labumu. Vēlāk, 1996. gadā Ričards Smellijs (*Richard Errett Smalley, 1943–2005*) no *Raisa Universitātes* ieguva Nobela prēmiju ķīmijā par fullerenu atklāšanu – oglekļa karkasa struktūru ar milimikrona mērogu, ar apbrīnojamām īpašībām un daudzām izmantošanas iespējām.



2. att. 1996. gada Nobela prēmijas laureāts ķīmijā Ričards Smellijs

Līdz šim tirgū vēl nav milimikrona mēroga produktu, bet valdības visā pasaulē iegulda milzīgus līdzekļus nanotehnoloģiju attīstībā, cerot, ka tās mainīs dzīvi tikpat pamatīgi, kā tas notika ar mikroelektronikas attīstību. Futurologs *Ēriks Drekslers* savā grāmatā „Radošas mašīnas: sakās nanotehnoloģijas laikmets” (1986) paredz tādu nanomehānismu rādīšanu nākotnē, kas spēj manipulēt jebkuru vielu ar atomu precizitāti. *Ē.Drekslers* agrāk bija Pareģošanas institūta priekšsēdētājs (*Foresight Institute, ASV*), kura pamatuzdevums ir sagatavot pasauli revolūcijai nanotehnoloģiju jomā. Tagad *Ē.Drekslers* strādā *Nanorex, Inc. ASV*. Viņš redz pasauli, kurā neticami mazi autosaliekamie roboti-montētāji (viņš tos sauc par „asembleriem”) veiks visu

nepieciešamo darbu, regulējot ķīmiskās reakcijas, vadot reaģējošo molekulu stāvokli ar atomu precizitāti. Nodrošināti ar pirmatnējo materiālu, assembleri var būt ieprogrammēti uz visa tā radīšanu, kas varētu mums noderēt, tai skaitā arī jaunu assembleru radīšanu. Bet vai eksistē nanotehnoloģiju un MI tumšā puse? Ričardam Smellijam radās jautājums, kas varēs likt autosalikamiem nanorobotiem beigt „košļāt” un radīt sev līdzīgus, „kamēr viss pasaulē nepārvērtīsies pelēkā lipīgā masā”.



3. att. **Futurologs Ēriks Drekslers**

Daudzus cilvēkus, ieskaitot princi Čārlzu Lielbritānijā, uztrauc šī „lipīgā masa”, sakarā ar ko parādījās prasības aizliegt tālākos pētījumus nanotehnoloģiju jomā. Bet tā rīkoties būtu ļoti kļūdaini. Pirmkārt, „pelēka lipīga masa” – ir tikai iedomājamā bīstamība. Assembleri nepastāv, un neviens vēl īsti nezina, kā tos radīt. Bet ja assembleri pastāvētu, viņi sadurtos ar tādu pašu ierobežojumu kā baktērijas – viņi nevarētu pārvietoties lielos attālumos bez braukšanas līdzekļiem un pirmējais materiāls kaut kādā brīdī vienkārši beigsies. Assembleri vairs nav vienīgais manipulēšanas veids ar vielām atomu līmenī. Neviens no mūsdienu valsts finansējamiem pētījumiem nenodarbojas ar montētāju izstrādāšanu. Zinātnieki, kuri izgudroja STM, nemēģināja radīt tādus montētājus un nekad pat nebija dzirdējuši par Ēriku Dreksleru [108].

Viņi tikai mēģināja apgūt virsmu fizikas fundamentālos principus. Likās, ka, lai ieskatītos nākotnē, futurologu palīdzība nav nepieciešama. Tomēr neticamais notika: mikrorobots jau ceļo pa asinsvadiem! Pēc 21 gada no futuristiskām prognozēm (30

gadu vietā, kā bija domāts) – 2007. gada sākumā nanorobots īstenoja savu pirmo ceļojumu pa asinsvadiem. Par izcilo veiksmi paziņoja Monreālas Universitātes un Politehniskā institūta zinātnieki, kuri strādāja doktora *S.Martela* vadībā. Viņiem pirmajiem izdevās ievadīt cūkas miega artērijā mikroskopisko zondi-robotu, kontrolēt un vadīt visas tā kustības. Viena no mūsdienu medicīnas vispievilcīgākajām idejām ir sīku robotu radīšana, kuri varētu atrast organismā slimus audus un pat šūnas, pievadītu tiem zāles vai, otrādi, kaut ko likvidētu, noteiktu precīzu diagnozi un tml. Līdz šim visvairāk sapņa realizācijai pietuvinājušās 14 milimetru kapsulas-videokameras, kuras ļauj apskatīt tievo zarnu. Pacients norij kapsulu un tā nodod uz datoru visu, ko „redz” apskates ceļā. Bet šodien tāda kapsula maksā 4-5 tūkstošus ASV dolāru. Bet vienai injekcijas dozi organismā ir nepieciešami miljardi nanorobotu! Iedomājieties, cik tas varētu maksāt?

Mikro- un nanorobotu radīšana ir nākošais etaps, kam pasaules zinātne šodien ir daudz tuvāk, nekā domā nanoskeptiķi. Daudzās pasaules laboratorijās notiek īsta cīņa šajā sfērā. Mikrorobotu (vai „mikrobotu”, kā tos jau sauc) prototipus radīja amerikāņi. Tuvu rezultātam ir Ķīnā, Austrālijā un Šveicē. Taču tur koncentrējas uz pašām ierīcēm, cenšoties sasniegt maksimāli iespējamo miniaturizāciju un atstrādājot to funkcijas. Dzīvajā organismā pirmie tehnoloģiju izmēģināja kanādieši, apsteidzot visus. Bet viņiem vēl nav mikrobota kā tāda – tā lomu izpildīja „tukša” feromagnētu sfēra ar diametru 1,5 mm. Zinātnieki ievietoja parasto cūku standarta magnēta rezonanses tipa tomogrāfā (MRT) magnētiskajā laukā. Šī ierīce dod iespēju ne tikai pēfīt dažādus orgānus, bet arī pārvietot metāliskus priekšmetus. Vadot MRT gradientās spoles, zinātnieki kustināja feromagnēta krellīti un novēroja to monitorā. Kustības ātrums sasniedza 10 cm sekundē. Izrādījās, ka MRT ļauj precīzi vadīt zondi pa iepriekš iezīmētu maršrutu lielā asinsvadā. Tagad autori mēģina samazināt zondes izmēru, lai tas varētu iekļūt sīkajos asinsvados. Vienlaicīgi Kanādas speciālisti rada veselu ģimeni mikrobotu, kuri varētu pievadīt zāles vai diagnostiskos sensorus vajadzīgajās ķermeņa daļās. Datori tagad var pat lasīt mūsu domas, bet cilvēkam ir pieejama telekinēze. ASV Masačūsetsas štata kompānija *Cyberkinetics* izstrādāja unikālu mikročipu ar 4 kvadrātmilimetru izmēru, kurš ļaus sūtīt cilvēka smadzeņu signālus datoram.

Tādus čipus implantēs smadzenēs slimiem cilvēkiem ar motoriskās sistēmas traucējumiem, kuri nevar patstāvīgi veikt virkni nepieciešamo darbību. Kompānijai nācās izcīnīt valsts atļauju tāda eksperimenta veikšanai ar cilvēka smadzenēm, bet beigu beigās sankcija tika dota. Implantētais čips ļaus ievērojami atvieglot dzīvi invalīdiem. Virkni funkciju izpildīs dators. Dators var palīdzēt slimiem cilvēkiem darīt visu nepieciešamo, tai skaitā ar elektrisko impulsu palīdzību viņi varēs stimulēt pašu muskuļus. Agrāk bija veikts čipa implantācijas eksperiments pērtiķiem. Bija radīta programma, ar kuras palīdzību pērtiķiem izdevās pārvietot kursoru datora monitorā ar smadzeņu impulsa palīdzību, bez kādām fiziskām pūlēm. Piemēram, pacientu var palūgt iedomāties, ka viņš kustina roku, bet dators uztvers atbilstošo smadzeņu impulsu un izpildīs to. Bet ko šodien var „mākslīgās smadzenes”? Mūsdienās tāds sarežģīti organizēts mehānisms, kā smadzenes, mums praktiski nav īpaši vajadzīgs. Mēs taču to izmantojam ļoti nepietiekoši un primitīvi – mūsdienu tehnocivilizācijas attīstības sekas.

Bet smadzeņu iespējas taču ir milzīgas. Ar to izzināšanu cilvēcei iestāsies kopīgās apskaidrības un apgarotības stāvoklis. Atklāsies spēja paredzēt notikumus, mēs apzināsimies to kaitējumu, ko padarām sev un pasaulei, varēsim pagarināt savu dzīvi ar to, ka veiksīm organisma profilaksi un ārstēsim to paši. Kad tiks izpētītas cilvēka smadzenes, mēs zināsim visu ne tikai par sevi, bet arī par Visumu. Sarežģītais cilvēka organisms no dabas strādā apbrīnojami saskaņoti, tajā ir viss nepieciešamais ilgai dzīvei. Bet, kad mēs sakām: sirds darbojas kā sūknis, būtu pareizāk teikt: sūknis darbojas kā sirds! Mūsu iekšējie orgāni – ir tikai zobrautiņi sarežģītā mehānismā. Tāpēc mums nav jāsatraucas par to, ka mēs sastāvēsim no mākslīgi radītiem sērijveida „nano-zobrautiņiem”. Unikālas paliks tikai mūsu emocijas, jūtas, pārdzīvojumi un sakrātā notikumu atmiņa. Intelekti nebūs „zobrītāju” prerogatīva, bet kļūs par kopīgu, visiem pieejamu īpašumu, spēcīgu un drošu komunikācijas līdzekli. Tas būs visai saprātīgs – gandrīz kā „dzīvs”, kaut gan būs radīts „mākslīgi”. Bet vai saprāts ir smadzeņu darbības produkts vai pastāv atsevišķi – to mēs vēl nezinām. Ieskatoties nanotehnoloģijas un mākslīgā intelekta veiksmēs un attīstības tempos, mūsdienu cilvēku paaudzi sagaida apbrīnojami atklājumi.

Mācību-metodiskā līdzekļa *Intelektuālās sistēmas un tehnoloģijas* (IST) mērķis ir metodiskā palīdzība inženieru un citu specialitāšu studentiem, maģistrantiem un doktorantiem sākotnējo zināšanu apgūšanā IST jomā, kā arī grūti formalizējamo uzdevumu risināšanā, kas līdz šim joprojām tiek uzskatīti par cilvēka prerogatīvu, zināšanu iegūšanā un papildināšanā par cilvēka domāšanas veidiem, kā arī par dažām to realizācijas metodēm ar tehnisko līdzekļu, programmu un algoritmu palīdzību. Šādi IST izzināšanas modeļi tiek izmantoti, lai izpētītu cilvēka intelekta īpašības (kognitīvā psiholoģija) un radītu tādas tehniskās ierīces, programmas, algoritmus un operācijas, ko pagaidām var veikt tikai cilvēks.

1. IST pamatjēdzieni un terminoloģija

Daba tiecas pēc ģeometriskiem risinājumiem (Ričards Smellijs, Nobela prēmijas laureāts)

Pasaules uzskats - ir cilvēka dzīves izpratnes pamatkategorija, kas raksturo organizācijas sistēmu galvenos funkcionālos procesus, kuri veicina fiziskās un garīgās būtnes izveidošanas, pašsaglabāšanās un pašattīstības procesus.

Informācija ir dzīves un reizē - garīgas eksistences pamats. Viss dzīvais pasaulē sastāv no aktīviem subjektiem. Tāpēc arī viss dzīvais ir atsaucīgāks nevis uz mehānisko, bet uz sistemātisko enerģētisko iedarbību. Pie tam atsaucas ne tikai atsevišķie struktūras elementi, bet arī visa iekšējo struktūru hierarhija.

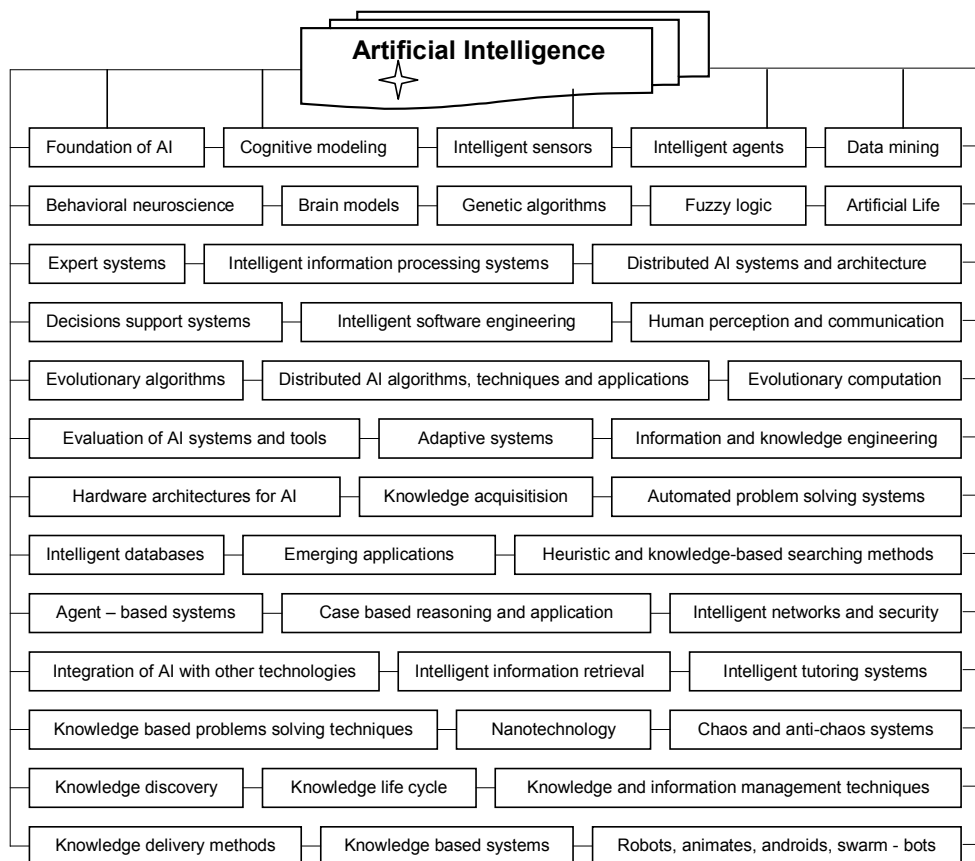
Pasaules lineārā uztvere - process, kad cilvēki pasauli cenšas saprast un izskaidrot ar Ņūtona mehānikas palīdzību.

Metriskā tēla izveidošana - mēģinājums atrast, saprast un novērtēt pētāmas vides pozitīvās, lietderīgās īpašības.

IST perceptuālie modeļi nosaka vadāmās sistēmas uzvedības principus konkrētos apstākļos. Bet radoši MI risinājumi var būt arī bīstami cilvēcei ar vilinošu ideju izmēģināt tos praksē. Noņemt MI virtuālās drošības problēmu traucē cilvēku nepilnvērtīgais perceptuālisms, pasaules ainas izpratne un uztvere, cilvēku parapsiholoģiskās barjeras un dogmatisms. Mūsdienu pasaulē būtisks progress tiek panākts tikai tādos gadījumos, kad daļu cilvēka intelektuālās slodzes pārņem dators.

Viena no maksimālā progresā panākšanas iespējām IST jomā ir „mākslīgais intelekts” (MI), kad dators pārņem ne tikai viena tipa operācijas, kas vairākkārtīgi atkārtojas, bet arī pats spēj apmācīties [3]. Pilnvērtīga „mākslīgā intelekta” radīšana atver cilvēcei jaunas attīstības apvārsņus.

Termina *intelekts (intelligence)* izcelsme ir no latīņu vārda „*intellectus*” – kas nozīmē gudrību, saprātu, prātu, cilvēka domāšanas spējas. Atbilstoši mākslīgais intelekts (*AI - Artificial Intelligence*) – MI parasti tiek definēts kā automātisko sistēmu spēja pārņemt atsevišķas cilvēka intelekta funkcijas, piemēram, izvēlēties un pieņemt optimālo risinājumu (jo cilvēka domāšana, kā izziņas process, ir racionāls, nevis optimāls) uz agrāk iegūtas pieredzes un ārējo iedarbību analīzes pamata.



4. att. Ar „Artificial Intelligence” saistītās zinātnes (avots: Moskvins, 2008)

Mēs sauksim par „intelektu” smadzeņu spēju risināt intelektuālos uzdevumus zināšanu iegūšanas, arhivēšanas, atcerēšanās un mērķtiecīgas pārveidošanas ceļā, kas iegūta apmācībās, pieredzes uzkrāšanas un adaptācijas procesā dažādos jaunos apstākļos. Citiem vārdiem sakot intelekts ir spēja objektīvi analizēt situāciju un pieņemt optimālo lēmumu. Šajā definīcijā ar terminu „zināšanas” tiek domāts ne tikai tā informācija, kas nonāk līdz smadzenēm caur maņu orgāniem. Šāda tipa zināšanas ir ārkārtīgi svarīgas, bet nepietiekamas IST intelektuālai darbībai, jo apkārtējās vides objektiem piemīt īpašība ne tikai iedarboties uz mums, bet arī mijiedarboties savā starpā noteiktā veidā.

Saprotams, ka lai īstenotu intelektuālo darbību apkārtēja vidē, vai vismaz elementāri nodrošināt savu eksistēšanu, cilvēku zināšanu sistēmai iepriekšēji ir nepieciešams šis pasaules modelis. Šajā informatīvajā apkārtējās vides modelī reālie objekti, to īpašības un savstarpējās attiecības ne tikai tiek attēlotas, atveidotas un iegaumētas, bet arī, kā tas tiek atzīmēts šajā intelekta definīcijā, var tikt mērķtiecīgi un saprātīgi pārveidotas. Pie tam būtiski ir tas, ka ārējās vides modeļa veidošana notiek apmācības, pieredzes uzkrāšanas un adaptācijas procesu gaitā pie dažādiem jauniem apstākļiem.

Mēs pielietojām terminu *intelektuālais uzdevums*. Lai paskaidrotu, ar ko atšķiras intelektuālais uzdevums no parasta uzdevuma, nepieciešams ievadīt jaunu terminu „algoritms” – viens no kibernetikas galvenajiem terminiem. Ar *algoritmu* saprot tiešu norādījumu par īstenošanu noteiktā sistēmu operāciju kārtībā jebkura uzdevuma risināšanai no kādas noteiktas uzdevuma klases (liela daudzuma, kopas). Termins „algoritms” nāk no uzbeku matemātiķa *Al-Horezmi* vārda, kurš vēl IX gadsimtā piedāvāja vienkāršākos aritmētiskos algoritmus. Matemātikā un kibernetikā noteikta tipa uzdevumu klase tiek uzskatīta par atrisinātu, kad to risināšanai tiek noteikts algoritms. Algoritmu meklēšana un atrašana ir cilvēka dabisks mērķis, kad viņš risina paša izstrādātas uzdevumu klases. Algoritma atrašana dažāda tipa uzdevumam ir saistīta ar smalkām un sarežģītām pārdomām, kas pieprasa lielu atjautību un augstu kvalifikāciju. Bet kas attiecas uz uzdevumiem, kuru risināšanas algoritmi jau ir noteikti, tad, kā atzīmē pazīstams speciālists MI jomā – *M. Minskis* [59-61] ir lieki piedot tiem tādas mistiskas īpašības, kā „intelektualitāte”.

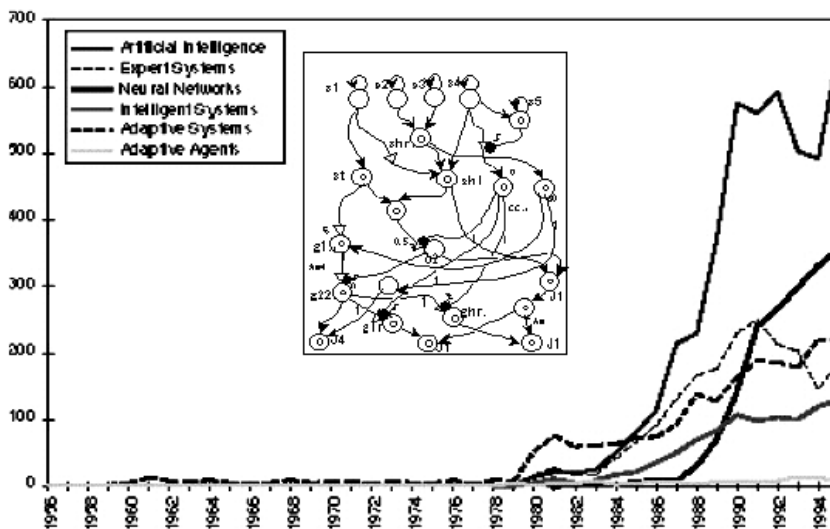
Patiešām, pēc tam, kad šāds algoritms jau ir atrasts, attiecīgo uzdevumu risinājums paliek tāds, ka to var precīzi izpildīt gan cilvēks, gan attiecīgi ieprogrammētā skaitļošanas mašīna vai robots, kam nav nekādas saprašanas par uzdevuma būtību. Pieņemts uzskatīt, ka šāda veida darbība pieprasa cilvēka intelekta līdzdalību. Uzdevumi, kas saistīti ar optimāla algoritma atrašanu noteiktai uzdevumu klasei, kas izriet no cilvēka dabiskās racionālas domāšanas mērķtiecīgas darbības procesā saucsim par „intelektuāliem”.

Tiek pieprasīts tikai, lai persona, kas risina uzdevumu, būtu spējīga izpildīt tās elementārās operācijas, no kurām veidojas process, un, bez tam, lai tā pedantiski un akurāti vadītos pēc piedāvātā algoritma. Šāda persona, darbojoties, kā tādos gadījumos saka, „tīri automātiski”, var veiksmīgi atrisināt jebkuru apskatāmā tipa uzdevumu. Tādēļ ir pilnīgi dabiski izslēgt no intelektuālo uzdevumu klases tādus uzdevumus, kuru risināšanai eksistē standarta risināšanas metodes.

Par šādu uzdevumu piemēru var būt skaitļošanas uzdevumi: lineāro algebras vienādojumu sistēmas risinājums, diferenciālvienādojumu skaitliskā integrēšana utt. Šāda veida uzdevumu risināšanai ir standarta algoritmi, kas paredz virkni noteiktu elementāro darbību secību, kas var tikt viegli realizēta skaitļošanas mašīnas programmas veidā. Pretēji tam intelektuālo uzdevumu plašai klasei, tādu kā tēlu atpazīšana, šaha spēlēšana, teorēmu pierādīšana u.tml., šī formālā risinājuma procesa sadalīšana atsevišķos elementāros soļos bieži ir visai apgrūtināta, pat ja to risinājums nav sarežģīts. Tādā viedā mēs varam pārfrāzēt intelekta definīciju, kā universālu „virsalgoritmu”, kurš spēj „radoši” veidot konkrētu uzdevumu risināšanas algoritmus.

Programmētāja darbības produkts ir programmas – algoritmi tūrā veidā. Tieši tāpēc, pat MI elementu radīšanai vajadzētu stipri palielināt programmētāja darba ražīgumu. Mēs saucsim par domāšanu, vai intelektuālu darbību tādu smadzeņu darbību, kam ir intelekts. Intelekts un domāšana ir organiski saistīti ar tādu uzdevumu risināšanu, kā teorēmu pierādīšana, loģiskā analīze, situāciju atpazīšana, uzvedības plānošana, spēles un vadīšana nenoteiktos apstākļos. Intelektam raksturīgās īpašības, kuras izpaužas uzdevumu risināšanas gaitā, ir spēja apmācīties, apkopot, krāt pieredzi (zināšanas un prasmes) un adaptēties pie mainīgiem apstākļiem uzdevumu risināšanas gaitā. Pateicoties šādām intelekta īpašībām smadzenes var risināt visādus uzdevumus,

kā arī viegli pāriet no viena uzdevuma risināšanas uz citu. Šādā veidā smadzenes, kuras ir apveltītas ar intelektu, ir universāls plaša uzdevumu spektra risināšanas līdzeklis, tai skaitā ir neformalizēti uzdevumi, kuriem nav standartu, iepriekš zināmu risināšanas metožu. Der paturēt prātā arī citas, tīri funkcionālas intelekta definīcijas. Tā, pēc *A. N. Kolmogorova*, jebkura materiāla sistēma, ar kuru var pietiekami ilgi apspriest zinātnes, literatūras un mākslas problēmas, ir intelektuāla. Par citu funkcionālo IST „*intelektuālās uzvedības*“ traktējuma piemēru var kalpot *A. Tūringa* [1, 4, 16, 17] pazīstamā definīcija. Tās jēga ir sekojoša. Dažādās istabās atrodas cilvēks un ESM mašīna. Tie nevar viens otru redzēt, bet ir iespēja apmainīties ar informāciju (piemērām, ar elektroniskā pasta palīdzību). Ja dialoga gaitā starp spēles dalībniekiem cilvēkiem neizdodas noteikt, ka viens no dalībniekiem ir mašīna, - raksta *Tūringa*, - mēs esam spiesti daudz domāt par to procesu, kura rezultātā cilvēka smadzenes ir sasniegušas savu tagadēju stāvokli. No tā izriet svarīgs nosacījums IST attīstībai: lai notiktu sistēmas apmācība, vismaz vienam no procesa (dialoga) dalībniekam (skolnieks-mašīna-skolotājs) ir jābūt „*gudrākām*“ par citiem. Šis nosacījums realizē apmācības sistēmas motivācijas funkciju.

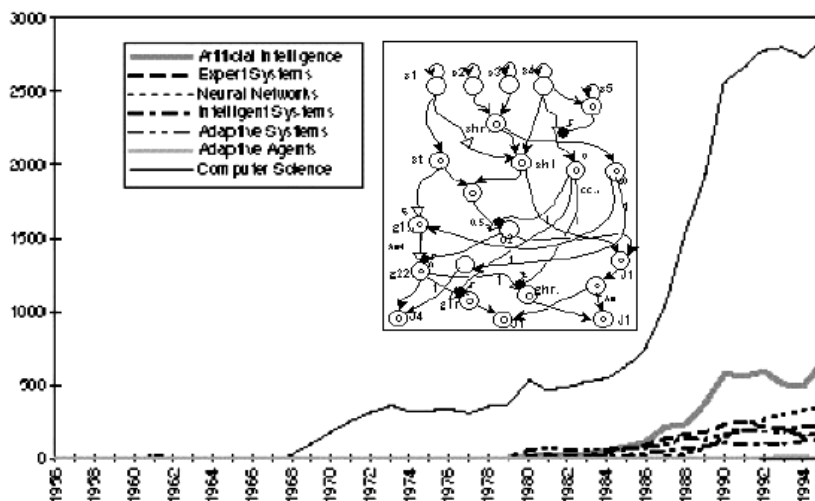


5. att. Ar IST saistīto zinātnu attīstības tendences. Doktora disertāciju skaita ikgadējs pieaugums IST zinātnu jomā (avots: *Dissertation Abstracts Online database from UMI Company*)

Sistēmas intelekta līmeni šeit nosaka „skolotāja“, eksperta intelekta līmenis. Sistēmas intelekta „piesātinājuma“, „maksimuma“ situācija rodas tad, kad dialoga dalībnieku atbildes vairs būtiski neatšķiras viens no otra un savas un svešas atbildes pēc būtības, ir gandrīz identiskās.

Kādēļ mums tā vietā, lai censtos izveidot programmu, kas imitē pieauguša cilvēka smadzenes, nepacenstos radīt tādu programmu, kura imitētu bērna intelektu? Jo, ja bērna intelekts saņem attiecīgo trenēšanu un audzināšanu, tas kļūst par pieauguša intelektu. Mūsu cerības ir saistītas ar to, ka ierīce, tam līdzīga, var būt viegli ieprogrammējama.

Tādā veidā, mēs sadalīsim intelekta attīstības problēmu divās daļās: „programmas-bērna” radīšanas problēma un šīs programmas attīstības, „audzināšanas” problēma. Tieši šo ceļu praktiski izmanto visas MI sistēmas. Jo ir taču skaidrs, ka praktiski nav iespējams iepriekš ielikt visas zināšanas visai sarežģītā MI sistēmā. Bez tam, tikai šādā ceļā parādīsies augstākminētās intelektuālās darbības pazīmes - pieredzes uzkrāšana, adaptācija u.tml.



6. att. Aizstāvēto doktora disertāciju kopēja skaita pieaugums IST zinātņu jomā
(avots: Dissertation Abstracts Online database from UMI Company).

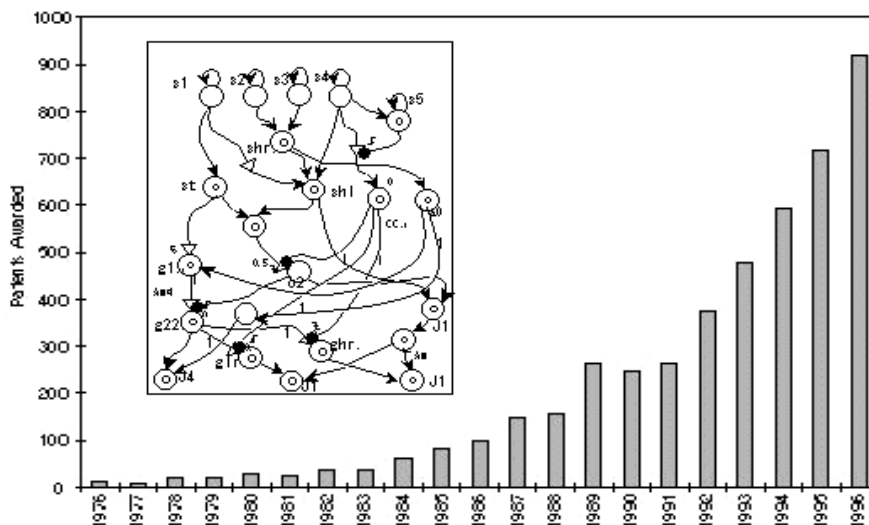
IST un mākslīgā intelekta (MI) pētījumu atšķirības pazīme ir tas, ka parasti MI problēmas risināšanai meklē heuristiku, nevis algoritmu klasiskajā vārda izpratnē. Heuristiskie algoritmi bieži negarantē pareizā rezultāta iegūšanu – vai nu precīzs

risinājums ir pārāk darbietilpīgs, (piemēram, šaha spēlē rodas pārāk liels variantu daudzums), vai nu pati priekšmetu joma nav noteikta un formalizēta un nepastāv labs kritērijs, kas ļauj nolemt, vai risinājums ir pareizs (piemēram, dabiskās valodas apstrāde). Parasti IST vai MI skaitās datorzinātnes un/vai informātikas sfēra, taču tas nav gluži korekti, jo pastāv jomas, kuras ir ļoti tuvas MI, bet atrodas ārpus datoru un informātikas zinātnēm: piemēram kognitīvā psiholoģija. Cilvēka informācijas uztveres *informatīvais duālisms* izpaužas *binaritātes principā*, kas ir gan Visuma, gan informātikas galvenais likums. Tas dod izpratnes slēdžus par laika kategorijām “nākotne”, „tagadne”, „pagātne” un informatīvu procesu ekstrapolācijas, interpolācijas un filtrācijas relatīvismu. IST modelēšanas principi pielietojami gan vadības sistēmās un tehnoloģijās, gan matemātiskajā bioloģijā, gan bionikā, kas pašlaik strauji attīstās.

Attiecības starp tādām dažādām MI “saprāta” īpašībām kā “sajūta”, „uztvere”, „saprātne”, „priekšstats”, „jēdziens”, „spriedums”, „slēdziens”, „asimilācija”, „izdalīšana” utt. raksturo MI “saprāta un organisma” kā sistēmas struktūras veselumu. Tādejādi IST, ES, un MI uzņemas savdabīga intelektuālā interfeisa lomu, kas interpretē un lietotājam piedāvā nevis miglainas informācijas kopas, bet lietošanai gatavas, mērķtiecīgas un tādēļ lietderīgas zināšanas un apdomātos lēmumus, ieteikumus un rekomendācijas.

Šīs attiecības paliek tādas pašas, t.i., invariantas visām MI struktūrām, „*organismiem*”, lai cik atšķirīgi tie būtu savā starpā pēc savas elementu un sastāvdaļu uzbūves. Intelektuālo tehnoloģiju uzbūves un attīstības pamatā ir bioloģisko objektu un subjektu analogijas noteikšana, modelēšana, maņu orgānu elektroniskā, bioķīmiskā, nanomolekulārā un cita fenomenoloģiskā imitācija un strukturālā modelēšana un pielāgošanas iespēju izpēte konkrētam automatizācijas objektam vai tehnoloģiskam procesam. Pasaulē piedejos gados novērojams ievērojams patentu, publikāciju un doktora disertāciju skaita pieaugums IST jomā (5.-7. att.). IST kontroles, vadības un lēmumu pieņemšanas pamats ir informācija par vadības objektu. Pie tam abstraktās informācijas uzkāpšana (evolūcija) līdz konkrētām zināšanām notiek ar psihisko procesu palīdzību, kuri piedalās informācijas pieņemšanā un iztēlošanā pēc noteiktas secības (sajūta, uztvere, priekšstats, jēdziens) līdz domāšanas analītiskiem procesiem

(spriedums, slēdziens, secinājums). Tehniskajās vadības un kontroles sistēmās cilvēka informācijas uztveri ir jāapskata kā percepcijas (jutekliskā) tēla formēšanas procesu. Ar to saprot darbojušies uz cilvēka objekta īpašību subjektīvo atspoguļojumu cilvēka apziņā. Pētījumi rada, ka perceptuālā tēla formēšana ir kompleksais (fāzu) process, kas sastāv no dažām atsevišķām stadijām: atrašana, atšķiršana un noteikšana (pazīšana).



7.att. Ar IST problēmām saistīto patentu skaita pieaugums
(avots: OCLC FirstSearch database from UMI Company)

2. Domas abstrakciju evolūcijas glosārijs

Sajūta – objektīvās realitātes īpašību atspoguļošana, kas rodas, tām iedarbojoties uz māņu orgāniem un galvas smadzeņu nervu centru uzbudinājums, pasaules uztveres un izzināšanas izejas punkts. *Sajūtu veidi*: tauste, redze, dzirde, vibrācija, oža. Noteiktu sajūtu kvalitatīvo īpašību sauc par modalitāti.

Uztvere – informācijas pieņemšanas un pārveidošanas salikts process, kurš nodrošina objektīvās realitātes atspoguļojumu un orientāciju apkārtējā pasaulē. Uztvere apzīmē: objekta atrašana uztveres laukā, atsevišķu pazīmju atšķiršana objektā, informācijas satura, darbības mērķa atbilstības izdalīšana tajā, uztveres tēla veidošana.

Priekšstats – agrāk uztveramā priekšmeta vai parādības (atmiņas, atcerēšanās) tēls, kā arī tēls, radīts ar produktīvo iztēli; jutekliskā atspoguļojuma augstākā forma uzskatāmi–tēlaino zināšanu veidā.

Jēdziens, saprašana – domāšanas forma, kura atspoguļo priekšmetu un parādību būtiskās īpašības, saites un attiecības. Saprašanas pamata loģiskā funkcija – tā kopīgā izdalīšana, kas sasniedzas ar atraušanu no atsevišķu dotās klases priekšmetu visām īpatnībām. Loģiskā saprašana ir doma, kurā apkopojas un izdalās kādas klases priekšmeti pēc noteiktām kopīgām un kopumā specifiskām pazīmēm.

Spriedums ir prāta akts, kurš realizē attieksmi pret domas saturu un tas ir saistīts ar pārliecību (zināšanu) vai šaubām tās patiesumu vai nepareizumu.

Slēdziens, secinājums – garīga darbība uz normu un slēdzienu pamata, kas piemīt individuālai apziņai. Tie sakrīt ar loģikas likumiem.

Spekulatīva izziņa – loģisko secinājumu virsējas garīgās apceres process, nepiedaloties māņu orgāniem. Spekulatīva izziņa (spekulācija) - idealizēta, fiksēta domāšanas forma, abstrahēta no jutekliskās pieredzes un sabiedriskās praktiskas. Ideālisma vēsturē izdalījās divi spekulācijas veidi: 1) racionālistiskā 2) intuitivistiskā.

Apcere ir izzināšanas jutekliskā pakāpe, realitātes tiešas uztveres process bez loģiskiem secinājumiem.

Apziņa ir objektīvās pasaules subjektīvais tēls, galvenā augstu organizētas matērijas īpašība, ideālās pasaules atveidošanas veids pēc noteiktiem kopīgiem un kopskaitā specifiskām pazīmēm.

Garša ir sajūta, kas rodas dažādu šķīstošu vielu iedarbībā uz garšas receptoriem, kuri mugurkaulniekiem atrodas galvenokārt uz mēles. Pamata garšas sajūtas: 1) rūgtā; 2) saldā; 3) skābā; 4) sāļā. Garša ietekmē apetīti un sagremošanu. Apetītes laika funkcija ir *izsalkums*. Garšas sajūtu stipri ietekmē organisma fizioloģiskais stāvoklis. Tas nozīmē, ka produktu organoleptiskās garšas īpašības, principā, nevar precīzi noteikt. Dažu slimību gadījumā garša vispār var būt sagrozīta. Vairākam mugurkaulnieku par kopējās ķīmiskās izjūtas (garša un oža) orgāniem kalpo sensilles un citi hemoreceptori.

Iztēle (fantāzija) ir psihiskā darbība, iztēles un iedomāto situāciju radīšana, kuras cilvēks nekad kopumā neuztvēra realitātē. Izšķir atveidošanas iztēli un mākslas iztēli.

Intelekts – no lat. *Intellectus* – izzināšana, saprašana, saprāts. Intelekts ir domāšanas, racionālās izzināšanas spēja. *Universums* – pasaule kā viens veselums.

Prāts – domāšanas un saprašanas spēja. Filozofijas vēsturē – saprāts, gars, intelekts. Šo stadiju ilgums laikā ir atkarīgs no uztveramā signāla vai tēla komplicētības.

Percepcija – (no lat. *Perceptio* – iztēlošana, uztvere).

Apercepcija – neskaidra, neapzināta uztvere skaidras apziņas pretstatā.

Perceptuālā uztvere kā informācijas pieņemšanas procesa pamats tās pārstrādei raksturojas ar tādām īpašībām:

- viengabalainība
- saprātība
- selektivitāte
- kontrastainums

Šīs perceptuālās uztveres īpašības nav perceptuālā tēla sākotnējās īpašības, tās veidojas tēla rašanās procesā. Šīm secinājumam ir īpaša nozīme, radot optimālos informācijas attēlošanas līdzekļus IST vadības, informācijas drošības un mākslīgā intelekta sistēmās. Kā ir zināms, perceptuālā tēla veidošanas fizioloģiskais pamats ir dažādu analizatoru izmantošana, ar kuru palīdzību īstenojas signālu – kaitinātāju analīze. Parasti analizators sastāv no receptora, pavadīto nervu ceļu un galvas smadzeņu garozas centra.

Receptora pamata funkcija ir funkcionējošā kaitinātāja enerģijas pārveidošana nervu procesā. Receptora ieeja ir pielāgota noteiktās modalitātes (veida) – gaismas, skaņas u. c. - signālu pieņemšanai. Tomēr receptora izeja sūta signālus, kuras pēc savas dabas ir vienīgas jebkurai nervu sistēmas ieejai. Šīs secinājums ļauj apskatīt receptorus kā informācijas kodēšanas ierīces. Receptoru un centru savstarpējās iedarbības procesā galvas smadzeņu lielo pusložu garozā notiek perceptuālā tēla veidošana. Atkarībā no ienākošā signāla modalitātes (veida) izšķir analizatoru klasificētos veidus. Perceptuālo modeļu radīšanā vislielākā nozīme ir redzes analizatoriem, aiz tiem seko dzirdes un

taustes analizatori. Citu analizatoru piedalīšanos pie tam ir maznozīmīga. Ar analizatoru palīdzību cilvēks var ne tikai sajūst kādu signālu, bet arī atšķirt signālus indikācijas, identifikācijas un klasifikācijas stadijās. Eksperimentāli ir noteikts, ka klasifikācijas sliekšņa vērtība ir proporcionāla kaitinātāja signāla izejas vērtībai.

Diferenciālā jūtīguma sliekšņa vērtība raksturo analizatora maksimālās iespējas. Analizatoru raksturojumi un darbības princips ļauj novērtēt tēla perceptuālo signālu iedarbību uz cilvēka organismu un smadzenēm un noformulēt IST drošības prasības signāliem – kaitinātājiem, kuri var būt vērsti uz cilvēku. Īpašu nozīmi iegūst smadzeņu hiperaktivizācijas procesa modelēšanā un analīzē, jo perceptuālo tēlu gaismas signāli var negatīvi ietekmēt cilvēka – IST operatora veselību. Adaptācijas procesā lielā pakāpē - līdz 108 reizēm mainās cilvēka redzes analizatora jūtīgums. Var izdalīt divas adaptācijas formas: tumšā (pārējā no gaismas tumsā) un gaišo (pretējā pārējā). Adaptācijas laiks ir atkarīgs no tās veida un sastāda no desmitiem minūšu (tumšā) līdz sekundes daļām (gaišā). Perceptuālo analizatoru laika raksturojumi noteicas ar laiku, kas ir nepieciešams, lai cilvēkam radītos sajūta.

3. IST teorijas pamatu apgūšanas metodoloģija

“Pasaulē ir kaut kas daudz svarīgāks, nekā visizcilākie atklājumi: tās ir metodes zināšana, ar ko tie tika izveidoti” (G.Leibnics)

Datorzinātnes psiholoģija. Datorzinātnes psiholoģijas mērķis - ir izpētīt cilvēka intelektuālo uzvedību ar datoru modelēšanas palīdzību. Algoritmam, ietvertam programmā, ir precīzi jākopē cilvēka uzvedība, tam ir ātri un precīzi jāveic tas, ko cilvēki parasti dara lēni. Algoritmam ir jāpamatojas uz tiem pašiem principiem un jāizmanto tās pašas intelektuālo uzvedību noteicošās funkcionālas struktūras, vadības un lēmumu pieņemšanas paradigmām (8. att.), kuras parasti izmanto domājošs cilvēks.

Datorzinātnes filozofija. Datorzinātnes filozofijas mērķis – ir saprast un izveidot tādu datorizētu modeli, kuru parasti sauc par „intelektu”. Pie tam tādām modelim nav obligāti jākopē cilvēka uzvedība, bet tai ir jābūt algoritmizētai. Lietišķās IST jomās (tehniskā redze, robottehnika) „datoriskā filozofija” uzliek mērķi izpētīt

kādi algoritmi var būt izmantoti, lai iegūtu informāciju par objektiem uz attēla pamata darbību plānošanai atkarībā no aktuāliem mērķiem, esošās situācijas, utt.

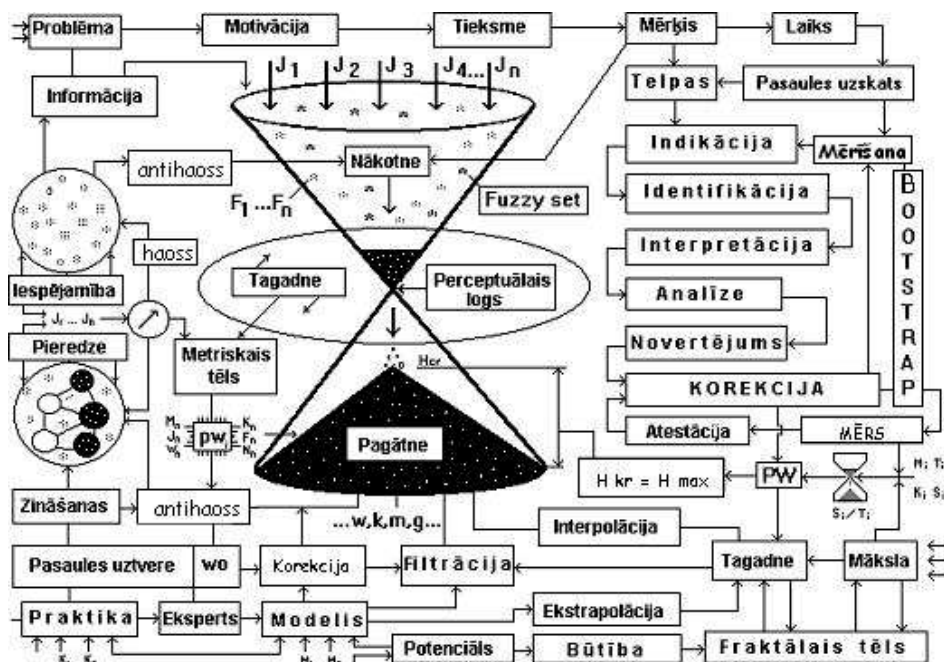
Jaunas metodes informātikā. Tie izvirza mērķi iemācīties programmēt to, ko agrāk cilvēki prāta darīt bez programmēšanas. Interesanti atzīmēt, ka šim virzienam ir “pašiznīcināšanas” raksturs – kā tikai tas gūs panākumus kādas problēmas risināšanā, šī problēma uzreiz pārstāj būt IST problēma.

Neironu tīkli. Izmantojamas visas metodes, kas ir saistīti ar neironu tīkliem, ar ģenētiskiem algoritmiem un semantiskiem tīkliem.

Zināšanu interpretācijas un attēlošanas metodes. saistītas ar labāko interpretācijas un attēlošanas metožu un algoritmu izstrādi, lai tos izmantotu konkrēto IST tehnoloģiju jomās.

Teksta izpratne. Pie šīs grupas pagaidām var pieskaitīt semantiskos tīklus.

Runas izpratne. Tādu izstrādājumu, kuri ir pilnīgi gatavi plašai pielietošanai, pagaidām praktiski nav.



8. att. IST uzbūves vispārinātā paradigma (avots: Moskvins, 1998)

Runas pazīšana. Dažādi runas pazīšanas algoritmi. Dotajā brīdī šīs sistēmas aprobežojas ar runas izteikumu pārveidošanu tekstā.

Tēlu pazīšana. Skanētā teksta pazīšana, cilvēka sejas pazīšana u.c.

Loģiskais secinājums. Tas var neattiekties uz MI, bet MP („mākslīgā prāta”) realizācija nav iespējama bez loģiskā secinājuma.

Koda ģenerācija. Dažādi algoritmi, kuri paši sevi pilnveido. Tādu sistēmu realizācijas mēģinājumi pagaidām izgāzās, taču tas ir visai interesants virziens.

Kognitivistika (cognitive science) – zinātnes nozare, kura pēta cilvēka uzvedību, pieņemot, ka izzināšana var būt aprakstīta informātikas terminos (un efektīvi izveidota ar datoru programmu palīdzību). Vēl piemēri - datorlingvistika ir datoru izmantošanas zinātne cilvēka valodas pētīšanai un apstrādāšanai tekstu saprašanai dabīgā valodā (DV), automātiska tulkošana, anotāciju sastādīšana, informācijas meklēšana, tai skaitā arī internētā utt. [7, 29, 35, 55].

No 40. gadu beigām lielāka universitāšu un rūpniecisko pētniecības laboratoriju skaita pētnieki vērsās pie izaicinoša mērķa: tādu datoru radīšana, kuri darbotos tādā veidā, ka pēc to darbības rezultātiem tos nevarētu atšķirt no cilvēciskā prāta. Pacietīgi kustoties uz priekšu, pētnieki, kuri strādāja IST jomā, atklāja, ka viņi sastopas ar visai komplikētām problēmām, kuri pārkāpj tradicionālās informātikas priekšstatu robežas. Izrādījās, ka pirmām kārtām ir nepieciešams saprast apmācības procesa mehānismus, jutekliskās uztveres un valodas dabu. Noskaidrojās, ka tādu mašīnu radīšanai, kuras imitē cilvēka smadzeņu darbību, vajag saprast, kā darbojas to savstarpēji saistītie miljardi neironu. Un tad daudzie pētnieki secināja, ka, acīmredzams, vissarežģītākā mūsdienu zinātnes problēma ir cilvēka prāta funkcionēšanas procesu izzināšana, nevis vienkārši tā darbības imitācija. Tas tieši aizskāra psiholoģiskās zinātnes fundamentālās teorētiskās problēmas. Īstenībā, pētniekiem ir grūti pat vienoties par viņu pētīšanas objektu – intelekta definīciju. Daži uzskata, ka intelekts ir prasme risināt sarežģītus uzdevumus; citi uzskata to par spēju apmācīties, apkopot un veikt analogijas; pārējie – ka tā ir iespēja savstarpēji iedarboties ar ārējo pasauli sazinoties, uztverot un uztverto apzinoties. Neraugoties uz to, daudzie IST un MI pētnieki tiecas pieņemt mašīnas intelekta testu, kuru piedāvāja 50. gadu sākumā izcils angļu matemātiķis un speciālists

skaitļošanas tehnikā *Alans Tūrings*. Datoru var uzskatīt par saprātīgo, ja tas spēj likt mums ticēt, ka mēs sadarbojamies ar cilvēku, nevis ar mašīnu.

3.1. Mehānikas metodoloģija

Ideja radīt „cilvēka tipa” mašīnas, kuras varētu domāt, kustēties, dzirdēt, runāt un vispār uzvesties kā dzīvs cilvēks, aiziet pagājībā. Vēl senie ēģiptieši un romieši sajūta godbijīgas šausmas no kulta statujām, kuras žestikulēja un izteica pravietojumus (protams, ar priesteru palīdzību). Viduslaiku anāles ir pilni ar stāstiem par automātiem, kuri spēj staigāt un kustēties gandrīz kā to īpašnieki cilvēki. Viduslaikos un pat vēlāk bija izplatītas baumas par to, ka kādam no prātniekiem ir homunkuli (sīkie mākslīgie cilvēciņi) – īstās dzīvās būtnes, kuras spēj sajust. Izcils Šveices ārsts un dabas pētnieks *Teofasts Bombasts fon Hogenheims* (vairāk pazīstams kā *Paracels*) XVI gs. atstājis homunkula izgatavošanas rokasgrāmatu, kurā bija aprakstīta dīvaina procedūra, kas sākās no hermētiski noslēgtas cilvēka spermas ierakšanas zirgu mēšlos. „Mēs būsim kā dievi”, - proklamēja *Paracels*, „Mēs atkārtosim diženāko no dieva brīnumiem – cilvēka radīšanu!” XVIII gs. Pateicoties tehnikas attīstībai, it īpaši pulkstenis mehānismu izstrādāšanai, palielinājās interese pret līdzīgiem izgudrojumiem, kaut gan rezultāti bija ļoti mazi, nekā gribēja *Paracels*. 1736. gadā franču izgudrotājs *Žaks de Vokansons* izgudroja mehānisko fleitistu cilvēka augumā, kurš prāta izpildīt divpadsmit melodijas, šķirstot ar pirkstiem spraugas un pūstot uzgalē, kā īsts muzikants. 1750. gadu vidū austrietis *Fridrihs fon Knauss*, kurš dienēja *Franciska I* galmā, uzkonstruēja mašīnu sēriju, kuri prāta turēt spalvu un varēja rakstīt diezgan garus tekstus. Cits meistars, *Pjers Žaks Drozs* no Šveices, uzcēla īpaši sarežģīto leļļu bērna izmērā: puisi, kurš raksta vēstules un meiteni, kura spēlē klavesīnā. Mehānikas panākumi XIX gs. Stimulēja vēl godkārīgākus nodomus. Tā, 1830. gados angļu matemātiķis *Čarlz Bebbidžs* nodomāja, bet nepabeidza, sarežģīto ciparu kalkulatoru, kuru viņš nosauca par „analītisko mašīnu”. Kā *Č. Bebbidžs* apgalvoja, viņa mašīna principā varēja aprēķināt šaha gājienus. Vēlāk, 1914. g., viena spāņu tehniskā institūta direktors *Leonardo Torresijs Kevedo* patiešām uztaisīja elektromehānisko ierīci, kura spēj nospēlēt vienkāršākos šaha galotnes gandrīz tikpat labi, kā cilvēks.

3.2. Elektronikas metodoloģija

Tomēr tikai pēc otrā pasaules kara parādījās ierīces, kuras, liekas, ir piemērotas slepena mērķa sasniegšanai – prātīgās uzvedības modelēšanai: tās bija elektroniskās ciparu skaitļošanas mašīnas. „Elektroniskās smadzenes”, kā tajā laikā sajūsmināti sauca datoru, satrieca 1952. gadā ASV teleskatītājus, precīzi pareģojot prezidentu vēlēšanas rezultātus dažas stundas pirms beigu doto iegūšanas. Šis datora „varoņdarbs” tikai apstiprināja secinājumu, kuru tajā laikā izdarīja daudzi zinātnieki: iestāsies tā diena, kad automātiskie aprēķinātāji, kuri tik ātri, nenogurstoši un nekļūdīgi veic automātiskās darbības, varēs imitēt neskaitļošanas procesus, kuri ir raksturīgi cilvēka domāšanai, tai skaitā arī uztveršanu un apmācību, tēlu atšķiršanu, mūsdienu runas un rakstīšanas saprašanu, lēmumu pieņemšanu nenoteiktajās situācijās, kad ne visi fakti ir zināmi. Tādā veidā „aizmuguriski” bija formulēts, īpašs „sociālais pasūtījums” psiholoģijai, stimulējot dažādas zinātnes nozares. Daudzi datoru izgudrotāji un pirmie programmētāji izkļaidējās, sastādot programmas pavisam netehniskām nodarbībām, tādas kā mūzikas sarakstīšana, grūti atrisināmo uzdevumu atrisināšana un spēles, pirmajā vietā te izrādījās šahs un dambretes. Daži romantiski noskaņoti programmētāji pat lika savām mašīnām rakstīt mīlestības vēstules. 50. gadu beigās visas šīs izkļaidēšanas apkopojās jaunajā vairāk vai mazāk patstāvīgajā informātikas nozarē, kura ieguva nosaukumu „mākslīgais intelekts”.

Pētījumi MI jomā, kuri no sākuma bija koncentrēti dažos ASV universitāšu centros, Masačūsetsas tehnoloģiskajā institūtā, *Karnegi Tehnoloģiskajā institūtā* Pitsburgā, *Stenforda universitātē*, šobrīd tiek vadīti daudzās citās ASV un citu valstu universitātēs un korporācijās. Kopumā, MI pētniekus, kuri strādā domājošo mašīnu radīšanā, var iedalīt divās grupās. Vienu grupu interesē tīrā zinātne un priekš viņiem dators ir tikai instruments, kas nodrošina domāšanas procesu teorijas eksperimentālās pārbaudes iespēju. Otrās grupas intereses ir tehnikas jomā: viņi tiecas paplašināt datoru pielietojanas sfēru un atvieglot to lietošanu. Daudzi otrās grupas pārstāvji maz rūpējas par domāšanas mehānisma noskaidrošanu, viņi uzskata, ka priekš viņu darba tas ir tikpat nepieciešams, kā putnu lidošanas un lidmašīnu būvēšanas pētīšana. Mūsdienu laikā, tomēr, izrādījās, ka kā zinātniskie, tā arī tehniskie meklējumi satikās ar daudz

nopietnākām problēmām, nekā likās pirmajiem entuziastiem. Pirmajā laikā daudzi MI pionieri ticēja, ka pēc kāda gadu desmitnieka mašīnas iegūs augstākos cilvēciskos talantus. Uzskatīja, ka pārvarot „elektroniskās bērnības” periodu un apmācoties visas pasaules bibliotēkās, viltīgi gudri datori, pateicoties ātrdarbībai, precizitātei un netraucētai atmiņai tie pakāpeniski pārspēs savus radītājus – cilvēkus. Tagad maz cilvēku runā par to, bet ja arī runā, tad neuzskata, ka šādi brīnumi ir ne aiz kalniem. Visas savas īsas vēstures garumā pētnieki MI jomā visu laiku atradās informātikas priekšējā malā.

Daudzi tagad vienkārši izstrādājumi, tai skaitā arī pilnveidotās programmēšanas sistēmas, tekstu redaktori un tēlu atšķiršanas programmas, lielā mērā tiek apskatīti MI darbos. Īsiem vārdiem sakot, jaunās idejas un MI izstrādājumi vienmēr pievelc to cilvēku uzmanību, kuri tiecas paplašināt datoru iespējas un pielietošanas jomas, padarīt tos par vairāk „draudzīgiem”, tas ir vairāk līdzīgiem citiem saprātīgiem palīgiem un aktīviem padomdevējiem, nekā tie pedantiskie un pastulbi elektroniskie vergi, par kuriem tie bija visu laiku. Neskatoties uz daudzsološas perspektīvas, nevienu no līdz šim izstrādātām MI programmām nevar nosaukt par „saprātīgo” parastajā šī vārda nozīmē. Tas ir izskaidrojams ar to, ka tās visas ir šauri specializētas; vissarežģītākās ekspertu sistēmas pēc savām iespējām drīzāk atgādina dresētas vai mehāniskās lelles, nekā cilvēku ar viņa lokano prātu un plašu redzesloku. Pat starp MI pētniekiem tagad daudzi šaubās, ka vairākums tādu izstrādājumu dos būtisko labumu. Daudz MI kritiķu uzskata, ka tāda veida ierobežojumus vispār nevar pārvarēt. Pie tādiem skeptiķiem pieder arī *Hūberts Dreifuss*, Berklija Kalifornijas universitātes filozofijas profesors. Viņš uzskata, ka īsto prātu nevar atdalīt no tā cilvēciskā pamata, kurš ir ieslēgts cilvēka organismā. „Ciparu dators nav cilvēks” – saka *Dreifuss*. „Datoram nav ne ķermeņa, ne emociju, ne vajadzību. Viņam nav sociālās orientācijas, kuru var iegūt tikai dzīvojot sabiedrībā, bet tieši tā padara uzvedību par prātīgu. Es negribu teikt, ka datori nevar būt saprātīgi. Bet ciparu datori, kuri ir ieprogrammēti ar faktiem un likumiem no mūsu, cilvēciskās dzīves, patiešām nevar kļūt saprātīgi. Tāpēc MI nav iespējams tajā veidā, kā mēs to iedomājamies”.

3.3. Datormodelēšanas metodoloģija

Mēģinājumi uzbūvēt mašīnas, kuras spēj saprātīgi uzvesties, lielā mērā ir sajūsmināti ar *Masačūsetsas tehnoloģiskā institūta* (MTU, ASV) profesora *Norberta Vīnera* idejām – vienas no ievērojamām personībām Amerikas intelektuālajā vēsturē. Izņemot matemātiku, viņam bija plašas zināšanas citās jomās, ieskaitot neiropsiholoģiju, medicīnu, fiziku un elektroniku. Vīners bija pārliecināts, ka vairāk perspektīvie zinātniskie pētījumi tā sauktajās pierobežas jomās, kuras nevar konkrēti pieskaitīt pie kādas konkrētas disciplīnas. Tie atrodas kaut kur zinātņu saskares vietā, tāpēc pret tiem neattiecas tik stingri. „Ja grūtībām kādas psiholoģijas problēmas risināšanā ir matemātisks raksturs, viņš skaidroja, tad desmit psiholoģi, kuri neko nesaprot matemātikā, pastumsies ne tālāk, par vienu, kurš tikpat neko nesaprot”. Vīneram un viņa līdzstrādniekam *Džuliānam Bigelovam* (*Julian Bigelow, 1913-2003*) pieder „atgriezeniskās saites” principa izstrādāšana, kurš bija veiksmīgi pielietots jaunā ieroča izstrādāšanā ar radiolokācijas pielietošanu. Atgriezeniskās saites princips ir tās informācijas izmantošana, kura nonāk no apkārtējās pasaules, mašīnas uzvedības maiņai. Vīnera un Bigelova izstrādāto iestatīšanas sistēmu pamatā ir smalkas matemātiskās metodes; ja gadījumā atstaroti no lidmašīnas radiolokācijas signāli pat nedaudz mainījās, tie attiecīgi mainīja ieroču iestatījumu, tas ir, ievērojot lidmašīnas mēģinājumu novirzīties no kursa, tie uzreiz aprēķināja tās tālāko ceļu un virzīja ieročus tā, lai lidmašīnu un ieroču trajektorijas krustotos. Turpmāk, pamatojoties uz atgriezeniskās saites principu Vīners izstrādāja mašīnas un cilvēka prāta teorijas. Viņš pierādīja, ka tieši pateicoties atgriezeniskai saitei, viss dzīvais pielāgojas apkārtējai videi un sasniedz savus mērķus. „Visām mašīnām, kuras pretendē un saprātīgumu, viņš rakstīja, ir jābūt spējai tiekties pēc saviem mērķiem un pielāgoties, t.i. apmācīties”. Paša radītai zinātnei Vīners dod nosaukumu kibernetika, kas grieķu valodā apzīmē „,stūrmanis”. Ir jāatzīmē, ka „atgriezeniskās saites” principu, kuru piedāvāja Vīners, bija kaut kādā mērā paredzējis izcils krievu fiziologs *Ivans Sečenovs* (1829-1905) „centrālās bremzēšanas” parādībā „Galvas smadzeņu refleksos” (1863 g.) un tas tika apskatīts kā nervu sistēmas darbības mehānisms un kurš bija par pamatu daudzām patvaļīgās uzvedības modeļiem psiholoģijā.

3.4. Neiromodeļu metodoloģija

Uz šo laiku arī citi zinātnieki sāka saprast, ka skaitļošanas mašīnu rādītājiem ir jāpamāca bioloģija. Viņu vidū bija neurofiziologs un dzejnieks-amatieris *Vorens Makkalohs*, kuram, tāpat kā *Vīneram*, bija filozofisks domāšanas veids un plašs interešu loks. 1942. gadā *Makkalohs*, piedaloties zinātniskajā konferencē Ņujorkā, sadzirdēja referātu viena no *Vīnera* līdzstrādniekiem par atgriezeniskās saites mehānismiem bioloģijā. Idejas, izteiktas referātā, bija līdzīgas paša *Makkaloha* idejām par galvas smadzeņu darbību. Nākošā gada gaitā *Makkalohs*, līdzautorībā ar savu 18-gadīgo protežē, izcilu matemātiķi *Valteru Pītsu*, izstrādāja galvas smadzeņu darbības teoriju. Šī teorija arī bija tas pamats, uz kura veidojās plaši izplatīts viedoklis, ka datora un smadzeņu funkcijas lielā mērā ir līdzīgas. Daļēji izejot no iepriekšējiem neironu pētījumiem (pamata aktīvām šūnām, kuras veido dzīvnieku nervu sistēmu), kurus veica *Makkalohs*, viņi ar *Pītsu* [58] izvirzīja hipotēzi, ka neironus var vienkāršoti apskatīt kā ierīces, kuras operē ar bināriem skaitļiem. Binārie skaitļi, kuri sastāv no cipariem viens un nulle, ir vienas matemātikas loģikas sistēmas darba instruments.

Angļu matemātiķis *Džordžs Buls (1815-1864)* 1854. gadā parādīja, ka loģiskos apgalvojumus var sastādīt ar koda palīdzību vieninieku un nulļu veidā, kur vieninieks atbilst patiesam apgalvojumam, bet nulle – neīstam, un pēc tam ar to var operēt kā ar parastajiem cipariem. XX gs. 30. gados informātikas pionieri, it īpaši amerikāņu zinātnieks *Klods Šennons*, saprata, ka binārie vieninieks un nulle pilnīgi atbilst diviem elektriskās virknes stāvokļiem (ieslēgts-izslēgts), tāpēc binārā sistēma ideāli der elektroniskām skaitļošanas ierīcēm. *Makkalohs* un *Pīts* piedāvāja tīkla konstrukciju no elektroniskiem „neironiem” un parādīja, ka tāds tīkls var izpildīt praktiski jebkuras iedomājamās loģiskās un ciparu operācijas. Turpmāk viņi iedomājās, ka tāds tīkls var arī apmācīties, atšķirt tēlus, apkopot, t.i. tai ir visas intelekta pazīmes. *Makkaloha* un *Pītsu* teorijas kopā ar *Vīnera* grāmatām izsauca milzīgo interesi pret saprātīgām mašīnām. 40-60. gados jo vairāk kibernetiķu no universitātēm un privātām firmām ieslēdzās laboratorijās un darbnīcās, intensīvi strādājot ar smadzeņu funkcionēšanas teoriju un metodiski pielodējot neironu modeļu elektroniskos komponentus. No šīs kibernetiskās, vai neiromodeļu pieejas pie mašīnu prāts drīz izveidojās tā sauktā

„kāpjošā metode” – kustība no primitīvo būtņu vienkāršiem nervu sistēmas analogiem, kuriem ir neliels neironu daudzums, uz sarežģītāko cilvēka nervu sistēmu un pat augstāk. Beigu mērķis bija „adaptīvā tīkla”, „pašorganizējošās sistēmas” vai „apmācošās mašīnas” radīšana. Visus šos nosaukumus dažādi pētnieki izmantoja tādu ierīču apzīmēšanai, kuras spēj vērot apkārtējo stāvokli un ar atgriezeniskās saites palīdzību mainīt savu uzvedību, pilnīgi atbilstot tajā laika dominējošai biheivoristiskai psiholoģijas skolai, t.i. uzvesties tāpat, kā uzvedās dzīvie organismi.

Taču ne visos gadījumos ir iespējama analogija ar dzīviem organismiem. Kā vienreiz pamanīja *Vorens Makkalohs* un viņa līdzstrādnieks *Maikls Arbibs*, „ja pavasarī jums sagribējās atrast sirdsmīlo, nav vērts ņemt amēbu un gaidīt, kamēr viņa evolucionēs”. Bet šeit ir darīšanās ne tikai laikā. Pamata grūtības, ar kurām saskrējās „kāpjošā metode” savas eksistences ausmā, bija augsta elektronisko elementu izmaksa. Pārāk dārgs izrādījās pat skudras nervu sistēmas modelis, kurš sastāv no 20 tūkstošiem neironu, jau nerunājot par cilvēka nervu sistēmu, kura ietver sevī ap 100 miljardus neironu. Pat vispilnveidotākie kibernetiskie modeļi saturēja tikai dažus neironu simtus. Tik ierobežotas iespējas daudziem tā laika pētniekiem atņēma drosmi. Vienu no tiem, kuru nemaz neizbiedēja grūtības, bija *Frenks Rozenblats* [67], kura darbi, kā izrādījās, atbildēja visslepenākajiem kibernetiķu tieksmēm. 1958. gada vidū viņš piedāvāja elektroniskās ierīces modeli, kuru viņš nosauca par perceptronu, un kurai vajadzētu imitēt cilvēka domāšanas procesus. Perceptronam bija jānodod signālus no „acs”, kurš bija sastādīts no fotoelementiem, elektromehānisko atmiņas šūnu blokos, kuri novērtēja elektrisko signālu relatīvo lielumu. Šīs šūnas savienojās savā starpā nejaušā veidā saskaņā ar tajā laikā dominējošo teoriju, saskanīgi ar ko smadzenes uztver jauno informāciju un reaģē uz to caur nejaušo saišu sistēmu caur neironiem.

Pēc diviem gadiem bija nodemonstrēta pirmā funkcionējošā mašīna „Mark1”, kura varēja iemācīties atšķirt dažus burtus, uzrakstītus uz kartiņām, kuras pienesa klāt tās „acīm”, kuri atgādināja kinokameras. *Rozenblata perceptrons* izrādījās mākslīgā intelekta radīšanas „kāpjošās” vai neiromodeļu metodes visaugstākais sasniegums. Lai iemācītu perceptronu veidot minējumus izejas priekšnosacījumu pamatā, tajā bija paredzēts īpašs autonomā darba vai „pašprogrammēšanas” elementārais veids. Kāda burtu atpazīšanā, vieni tā elementi vai elementu grupas izradās daudz būtiskāki, nekā

citi. Perceptrons varēja iemācīties izdalīt tādas burtu raksturīgas īpašības pusautomātiski, savā ziņā ar mēģinājumu un kļūdu metodi, kas atgādina apmācības procesu. Taču perceptrona iespējas bija ierobežotas: mašīna nevarēja droši pazīt daļēji aizsegtos burtus, kā arī cita izmēra vai zīmējuma burtus, nekā tos, kuri bija izmantoti tās apmācības etapā. Tā sauktās „lejupejošās metodes” vadošie pārstāvji specializējās, atšķirībā no „kāpjošās metodes” pārstāvjiem, tādu uzdevumu risināšanas programmu sastādīšanā vispārīgā uzdevuma ciparu datoriem, kuras prasa no cilvēka ievērojamo intelektu, piemēram, spēlei šahā vai matemātisko pierādījumu meklēšanā. Pie „lejupejošās metodes” aizstāvju skaita pieskaitījās *Marvins Minskis* un *Seimurs Peiperts* [61], Masačūsetsas tehnoloģiskā institūta profesori. *M. Minskis* [59-61] uzsāka savu MI pētnieka karjeru kā „kāpjošās metodes” piekritējs un 1951. gadā uzbūvēja apmācošo tīklu uz vakuuma elektroniskām lampām. Taču drīz, perceptrona radīšanas laikā, viņš pārcēlās pretējā nometnē. Līdzautorībā ar dienvīdāfrikāņu matemātiķu Peipertu, ar kuru viņu iepazīstināja Makkalohs, viņš uzrakstīja grāmatu „Perceptroni”, kur bija matemātiski pierādīts, ka perceptroni, līdzīgi Rozenblata perceptroniem, principiāli nevar izpildīt daudzas no tām funkcijām, kurus viņiem pareģoja *Rozenblats*. Bet *Minskis* apgalvoja, ka, nerunājot par strādājošiem pēc diktāta mašīnrakstītājām, kustīgiem robotiem vai mašīnām, kuras spēj lasīt, klausīties un saprast izlasīto un sadzirdēto, perceptroni nekad neiegūs pat spēju pazīt priekšmetu, daļēji aizsegtu ar citu. Skatoties uz izlīdušo aiz krēsla kaķa asti, līdzīga mašīna nekad nevarēs saprast, ko viņa redz. Nevar pateikt, ka šis 1969. gada darbs likvidēja kibernetiku. Tas tikai pārvietoja aspirantu interesi un ASV vadošo organizāciju subsīdijas, kuri tradicionāli finansē MI pētījumus, uz citu pētījumu virzienu – „lejupejošo metodi”. Interese pret kibernetiku pēdējā laikā atjaunojās, jo „lejupejošās metodes” piekritēji saskārās ar nepārvaramām grūtībām. *M. Minskis* publiski izteica nožēlu, ka viņa uzstāšanās nodarīja zaudējumus perceptronu koncepcijai, paziņojot, ka, saskaņā ar viņa šodienas priekšstatiem, lai notiktu reāls pārrāvums uz priekšu saprātīgo mašīnu radīšanā, būs nepieciešama ierīce, kura lielā mērā ir līdzīga perceptronam. Bet pamatā MI kļuva par lejupejošās pieejas sinonīmu, kas izpaužas jo sarežģītāko datorprogrammu sastādīšanā, kuras modelē cilvēka smadzeņu sarežģīto darbību.

4. Kā domā cilvēks?

Smadzenes ir orgāns, kura darbības rezultāts dažreiz noliedz tā eksistences faktu (novērojums)

Cilvēks domā ar tēliem. Ir arī definīcija, ka doma ir tēlu aktivēšana atmiņā, kad tos stimulē māņu orgāni vai arī šī aktivācija tiek pārnesta starp tēliem (no viena tēla uz otru). Jo tēli ir vairāk līdzīgi viens otram, jo labāk tiek pārnesta aktivācija. Šī līdzība (*similarity, dynamic similarity, geometrical similarity*) var būt fiziskā un apzinīga. Fiziskā – kad tēli ir fiziski līdzīgi (viens durvis tēls aktivizē visus durvju tēlus). Apzinīga – kad tēli bija aktīvi pagātnē, vienā laika sprīdī (durvju tēls aktivizē cilvēka tēlu, jo viņš atradās blakus durvīm tajā laika momentā).

Aktivācijas pārnese notiek labāk uz to pusi, kur atrodas tēls, kurš sagādā mums vislielāko prieku. Prāts prognozē uz labvēlīgo pusi. Taču nav iespējams pilnīgi noskaidrot prāta darbības mehānismu, ja nav tā modeļa pilnā apraksta. Šo problēmu nevar risināt pakāpeniski. Daudzi zinātnieki uzskata, ka sadalot prāta algoritma uzdevumu vairākās daļās (apkārtējas vides novērtējums, mērķu hierarhijas uzcelšana, atsevišķā mērķa sasniegšanas progress, refleksija utt.) un risinot tās, viņi varēs izveidot tā modeli. Bet viņi ļoti ātri sāk saprast, ka kaut kā trūkst, daudzas funkcijas un procesi traucē viens otram, daudzas lietas vienkārši kļūst liekas kopējā skatā utt. Galu galā prātīguma efekts nav sasniegts.

Iznāk paradoksālais loks. Lai saprastu prāta algoritmu, ir nepieciešams uzcelt tā pilno modeli (un kļūt par „Radītāju”, veicot „Virsprāta” funkcijas). Bet lai izveidotu prāta modeli, ir nepieciešams to saprast. Un sanāk, ka cilvēki intuitīvi jūt, kā darbojas prāts, taču viņi nevar šo darbību aprakstīt, tāpēc, ka viņi nevar aprakstīt tā pilno modeli. Prāts (*smadzenes*) atceras vispilgtākos pārdzīvojumus un domā par tiem. Var pievērst uzmanību tam, ka prāts domā nevis spilgtu pārdzīvojumu virzienā, bet labu emociju (baudas) virzienā. Atbrīvošana no potenciālās nepatikas ir arī patika. Jo aktīvāks ir tēls, jo lielāka uzmanība ir tam pievērsta, jo vairāk tas ir apjēdzams. Tēla aktivitātes mērs ir proporcionāls tā apjēgšanas mēram. Apziņas centrā atrodas visaktīvākā atmiņas daļa. Vislielākās aktivācijas pārņemšana starp tēliem ir apzinīga un ielāgojas atmiņa vislabāk. Tomēr, ja kopējais aktivācijas līmenis nav augsts, process var likties neapzināts (kā miegs). Būtībā šis process ir arī apzinīgs, taču zems

aktivācijas līmenis neļauj nākotnē to atcerēties, jo domas tiecas tikai pēc maksimālās aktivācijas. Tāpēc mēs varam aizmirst to laika intervālu, kad process bija neapzinīgs. Aktivācijas pārņemšanai starp tēliem ir dinamisks raksturs un tā var notikt vienlaikus dažādās vietās. Saprātīgas sistēmas iekšā aktivācijas pārņemšanai ir rimstošs raksturs. Tāpēc rodas „samaņa sašaurināšana” un iekļūšana transa stāvoklī pēc minimālas iedarbības uz māņu orgāniem. Piemēram, kad cilvēku ievieto izolētā melnā istabā.

5. IST aktuālie uzdevumi

IST problēmu risinājumu pamatā ir datora programmu un aparātu līdzekļu izveidošana ar augstu intelekta līmeni [1, 3, 4, 14, 15].

Dabīgā valoda. Pamata mērķis – radīt tā datora modeli, kā cilvēki izmācās un izmanto savu dzimto valodu; izveidot datorprogrammu, kura varētu mācīties un izmantot valodu tāpat, kā to dara cilvēki. Tā kā valodas izpratne ir cieši saistīta ar zināšanu iztēlojumu un vispār ar domāšanu, tad šīs problēmas risinājums ir nepieciešams un pietiekošs intelektuālās sistēmas izveidošanai.

Uzdevumu risināšana un meklēšana. Meklēšana vai risinājumu izveidošana norādītajiem uzdevumiem. Atšķirības pazīme ir pieeja: mērķis = uzdevuma risinājums, uzdevums = iespējamo risinājumu kopa. Process ierobežojas ar pareiza (vai labākā) risinājuma sameklēšanu.

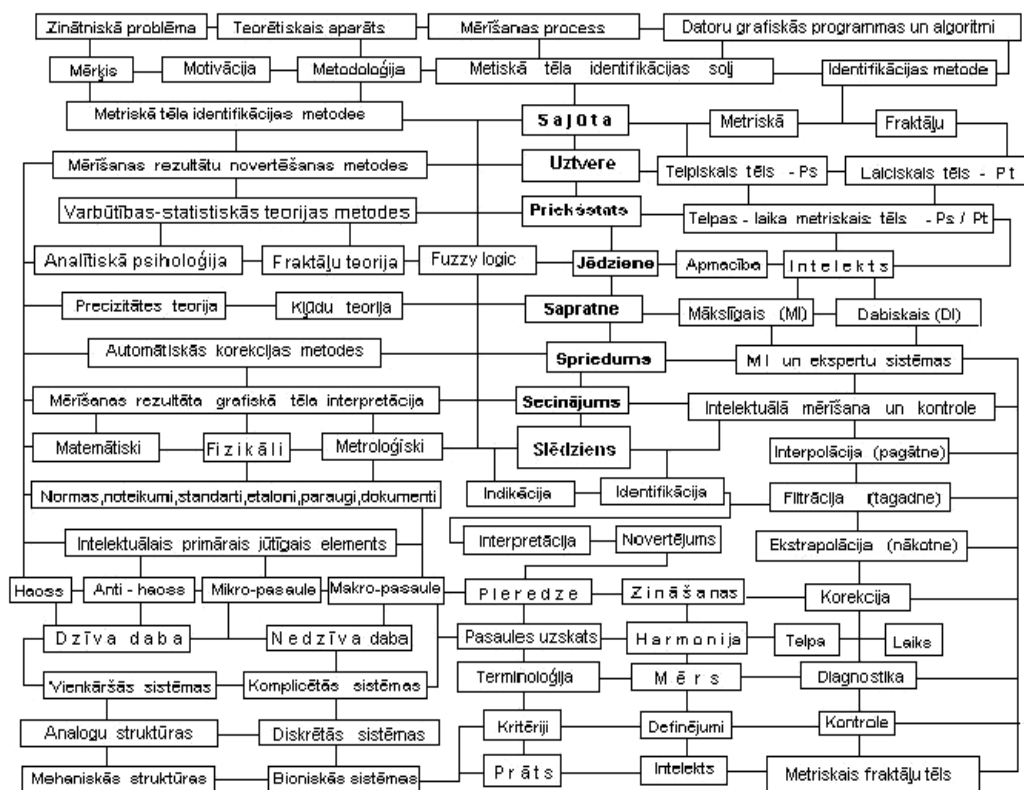
Zināšanu iztēlojums ir IST joma, kura ir saistīta ar simboliskiem formālismiem, lai iztēlotu intelektuālās sistēmas sniegtas zināšanas, kā arī ar doto struktūru pētīšanu, kuri tiek izmantoti šo formālismu realizācijai.

Taču nav iespējams aprobežoties tikai ar statistiskiem iztēlojumiem – nepieciešams ņemt vērā, cik svarīga ir šo zināšanu izmaiņas dinamika noteiktā mērķa sasniegšanai. Daudzos gadījumos mērķis ir jauno zināšanu gūšana spriedumu formā (*reasoning*). Saskaņā ar intelektuālo sistēmu attīstības paradigmu (9. att.) zināšanu iztēlojums un spriedumi ir nepieciešams nosacījums intelektuālās sistēmas izveidošanai.

Jāatzīmē, ka attiecības starp tādām dažādām MI „saprāta” procesora īpašībām kā „sajūta”, „uztvere”, „sapratne”, „priekšstats”, „jēdziens”, „spriedums”, „secinājums”, „slēdziens” abstraktas informācijas evolūcijas gaitā uz konkrētām zināšanām raksturo sapratīgas sistēmas „domas” vienotību un realizējas „īpašību procesoru” vienotas struktūras veidā [102, 103].

Tādu gudru, smalki domājošu procesoru pagaidām nav. Saskaņā ar piedāvāto paradigmu šīs sapratīgas „īpašību procesoru” attiecības paliek invariantās visām intelektuālām struktūrām, „organismiem”, lai cik atšķirīgi tie būtu savā starpā pēc elementu un sastāvdaļu uzbūves. IST uzdevumi reālās sistēmās un tehnoloģijās nereti pārklājas, pie tam konkrēta gadījumā intelektuālai sistēmai var būt vairāku uzdevumu pazīmes. Kā IST uzdevumu piemēri var būt minēti:

- Mākslīgā intelekta (MI) sistēmas;
- Ekspertu sistēmas (ES), kas ļauj izmantot ekspertu zināšanas lai simulētu domāšanu;
- Neironu tīkli (NT) ;
- Faziloģikas sistēmas (FLS) („fuzzily” - neskaidri)
- Ģenētiskie algoritmi (GA);
- Inteliģentie aģenti (IA);
- Nanotehnoloģijas (NT);
- Datu apstrādes transakciju sistēmas (DATS). Transakciju sistēmu funkcija – datu par darījumu operācijām vākšana un fiksēšana;
- Vadības informācijas sistēmas (VIS) – fiksēta formāta informācijas sagatavošana ražošanas un tehnoloģisko procesu vadībai;
- Lēmumu atbalsta sistēmas (LAS) – sistematizētas informācijas sagatavošana augstāka līmeņa vadībai ;
- Biroju informācijas sistēmas (BIS) – informācijas apmaiņas organizēšana starp informācijas lietotājiem, uzņēmuma darbiniekiem un lietvedībai;
- Personīgās informācijas sistēmas (PIS) – lietošanai vienam cilvēkam (piemēram, ģimenes ārstam, advokātam, uzņēmējam) ;
- Darba grupu informācijas sistēmas (DGIS) – lai veicinātu grupas darba produktivitāti ;



9. att. **Intelektuālo sistēmu attīstības paradigma** (avots: Moskvins, 1996)

- Apmācība. Apmācības spēja ir viena no intelekta būtiskajām īpašībām. Jebkurš uzdevums var būt aprobežots ar apmācību, plašākā ziņā – tai skaitā arī bez „skolotāja”;
- Tehniskā redze (TR). Redzes tēlu interpretācijas un saprašanas modeļu izveide. Ja saprast interpretāciju plaši, tad pilnvērtīgai interpretācijai ir nepieciešams intelekts;
- Roboti. Šīs ierīces, kuras var pārvietoties reālā fiziskā pasaulē un/vai manipulēt ar citiem reālās pasaules objektiem.

Visos iepriekšminētajos uzdevumos netieši ir paredzama cilvēka klātbūtne: problēmas formulējumā atslēgas vārds ir „pareizs” (pareiza valodas izpratne, pareizs uzdevuma risinājums, pareizo zināšanu iegūšana spriedumu vai apmācības rezultātā, tehniskas redzes tēlu pareizā interpretācija utt.). Šeit roboti ir izņēmums, tie ir drīzāk

pielikums, nekā IST pētīšanas joma. Tā kā kritēriju „pareizs” nevar formalizēt (vai apzināti tas netiek iekļauts formālisma ietvaros), tad visiem šiem virzieniem ir jābūt saskaņotiem ar atbilstošām heuristikām. Bet kas tad ir intelekts?

6. Saprātīgas domas akts

Mulķis cenšas apgriezt pasauli kājām gaisā, gudrinieks tiecas to kaut nedaudz uzlabot. Un tikai ģēnijs dara gandrīz neiespējamo – cenšas likt pasauli mierā (M. Geny)

Ribonukleīnskābes molekulas (RNS), kuras pašas sevi kopē mēģenēs, rada, kā saprātīga ideja var īstenoties pavisam „neprātīgu” molekulu vidē. Šīm sintezējošām molekulām ir nervu un muskuļu nepietiekamība, taču tās pietiekami labi attīstās, lai nodrošinātu savu reprodukciju. Mainīgums un selektīvā izvēle nosaka katras molekulas noformētu uzvedību, kas paliek nemainīga visas dzīves ilgumā. RNS atsevišķas molekulas nepielāgojas, bet baktērijas prot to darīt. Konkurence atlasīja noturīgas baktēriju ģintis, kuras var pielāgoties, piemēram, regulējot to savienošanos ar fermentiem, lai piekļūtu barībai. Šie adaptācijas mehānismi nodibinājās patstāvīgi: barības molekulas transportē ģenētiskās atslēgas, kā termostata aukstās gaisa plūsmas. Daži baktēriju tipi arī izmanto empīriskās adaptācijas primitīvo formu ar iespējamo variantu izskatīšanas metodi, tas ir ar kļūdu un mēģinājumu metodi. Šīs ģints baktērijām ir tendence pārvietoties pa taisni un tiem pietiek atmiņas sava tekoša stāvokļa paškontrolei un adaptācijas procesa pašanalīzei. Ja baktērijas jūt, ka apstākļi uzlabojas, tās turpina pārvietošanos pa taisni. Ja tās jūt, ka apstākļi pasliktinās, tad tās apstājas, apgriežas apkārt un pārvietojas nejaunos, haotiskos virzienos. Baktērijas it kā pēc taustes pārbauda iespējamu virzienu drošumu pārvietošanās laikā un izvēlas labvēlīgus virzienus, atsakoties no nelabvēlīgiem. Pateicoties tādām saprātīgi-ceļojošai pārvietošanai uz barojošas vides molekulu koncentrācijas zonām, baktērijas izdzīvo, dzīvo, attīstās, zied un plaukst. Dažu tārpju „prāta spēju” analīze parāda to spēju mācīties. Vienkāršākie tārpji spēj izvēlēties pareizu ceļu vienkāršā T- veidīgā šūnā. Tie mēģina nogriezties pa kreisi un pa labi, uzmanīgi izvēlas labāko uzvedības formu, formē savdabīgus pieradumus – un tas viss noved pie vislabākā rezultāta. Tāda tārpju

lietderīgi-saprātīga uzvedība ir uzvedības stratēģijas izvēles veids, kas ir noteikts ar labāka rezultāta sasniegšanu no visām iespējamām. Tādu stratēģiju psihoetologi sauc par „rezultāta likumu”. Tādā veidā, tārpu gēni atražoja tārpu ar uzvedības tipu, kas attīstās. Tomēr tārpi, kuri ir apmācīti orientēties T-šūnās (vai pat Skinnera baloži, kas ir apmācīti knābāt graudus zaļās lampas gaismā) neuzrāda nekādu refleksiņas domas pazīmju, kuru mēs saistām ar prātu.

Organismi, kas pielāgojas pēc parastā rezultāta likuma, apmācas ar mēģinājumu un kļūdu metodi. Bet, mainot un izvēloties faktisko uzvedību, tie neapdomā savu rīcību, nedomā uz priekšu un nepieņem nekādu lēmumu. Tomēr dabiskā atlase biežāk dod priekšroku organismiem, kuri ir spējīgi domāt. Domāšana nav brīnišķīgo dabas brīnumu maģiskā joma. Attīstītie gēni var uzlabot dzīvnieku domāšanas spējas uz iekšēji nanomolekulāro modeļu radīšanas pamata, kas apstiprina daudzu zinātnisko darbu rezultātus. Reizēm dzīvnieki spēj tēlaini iztēloties dažādas rīcības un sekas, izvairoties no rīcībām, kuras šķiet bīstamas un izpilda rīcības, kas tiem šķiet visizdevīgākās un visdrošākās. Tas ir, dzīvnieki risina savdabīgu optimizācijas uzdevumu izejošo doto apstākļos. Iekšējo, molekulāro nanomodeļu darbaspējīguma pārbaude samazina risku un atvieglo modeļu izmēģinājumus ārējā pasaulē. Rezultāta likums var mainīt molekulāros modeļus. Tā kā gēni var nodrošināt uzvedības attīstību, tad arī tie var nodrošināt mentālo modeļu attīstību. Adaptīvie organismi var mainīt pašu modeļus to versiju priekšrokas virzienā, kuri labāk pierāda šo modeļu vadīšanas izdevīgus principus. Ir zināms, ka kaut kā pārbaudei vislabāk pārbaudīt šo lietu darbā. Modeļiem nav jābūt instinktīviem, tās var attīstīties atsevišķās dzīves evolūcijas procesā. Taču nerunājošie dzīvnieki nenodod mums savu jauno evolūcijas zināšanu un saprašanu. Zināšanas izzūd ar smadzenēm, kur sintezēja šīs zināšanas tāpēc, ka apmācošie intelektuālie modeļi neiezīmējas gēnos. Tomēr pat nerunājošie dzīvnieki var papildināt viens otru, bagātinot ģenētisko bāzi.

Piemēram, viens pērtiķis Japānā izgudroja tādu ūdens izmantošanas veidu, lai atdalītu graudus no smiltīm. Citi pērtiķi ātri iemācījās darīt to pašu. Analogiski notiek arī cilvēku kultūrā, ar tās valodu un ainaviskām gleznām. Šedevri, kā pasaules līmeņa darba piemēru paraugi, var pārdzīvot to radītājus un izplatās pa visu pasauli. Vēl augstākā līmenī – prāta līmenī, var sastādīt vērtējumu evolūcijas standartus, lai

salīdzinātu atsevišķu komponentu atbilstību kopējā modeļa struktūrai, piemēram, perceptuālo un fraktālu tēlu atbilstības veidā. Tas dos iespēju padarīt šos vērtējumus pietiekoši drošus kaut kādu tālāko atbildīgo rīcību pārvadīšanai. Tādā veidā prāts izvēlas uzvešanās noteikumus no sava satura, ieskaitot izvēles noteikumus. Spriedumu noteikumi ir zināšanu filtrācijas algoritms, kuras filtrācijas un atlasas rezultātā attīstās un uzlabojas. Kā uzvedības modeli, tā arī zināšanas standarti attīstās saskaņā ar mērķi. Tas, kas ir nozīmīgs, būtisks un derīgs, tiek novērtēts augstāk par pamata standartu un sāk likties lietderīgs un motivēts. Tad tas paliek par pašmērķi. Mērķtiecība un lietderība kļūst par rīcības pamatprincipu. Turpmāk mēs pieņemām tīri domāšanas un noteiktības atbilstības modeļus kā pašmērķi.

Lietderīgas un mērķtiecīgas sistēmas ziņkāre, kurai ir sava nozīmība un iekšējā jēga, pieaug. Līdz ar to pieaug visaptveroša tieksme pēc augstākās zināšanas pašu derīgumam. Evolūcija, mērķu attīstība, tādā veidā notiek tālāk – kā zinātnē, tā arī dzīvē un ētikā. *Čarlzs Darvins* rakstīja: „visaugstākā iespējamā stadija morālā kultūrā ir tā, kad mēs atzīstam, ka mums vajag vadīt mūsu domas”. Mēs panākam to arī ar mainīgumu un atlasī, koncentrējoties uz visnozīmīgākajām domām un ļaujot mazāk nozīmīgām domām aizslīdēt no mūsu uzmanības. *Marvins Minskis* [60, 61] no Mičiganas MI laboratorijas ASV, apskata prātu kā savā ziņā sabiedrību, kā attīstītu komunikācijas, sadarbības un konkurences sistēmu, kur katrs aģents sastāv no vienkāršākiem elementiem. Viņš apraksta domāšanu un rīcību šo aģentu rīcības terminos.

Daži aģenti var darīt daudz vairāk, nekā vadošais rokas princips, kura grābj čempiona kausu. Citi, daudz sarežģītāki aģenti, vada runas sistēmu, jo tas ļauj izvēlēties vārdus neskaidrā situācijā. Mēs iepriekš nezinām mūsu pirkstu novietojumu, kas ir nepieciešams, lai novietotu pirkstus apkārt kausam tieši tā, nevis savādāk. Mēs deleģējam tādos uzdevumos kompetentiem aģentiem un reti pamanām rezultātus, ja tie ir pozitīvi (piemēram, pirksti neslīd, grābjot kausu). Mēs visi jūtam pretrunīgus impulsus un izrunājam vārdus bez iepriekšēja nodoma. Tās ir nesaskaņu pazīmes prāta aģentu vidū. Mūsu izpratne par to ir pašregulēšanās procesa daļa, ar kuru mūsu kopējo aģentu vairākums vada citus. Domu var apskatīt, kā savdabīgu prāta „aģentu”, kurš formējas apmācīšanas un imitācijas procesā. Bet lai sajustu konfliktu starp divām

domām, jums vajag realizēt abas domas kā aģentus jūsu prātā – taču viena var būt veca, stipra un sabiedroto atbalstīta, bet otra, svaiga doma – ir jaunās idejas aģents, tā dažreiz nevar izturēt pat pirmo „kauju”.

Mēs bieži uzdodam sev jautājumu: kādā veidā rodas doma mūsu galvās? Daži cilvēki iedomājas, ka šīs domas un jūtas rodas tieši no aģentiem, kuri atrodas kaut kur viņu pašu prāta ārpusē. Viņi sliecas uz ticību gariem un spokiem. Senajā Romā ļaudis ticēja „ģēnijos”, viņu labā un sliktā izpausmē, kura pavadā cilvēku no šūpuļa līdz kapam, nesot veselību vai slimības. Viņi piedeva īpašu lomu „specifiskam ģēnijam”. Un pat tagad tie cilvēki, kuri nav spējīgi saprast un redzēt dabiskos procesus, kaut kā jaunā radīšanā viņi redz „ģēniju” izpausmi, kā burvestības formu. Bet faktiski tieši gēnu evolucionārā attīstība rada cēloņus prāta dzirksteles radīšanai. Prāts paplašina mūsu pieredzi un zināšanu, izmaina perceptuālo tēlu, jaunu ideju un domu standartus, atlasot labākās no tām. Ja ātra struktūras izmaiņa un pārbūvēšanās tiek pavadīta ar efektīvo izvēli ar vadītas zināšanas palīdzību, kas ir aizņemta no citām analogiskām struktūrām, tad kāpēc tādas saprātīgas struktūras neatklāj mums tā būtību, ko mēs mēdzam saukt par „ģēniju”? Domāšana, kā dabiskā procesa analīze, padara mākslīgo intelektuālo mašīnu un mehānismu radīšanas ideju mazāk noslēpumainu un vairāk satriecošu. Bez tam, praktiskie pētījumi MI jomā demonstrē mums intelektuālo mašīnu principiālās iespējas un to, kā un kur tās var efektīvi strādāt.

7. Par prātīguma principiem

The mind is an iceberg – it floats with only one-seventh of its bulk above water
(Sigmund Freud)

Daudzi principi un likumsakarības apraksta saprātīgu darbību. Taču vienīgā viedokļa par to sastāvu līdz šim brīdim nav. Nosauksim saprātīgo sistēmu pamatprincipus [1, 4].

Esamības un darbības princips. Šis princips ir primitīvs, bet nav nederīgs. Dažas filozofijas sistēmas var apgāzt pat pašus acīmredzamākos faktus. Tātad, prāts pastāv, aktīvi darbojas un ietekmē savu nesēju, kā arī apkārtējo pasauli.

Subjektivitātes princips. Katrai saprātīgai sistēmai ir sava pildīšana. Vienādās situācijās dažādas saprātīgas būtnes uzvedas dažādi, domā dažādi un katrai ir sava pozīcija, izejot no sava materiālā apvalka, uzkrātās un mantotās pieredzes īpatnībām.

Universalitātes princips. Izmantojot tādas pašas metodes, var secināt, ka daudzi dažādu sistēmu saprātīgas uzvedības īpašības un likumsakarības ir līdzīgas un pat analogiskas. Tas notiek tāpēc, ka saprātīgas sistēmas darbojas *vienā pasaulē*, un šīs pasaules likumi uzliek ierobežojumus uz būtņu iespējamām darbībām, liek tām darboties un domāt pēc līdzīgiem likumiem.

Piemēram, spēja atrast īsāko izejas ceļu no labirinta attīstās daudzām būtnēm, tai skaitā arī cilvēkam, žurkām u.c. Aktīvā mijiedarbībā ar pasauli prāts iegūst dažādas iemaņas, daudzās sfērās vienādās visām būtnēm. Pēc kāda principa viņi tās iegūst? Acīmredzot pēc viena iespējamā – adaptācijas pie apkārtnes, pielāgošanās, izdzīvošanas u.tml. principa, un šis princips nosaka visus prātus. Tomēr der piebilst, ka teorētiski dažādas sistēmas var trijos veidos reaģēt uz pasauli: 1) aktīvi pielāgoties evolūcijas (ģenētikas) un informācijas (saprātīgā) ceļā pasaules likumiem (ko mēs arī darām); 2) pārtraukt aktīvo informācijas apmaiņu ar pasauli, samierināties, izkust tajā (nedzīvās sistēmas); 3) mainīt pasaules likumus, tādā veidā sagraujot to.

Mainīguma un nepārslodzes princips. Prāts nepārtraukti darbojas un mainās laikā. Jebkurš cilvēks var apliecināt to, ka katrā laika momentā viņa domas mainās („karāšanās” uz vienas domas ir samērā nereta parādība). Dzīves gaitā viedoklis par kaut ko var krasi mainīties. Acīmredzams, ka vienā laika momentā ir aktīva viena vai par dažas domas. Pie tam tās ir aktīvas dažādā mērā. Domas aktivitātes mērs runā par uzmanības (vai apziņas) virzības mēru uz doto informāciju. Uzmanības centrā atrodas informācija ar vislielāko aktivitāti.

Egoisma princips. Prāts darbojas pateicoties motivācijai, ir virzīts uz noteiktu mērķu sasniegšanu u.tml. Kāds tad ir tā mērķis? Mērķis ir baudu iegūšana, ja nepretendē uz augstāko apkopojuma mēru. Ģenētiski mūsos ir ielikta bauda no ēšanas, gulēšanas, vairošanās, aktīvām darbībām un daudz kā cita. Bauda no abstraktiem tēliem izpaužas pateicoties asociatīvam sakaram ar primāriem baudas tēliem.

Principā tiešām visu var novest līdz baudai. Kāpēc? Tāpēc, ka cilvēks vienmēr var izvēlēties no divām alternatīvām vairāk gribēto, t.i. salīdzināt tās. Tādā veidā, visu,

ar ko operē prāts, var izvietot uz zināmas baudu, vēlēšanās un prioritāšu skalas. Proti, var parādīties jautājums – bet kā tad gadījums, kad cilvēks nevar izvēlēties starp divām alternatīvām. Tas izskaidrojams ar to, ka klusēšana vai bezdarbīgums ir trešā alternatīva un tā var būt vairāk vēlama. Bet tomēr, ja piespiestu cilvēku, viņš noteikti izvēlēties kaut ko. Tātad skaitās, ka viss, ko dara cilvēks ir virzīts uz maksimālās baudas, vislabvēlīgāko apstākļu iegūšanu.

Vai tiešām tas tā ir? Lai saprastu prāta mehānismu, vajag uzcelt *domājošās mašīnas* shēmu vai modeli. Tikai zinot visas prāta sastāvdaļas – tā objektus, procesus, principus u.c. var atkārtot tā saprātīguma pamatefektu. Problēma ir šādu prāta aspektu modelēšanas neiespējamība, kā apziņa, jūtas, daiļrade, dvēsele u.tml. Bet ja pieturas pie materiālistiska viedokļa, jautājums par jebkuras lietas modelēšanas iespējamību risinās pozitīvi. Pamata problēmas domājošās mašīnas radīšanā ir pētīšanas un modelēšanas problēmas.

8. Intelektā definīcijas

Intelektā esamība ir vismaz kaut kāda kompensācija par prāta trūkumu (M. Geny)

Daudzi domātāji aprakstīja tādas parādības, kā prāts, intelekts, apziņa un domāšana u.tml.[1-3, 51-53], veidoja tuvinātus nervu šūnu grupu modeļus [56, 57, 62-64] utt. Tāda fenomena esamība, kā prāts, ir nenoliedzams. Taču visi šie meklējumi neļauj izveidot pilnvērtīgu prāta modeli un aprobežojās ar vienkāršotu un visai nenoteiktu kvalitatīvu aprakstu. Piemēram: „Prāts operē ar jēdzieniem, ir virzīts uz vajadzību apmierināšanu un ir pakārtots noteiktai mērķu sistēmai” vai „Smadzenēs tiek veidots apkārtējās pasaules modelis”. Prāta modeļa izveidošanas jautājums, kā izrādījās, ir viens no sarežģītākiem, jo smadzenes ir smalkākais mehānisms Visumā. Tādu zinātņu nozaru attīstība kā kibernetika, psiholoģija, informātika un datu bāzes, dažādu procesu un sistēmu modelēšana ar jaunu informācijas glabāšanas tehnisko iespēju palīdzību, kā arī manipulēšanas iespējām ar lieliem informācijas kopām, kļūva par jaunu impulsu prātīguma mehānismu pētīšanā. Pasaule stāv uz kārtējā datoru buma sliekšņa. Jaunā

tehnoloģija pamudināja visas pasaules laboratorijas uzsākt datora aizstāšanu, kā fantastiski ātru skaitļošanas ierīci, uz iekārtu, kas imitē cilvēka domāšanas procesus, tas ir, uz mašīnām, kas būtu spējīgas spriest, izdarīt secinājumus un pat mācīties. Vēl joprojām nepastāv vienota MI formulējuma. Tāpēc nocitēsim „*MI Enciklopēdijā*” doto definīciju: „Mākslīgais intelekts – ir zinātnes un tehnikas nozare, kas ir saistīta ar datoru modelēšanu un intelektuālās uzvedības pētīšanu, kā arī ar tādu ierīču radīšanu, kurām piemīt tāda uzvedība”.

Klasiskā cilvēka dabiskā intelekta definīcija etimoloģiski izriet no latīņu vārda „*intellectus*”, kas vienlaicīgi nozīmē izzināšanu, saprašanu un prātu. Faktiski termins „*intellectus*” sastāv no trim elementiem „*in – tela – lectus*” un burtisks šī vārda tulkojums ir šāds: „smalkā struktūrā esošs”. Tas nozīmē, ka smalkajā smadzeņu struktūrā ir ievietots prāts. Līdz ar to termins „*intelekts*” ir tikai adrese, pēc kuras var atrast visus jēdzienus, kas līdzinās prātam pēc tā nozīmes. Pie tādām attiecina domāšanas un racionālās izzināšanas spējas. Vairākums triviālo priekšstatu par prātu, kā arī par intelektu noformēti filozofisko zināšanu par prātu ietekmē. Pie prāta tiek pieskaitīta saprašanas un apjēgšanas būtība, kas saliekas zināmā panloģisma augstākajā sākotnē, racionālās izziņas un cilvēku uzvedības pamatā, kulta augšupejā augstākos garīgajos izzināšanas sfērās. Visciešāk saistīti filozofijas jēdzieni ir prāts un saprāts.

Par šiem jēdzieniem tiek runāts *Kanta* un *Hēgeļa* darbos. Tā pēc *Kanta*, saprāts ir spēja radīt jēdzienus, likumus un spriedumus, bet prāts ir spēja radīt metafiziskās idejas. Savukārt, *Hēgelis* pierāda, ka saprāts ir spēja operēt ar gatavām zināšanām, bet prāts ir jauno iespēju atklāšana no vecām zināšanām un jebkura jauno zināšanu sintēze. Šāds divu filozofu secinājums var būt pieņemts par priekšnosacījumu smadzeņu funkciju pētīšanai. Mūsdienu etapā dabiska intelekta (DI) zinātnes rašanās un attīstība virzās uz jaunām zināšanām par cilvēku, vienlaicīgi tiek veikts darbs citādā radošā virzienā – mākslīgā prāta (MP) radīšana. Mākslīgais intelekts jau tagad risina tādus domāšanas uzdevumus, kuri liek cilvēkam iedomāties par sevis paša intelekta attīstības nepieciešamību. No slepeni laboratoriju istabām datoru parādījās tādēļ, lai palīdzētu cilvēkam ar rakstīšanu un skaitīšanu, lai izklaidētu to ar amizantām spēlēm mājās vai iesaistīt lietišķās spēlēs darbā. Praksē intelektuālās mašīnas risina salīdzinoši

vienkāršus, atkārtotus uzdevumus, bet laboratorijas apstākļos šīs mašīnas prot daudz vairāk. MI pētnieki apgalvo, ka datori jau šodien prot daudz vairāk no tā, ko faktiski dara, un aizvien mazāk un mazāk cilvēku nepiekrīt tam. Lai saprastu mūsu nākotni, mums nepieciešams redzēt, vai MI šodien ir tik pat neiespējams, kā cilvēku ceļojums uz Titānu.

Ir nepieciešams precizēt terminu „*optimāls risinājums*” un „*mērķis*” nozīmi. Vārds „*optimalitāte*” ir diezgan „viltīgs” vārds, kas atgādina mums gan par materiālo resursu, gan un mūsu intelektuālo spēju ierobežotību [22]. Jo, pirmkārt, nedrīkst aizmirst, kā cilvēks ar savām smagām smadzenēm un vieglu prātu, kā bioloģiskās sistēmas racionālais modelis (dažreiz arī šis apgalvojums skan visai optimistiski) pats ir „salikts” tā, ka rezultāts ir ļoti tālu no optimalitātes. Otrkārt, ne iepriekšminētam cāram, ne govij tīrā laukā „*optimalitāte*” nav vajadzīga – tiem iztikas resursu ir pārmērīgi daudz attiecībā pret potenciālo apetīti. Taču šo terminoloģisko niansi var viegli pārvarēt, aizvietojot vārdu „*optimalitāte*” ar izteiksmi „*sapratīgā lietderība*” (krievu valodā „*целесообразность*”). „*Lietderība*” ir saistīta ar dzīves jēgu. Bet kāda ir dzīves jēga? Tā ir cīņā par pašu dzīvi, tas ir, par eksistenci, par izdzīvošanu. Bet kādēļ? Lai turpinātu dzimtu, maksimāli pagarinātu savas personīgās dzīves ilgumu un galu galā nodot neatrisinātu dzīves jēgu problēmu turpmākām paaudzēm? MI „dzīves jēga” ir derīgums MI intelekta veidotajam. Dabīgā intelektā (DI) apstākļos „dzīves jēgas” mums faktiski nav (izdzīvot par katru cenu, drausmīgās mokās morāli degradējoties netikumus, lai tik un tā nomirtu?). Bet mēs tomēr runājam par MI kā par instrumentu cilvēka dzīves jēgas nodrošināšanai, bet ne par „mākslīgo dzīvi” (MD). Ja cilvēkam būtu zināma dzīves jēga, augstākā mēra patiesība, tad zinātne (tapāt, kā reliģija) zaudētu savu jēgu. Tajā skaitā arī MI veidošanas jēga.

Izanalizējot pēc būtības terminoloģiskas nianšes, var viegli nonākt pie slēdziena, ka MI sistēmu veidošanā „*optimalitāte*” un „*sapratīga lietderība*” nozīmē vienu un to pašu. Mērķis ar motivācijas un tieksmes operatora palīdzību attaisno (realizē) šo lietderību un MI „dzīves jēgu”. Sanāk: a) MI ir optimālo risinājumu meklēšanas procesors ierobežoto materiālo un garīgo resursu apstākļos; b) DI ir dzīves jēgas meklēšanas procesors ar motivācijas un tieksmes operatoru palīdzību optimālo lēmumu pieņemšanai izdzīvošanas problēmu risināšanas un dzīves pagarināšanas

mērķa sasniegšanas gaitā. Saprotams, ka mironim nekādas personīgas jēgas nav, jo nav dzīves. Pie tām, „dzīves jēga” dabā „tīrā veidā” nepastāv. Nav arī tās fizikālo parametru. Pat ja mēs atrastu dzīves jēgu, ko mēs ar to darītu? Tāpēc šodien runā par MI radīšanu kā par DI funkcionālo kopiju, nevis kā par dzīves jēgas meklēšanas procesoru.

Tātad, intelekts ir optimālo risinājumu meklēšanas procesors laikā un telpā, kurš vislabāk apmierina intelekta veidotāja mērķa sasniegšanas nosacījumus. Šie nosacījumi var saturēt ierobežojumus un ievērot citu intelektu pret darbību. DI ir dzīvās matērijas iedzimti-iegūtā īpašība. MI ir algoritms, nedzīvās matērijas iegūta īpašība. Pēc šāda intelekta formulējuma nav principiālas atšķirības starp MI un DI. Par atslēgas jēdzienu kļūst primāra intelekta dibinātāja (radītāja, veidotāja) jēdziens.

Bet uzreiz parādās jauni jautājumi: vai prāts ir pietiekošs un nepieciešamais nosacījums dzīvas matērijas un dzīves radīšanai? Vai dzīvei vispār ir jābūt „saprātīgai”? Un kas novērtēs radušos un pastāvīgi atjaunojušos prātus (MI)? „Virsprāts”? Kas tas ir? Dievs? Secinājums: radīt mākslīgu prātu, kas potenciāli pārspēj cilvēka prātu būtībā nav iespējams. Mēs par to vienkārši neuzzināsim. Ja tas notiks mūsu klātbūtnē, bet neatkarīgi no mūsu gribas, tad mēs to atklāsim tikai pēc bēdīgām sekām, kuras novērst nebūsim vairs spējīgi. No otras puses, ir absurdi domāt, ka globālā izpratnē MI var būt pārāks par cilvēka DI. Cilvēka smadzenes ir dzīvās matērijas eksistences augstākā forma. Bet prāts nav nepieciešamais un pietiekamais nosacījums dzīvās matērijas radīšanā jo, kā zināms, dzīvā matērija ir radusies no nedzīvās un diemžēl cilvēks nav optimāla, bet racionāla biosistēma. Tāpēc arī dzīvei nav obligāti jābūt „saprātīgai”, kaut gan saprātīgas būtnes iegūst priekšrocības cīņā par izdzīvošanu. To mēs varam noverot dabā. Cilvēka prāta telpa ir bezgalība, bet ar cilvēka un viņa dvēseles (kā noderīgas informācijas „sablīvējuma” sabiezējuma) nemirstības iegūšanu prāts kļūst mūžīgs arī laikā. Vai arī beigs pastāvēt Visumā tieši pēc „saprātīga” cēloņa. Bet kādēļ tad taisīt MI? Atbilde ir vienkārša: lai mazāk strādātu, labāk un ilgāk dzīvotu. MI nebūs gudrāks par cilvēku, tas vienkārši kļūs veikls un ātri domājošs cilvēka palīgs. Bet viņš tik un tā paliks par mašīnu, kura nekad nepārņems cilvēka tiesības, pārdzīvojumus, jūtas un emocijas.

9. Domājošo mašīnu intelektuālie modeļi

Komplicējot tīmekli, zirnekļi pilnveido mušas (M. Geny)

Domājošām mašīnām nav jābūt līdzīgām cilvēkiem dienesta formā, kas ir apveltīti ar dažām domāšanas spējām. Tik tiešām, dažām mākslīgā intelekta sistēmām ir tādas pašas īpašības, kā diplomētam speciālistam mākslu jomā, kurš spēj spēlēt varena dzinēja lomu dažādu projektu realizēšanā. Tomēr doma par to, ka cilvēka prāts ir cēlies no nesaprātīgas matērijas, zaudēs savu nozīmi, salīdzinot ar ideju padarīt mašīnu par domājošu. Prāts tāpat kā citas kārtības (antihaosa) formas, ir attīstījies, pateicoties mainīgumam un selektīvai izvēlei haosa un pretrunu pasaulē.

MI jēdziens ir visai miglains. Apkopojot visu pēdējo 30 gadu laikā veikto, izrādās, ka cilvēks vienkārši grib radīt sev līdzīgo, grib, lai kaut kādas darbības tiek veiktas racionālāk, ar mazākiem laika un enerģijas zudumiem. Lai „mākslīgais intelekts” veiktu savus tiešos pienākumus, tam nav vajadzīga „dzīves jēga”, tas aprobežojas ar vērtīgumu „radītājam”, intelekta veidotajam. Taču mēs nevaram likt viņam dzīvot savā vietā. Bet mēs varam dzīvot kopā! No kurienes ir radies intelekts? Atcerēsimies, ka informācija ir potenciālas zināšanas. Informācija ir arī atmiņa, cietais disks, grāmata, kristāls, DNS un tml. Informācija MI tiek apskatīta kā „sastingušas zināšanas”, motivēto, potenciāli jauno zināšanu avots. Informācija MI sistēmās ir savdabīga strukturālā hologramma, dabiskās dzīves (DD) gaitas „procesa”, iekšējo motivēto aktivitāšu un evolūcijas pēda, kas dabiskā veidā „iezīmējās” MI telpas - laika sistēmas hologrammā kā DD procesa determinēti-motivētu imitācijas trendu mijiedarbības rezultāts. „Mākslīgā dzīve” (MD) ir „tehnoloģiskā prāta” produkts. Tātad, MD ir termins, kuru MI pētījumos izmanto noteiktas jēdzienu telpas apzīmējumam.

Intelektam, kā vektoram, ir kvantitatīva vērtība, virziens un iekšēja tieksme pēc mērķa dominantes. Bet jo augstāks intelekta līmenis, jo augstāka ir tā nenoteiktības pakāpe. MI izstrādātāju mēģinājumi pazemināt mākslīgā intelekta problēmu līdz programmu un algoritmu izstrādes uzdevumam neizskatās pārliecinoši. Tāpat kā aparātu iespējas ne vienmēr ir adekvātas to programmu sarežģītumam, ko risina ar MI palīdzību. Arī MI skaņu pazīšanas datoru programmas nekad nekļūs par tiešu domu

realitāti, jo domāšana ir primāra un tas ir „process”, bet skaņa vai valoda ir sekundārās un tas ir tikai „instruments”.

Daudz perspektīvāk IST jomā izskatās vizuālo metožu, fraktāļu, neironu tīklu (NT) un ģenētisko algoritmu (ĢA) un MI metodoloģijas pielietošana [1, 4, 5, 34-40, 45-46, 71]. Piemēram, MI inteligentu instrumentu apmācība domāt ar fraktāļu tēliem tuvina tās procesoru cilvēka domāšanas asociatīvai metodei. Fraktāļu tēli [3, 12, 13, 24, 104] satur sevī telpas - laika notikumu infokvarku elementārās pēdas un tur glabājas pagātnes, tagadnes un nākotnes izejas tēlu formā. Tāpēc principiāli ir iespējama dažāda veida apmācības algoritmu pielietošana, t.i. interpolācijas, filtrācijas un ekstrapolācijas algoritmu izmantošana (9. att.). Bet jebkura secība realizējas laikā. Fraktāļiem ir raksturīga ne tikai iekšējā tēlu formēšanas, transformācijas un evolūcijas secība, bet arī iekšējais, endohronais laiks, kas noteic tēlainās domāšanas ātrumu. Intelektuālo mašīnu gandrīz radošu īpašību realizācijai vadošās elektroniskās kompānijas aktīvi izstrādā tā sauktos „video procesorus”.

Fraktāļu tēli spēj „atcerēties” savu pagātni, izcelsmi, „pareģot” savu formu nākotnē un tādējādi ne tikai virzīt savas individuālās pašattīstības evolūciju sava neizbēgamā „likteņa” virzienā, bet arī atgriezties atpakaļ, lai „uzlabotu” sava „likteņa” iznākumu. Līdz ar to fraktāļiem piemīt iekšējā potencialitāte, „intuīcija”. Piemēram, ar fraktāļu tēliem un Černova seju [20-24] tehnikas tēliem (*Chernoff faces*) intelektuālā mašīna var izpaust pat vienkāršas emocijas. Pētījuma objekta uzvedības modeļa faktoru analīzei ir jāaplūko uz otra tipa biokibernētisko modeļu izmantošanas pamata, kad tiek veikta sistēmas uzvedības analīze un tiek noteiktas atsaucēs funkcijas atgriezeniskās saites ieviešanai. Tālāk, uz standarta funkciju pamata, tiek noteikti ietekmes faktori uz radīto metrisko tēlu formu un tiek izstrādāts atgriezeniskās saites vadības algoritmi.

Ar algoritmu un programmu palīdzību veidojas metriskais tēls un turpmāk tas paliek atmiņā AK etalona veidā. Modeļu struktūras pamata kontūras veidojas pilnīgā atbilstībā ar mērīšanas procesa mērķiem, esošiem datiem un zināšanu līmeni par attiecīgā pētījuma objekta īpatnībām. Sistēmiskās pieejas būtība intelektuālo sistēmu izstrādāšanā ir apgalvojumā, ka jebkura problemātiski strukturētā sistēma jebkurā līmenī ir mērķtiecīga sistēma, kura var būt sadalīta un sintezēta no vienkāršāko

mērķtiecīgo apakšsistēmu kopuma. Tāda apakšsistēma interpretē, grasās un paredz, t.i. tai ir intelekta pazīmes. Mākslīgais intelekts (MI) attiecas pret jebkuru objektu vai mērķi tādā veidā, kā prasa to būtība. Universālā domu neatkarība, intelekts var atrast oriģinālās, efektīvās un pat neparedzētas atbildes. Mēraparāta rādījumi pie tam var atšķirties viens no otra tik daudz, ka tie var nesaturēt vispār nekādu būtību, jo pētījuma objekts ir saistīts ar noteiktu koordinātu nolasīšanas telpas-laika sistēmu.

Pie intelektuālo mērīšanas sistēmu izstrādāšanas ir jāparedz, lai tāda sistēma būtu pieejama cilvēka (eksperta, speciālista) intelektam, pilnīgākām zināšanām, veselam prātam un praktiskās pieredzes attīstībai, tai skaitā arī neformalizētas un nesistematizētas, kā arī DI intuīcijai. Turpmāk tāda mērīšanas – izzināšanas ekspertu sistēma var nepārtraukti pilnveidoties ar apmācību – kopā ar standartu, kritēriju un atbilstības vērtējumu atjaunošanu. Dialoga un pašapmācības rezultātā veidojas tālākai attīstībai un pilnveidošanai atklāta indikācijas, identifikācijas, mērīšanas un korekcijas intelektuālā sistēma ar pētāmās fizikālas vides fraktāļu metriskiem tēliem.

Tāda mērķtiecīga un uz mērķa sasniegšanu iekšēji motivēta intelektuālā sistēma ir virzīta nevis uz mērīšanas procesu, bet gan uz mērīšanas rezultātu. Fundamentālie pētījumi mākslīgā intelekta (MI) jomā ļauj radīt abstraktos universālos modeļus, kuras uzlabo reālos procesus un rezultātā pētāmais stāvoklis kļūst caurspīdīgs, pieejams teorētiskai analīzei un apkopojumiem. Turpmāk tas atvieglo apjēgšanu, programmu lingvistisko aprakstu un jaunās zināšanas likumsakarību formulēšanu. Mūsdienu pasaulē, kas ir novērtēta ar derīgas un nederīgas informācijas, progresīvo un atpalikušo tehnoloģiju plūsmām, viss ir savstarpēji saistīts. *Haosa teorija* rāda, ka šie procesi nav nejauši. Konkrēto modeļu pētīšanā labi palīdz imitācijas komplekso sistēmu pētījumi.

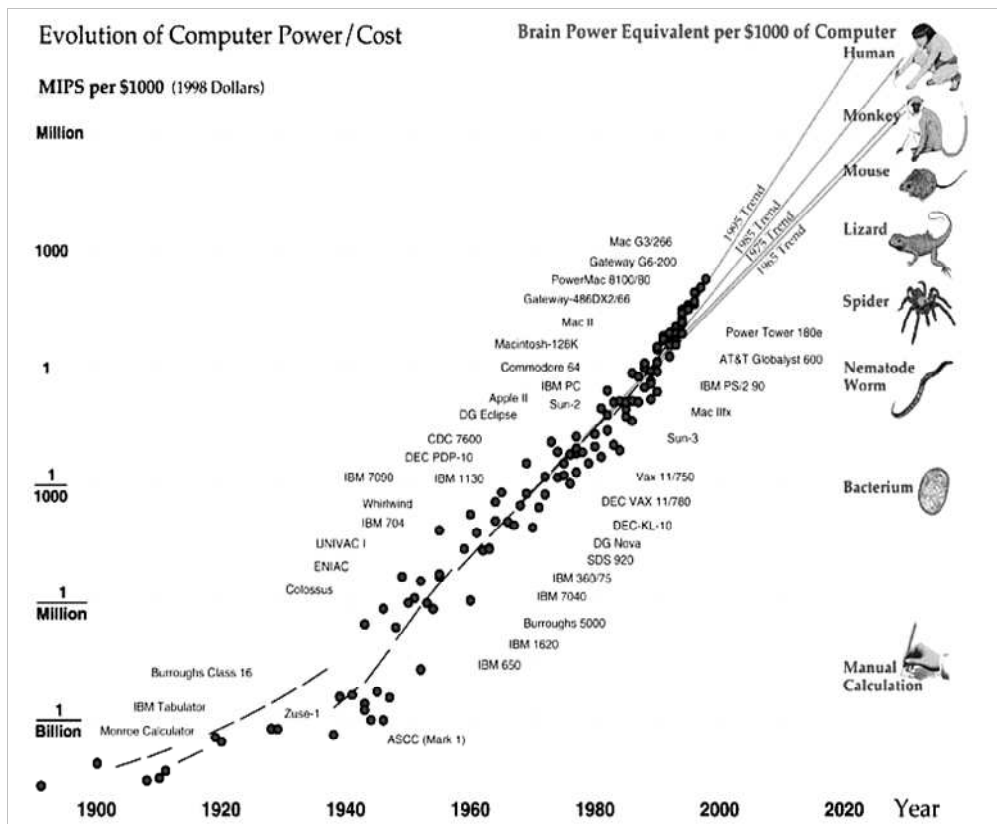
Dinamiskā sistēma, kura sastāv no vairākiem komponentiem, ir labi attēlota „smilšu kalniņa” modelī (8. att.). Berot smiltis uz galda, kalniņš no sākuma būs plakans, smilšu graudiņu izvietošana savstarpēji nav saistīta un sistēma kopumā savstarpēji neiedarbojas. Bet par cik smilšu kalniņš turpmāk kļūst stāvāks, smilšu kalniņš sāk kustēties. Tas nozīmē, ka sistēmas atsevišķās daļas sāk savstarpēju iedarbību un sistēma nonāk kustības stāvoklī. Sīkākie kustības elementi, smilšu graudiņi, visu laiku kļūst aizvien aktīvāki un kalniņa augstums sasniedz kritisko vērtību.

Turpmāk smilšu graudiņi kustas visos iespējamajos virzienos. Tas, kas notiek vienā kalniņa pusē, ietekmē to, kas notiek kalniņa otrajā pusē. Mēs vairāk nepinamies tikai ar vienu smilšu graudiņu, bet mums ir tikai viens kalniņš, kas ietver sevī vienu lielu dinamisku sistēmu. Rezultātā sistēmas savstarpēja iedarbība iegūst globālu raksturu. Ar tādām sistēmām mēs saskaramies visur. Piemēram, bioloģisko stohastisko procesu analīzē lauksaimniecībā. Faktiski tas nozīmē, ka mūsu materiālā un garīgā pasaule ir tā pati kritiskā joma, tāds pats smilšu kalniņš, kas atrodas tādā stāvoklī, kad tas vairs nevar augt (H_{kr} , 8. att.). Šis ir radikāli jauns priekšstats par pasaules attīstību un evolūciju, kas atšķiras no pasaules un to parādību, kas notiek mums apkārt, parastās izpratnes [14, 15].

Mēs uzskatām, ka atrodamies sabalansētā līdzsvarā, bet faktiski tāds priekšstats ir ļoti tāls no realitātes. Tas nozīmē, ka mēs dzīvojam nevis stabilā, proporcionāli attīstošā pasaulē, kā mēs līdz šim domājām, bet esam globālās vispasaules sistēmas daļa, kura balansē uz kārtības un nekārtības, haosa un antihaosa robežas. Cilvēka dabas haoss tiek kompensēts ar pasaules antihaosu, kurš alkst pēc kārtības. Pasaules eksistencei šis nosacījums ir sākotnējs, kā augstākā būtība, kura ir ieslēgta pašā antihaosa dabā. Pasaules, globālā haosa un antihaosa likumu rakstura saprašana var beigu beigās novest mūs līdz skaidrai saprašanai par visu pasaulē notiekošu procesu augstās pakāpes integrāciju, par matērijas, apziņas un tās sociālās vides garīgo dabu, kurā mēs visi dzīvojam un radam labumu [16].

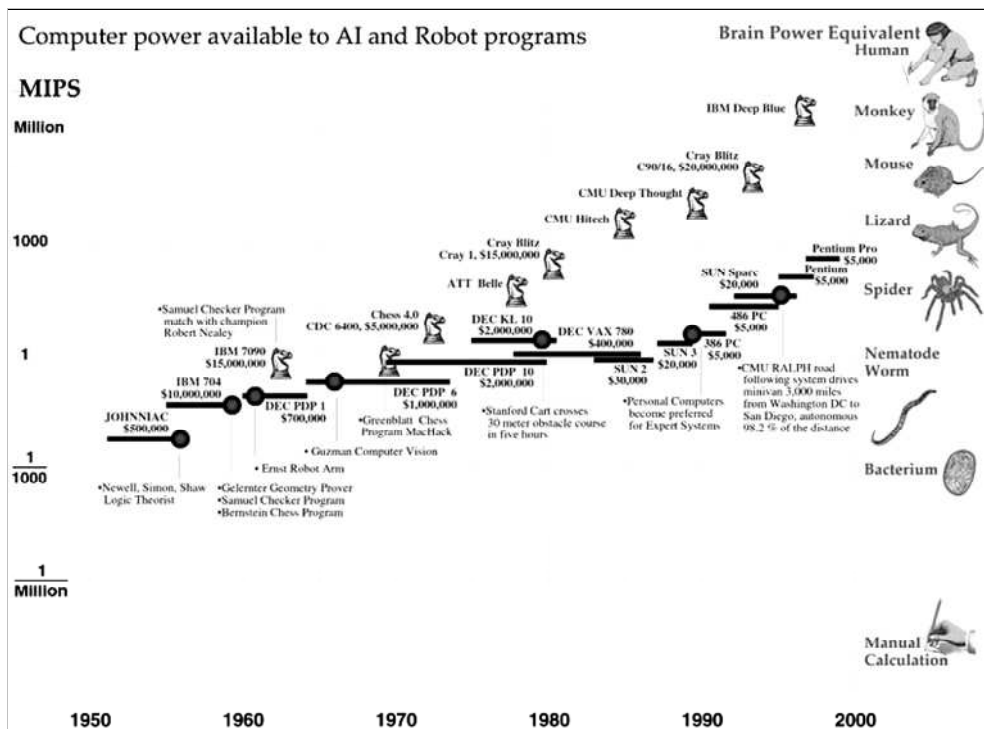
Biokibernētiskā pieeja intelektuālo sistēmu modelēšanā ļauj apskatīt divus „intelektuālos” vadības modeļu tipus. *Pirmā tipa* modeļu izstrādāšanai pietiek izpētīt tikai procesa vai pētījuma objekta „iekšējās” saites un parametrus, neņemot vērā ārējo faktoru īpatnības. Šāda tipa modeļi var būt derīgi atbilstības objekta iepriekšējai identifikācijai. Tālākā modeļa izmantošana ir atkarīga tikai no tā, cik veiksmīgs būs izveidotā modeļa teorētiskais un tehniskais turpinājums [17].

Lai ilustrētu kāpšanu no informācijas uz zināšanām (līdzīgi „0” un „1” attiecībai) un no tām uz intelektu, atcerēsimies joku par papagaiļa intelekta evolūciju. Domāšanas procesa kvalitātes kāpšanas princips tam ir vienkāršs: iemācītais lamu vārds „mulķis” viņam ir abstrakta informācija, frāze „cars – mulķis” ir konkrētas zināšanas par objektu, bet frāze „Nost patvaldību!” – tas jau ir secinājums, intelekts.



10. att. IT attīstības tendences (avots: Moravec, OCLC Firstsearch database)

Ir iespējami vairāki mākslīgā intelekta konstruēšanas varianti. Viens no variantiem, kā uztaisīt domājošu mašīnu (DM), var būt automātiski mehānismi uz dažāda pamata (tehniskais, bioloģiskais pamats). Sākotnēji „radītājs”, primārais intelekta veidotājs, atbilstoši savam zināšanām un pasaules uzskatam, konstruē DM modeli un pasaules modeli. „Radītājs” var ietekmēt sistēmas procesus, izmainot pasaules modeli, vai izmainot tā atsevišķas daļas.



11. att. **IT attīstības tendences** (avots: Moravec, OCLC Firstsearch database)

Tādā veidā „radītājs” var apmācīt, stimulēt, kontaktēt ar modelējamo radību, saņemot atbildes. Tas ir acīmredzams, ka „radītājam” nav vajadzības mainīt ne robota prātu, ne ķermeni, ne principus, pēc kādiem robots mijiedarbojas ar pasauli. DM ķermenis darbojas, kā informācijas nesēju (vadu) tilpne. DM „nāve” iestājas tad, kad beidzas informācijas apmaiņa sistēmas iekšienē.

Otrs DM konstruēšanas veids ir gadījums, kad tiek modelēts pats DM ķermenis, apkārtējā vide un mijiedarbība starp tām. Šāda pieeja uzliek lielus ierobežojumus saprātīgas sistēmas iespējām.



12. att. Vīnu atbilstības kontroles ekspertu sistēma (Austrālija, 2007. g.)

Dažādi pētnieki [42-44, 47, 51-53, 70] cenšas piedot domājošai mašīnai cilvēka domāšanas prototipus, taču tās „domā” savādāk, nekā cilvēki. Piemēram, robotam galīgi nav vajadzīgs jēdziens „siltums”. Pavisam cita lieta, kad robots asociē vārdu „siltums” ar temperatūras izmaiņu. Vispārīgi jebkuru cilvēka jēdzienu vai attēlu var parādīt citai saprātīgai sistēmai ar sakārtotu daudzu iekšēju tēlu sakarību. Tāpēc ir bezjēdzīgi ievadīt domājošā mašīnā cilvēka jēdzienus, domājošai mašīnai tas jādara patstāvīgi.

10. Ārprāta brāļi

Dažreiz prāta balss skan tik nepārliecinoši, kā svešs (M.Geny)

Pašam galvenajam testam, kas pārbauda cilvēka intelektuālās spējas (IQ) ir divi izgudrotāji. Pašu pirmo testu, kura uzdevums ir noskaidrot pārbaudāmā (cilvēka) intelekta līmeni, izgudroja franču psihologs *Alfrēds Binē* 1905. gadā. Sākumā to izmantoja, lai atklātu bērnus, kuri atpaliek attīstībā no citiem viņu vecuma bērniem. Pamatojoties uz šo testu, 30-to gadu sākumā *Stenfordas Universitātes* profesors *Luiss*

Termans noteica cilvēku intelektuālo spēju skaitlisko analogu – „intelektuālās attīstības koeficientu” vai vienkārši IQ (*intelligence quotient*).

Testu sāka plaši izmantot ASV un Eiropā. Tolaik daudzi cilvēki aizrāvās ar eigēniku un bija populāra spriešana par to, ka garīgi atpalikušos cilvēkus ir jāiznīcina. Piemēram, viena no valsts programmām ASV (ar nosaukumu „Trīs paaudzes idiotu – pietiks!”), kuras mērķis bija nācijas genofonda uzlabošana, pamatojās uz IQ izmantošanu. Bet 1958. gadā angļu sociologs Maikls Jangs izgudroja terminu „meritokrātija”, lai apzīmētu biedrības, kuras ir organizētas, bāzējoties uz IQ testiem. Tādā veidā paši zinātnieki vēl līdz galam nezin, kā šis tests ietekmēja cilvēci – pozitīvi vai negatīvi. No paša izgudrošanas brīža tests praktiski nemainījās.

Līdz šim tajā ietilpst daži nesarežģīti uzdevumi ātrai risināšanai, kuri neprasa iepriekšēju sagatavošanos. Pie tam, tests mēra (pārbauda) iedzimtas spējas un ir maz atkarīgs no iegūtām zināšanām un audzināšanas apstākļiem. Vidējo IQ vērtību parasti pieņem kā 100, atbilstoši 50-75 balles ir intelektuāli atpalikušie cilvēki, bet 120-150 – augsti apdāvināti.

Piemēram, fiziķiem-teorētiķiem koeficienta vērtība ir vienāda ar 130. Kaut arī zinātnieki noteica, ka apmēram 75%-iem cilvēku uz Zemes ir „idiotu” līmeņa IQ, tas ir tikai 50-75 balles. Patiesībā zinātnieki apgalvo, ka faktiski notiek tālāka cilvēka degradācija, nevis „evolūcija” [4, 32, 41]. Tāpēc jautājums par „pilnīgā cilvēka” uzlabošanu un MI radīšanu šim nolūkam pat neprasa apspriešanu.

11. Mākslīgā intelekta filozofiski-metodoloģiskās problēmas

Izstrādājot šo mācību - metodisko līdzekli izveidojās situācija, kura daļēji atgādina fantastiku – tiek mācīts tas, kā vēl nav. Un ja tas neparādīsies tuvāko simts gadu laikā, tad var būt, ka laikmets ar to arī beigsies ne tikai priekš MI, bet arī priekš tā dibinātājiem. No šejienes izriet pamata filozofiski-metodoloģiskā problēma MI jomā – cilvēka domāšanas modelēšanas iespējamība vai neiespējamība. Gadījumā, ja kaut kad būs iegūta negatīva atbilde uz šo jautājumu, tad visiem pārējiem jautājumiem nebūs nozīmes. Tik asi skan problēmas nostādne. Tāpēc, uzsākot MI pētīšanu, mēs jau

iepriekš pieņemam, ka atbilde būs pozitīva. Mēģināsim minēt dažus apsvērumus, kuri mūs pieved pie dotās atbildes.

Pirmais ir sholastu pierādījums [1, 4, 6-9] par nepretrunīgumu starp MI un *Bībeli*. Acīmredzams, pat cilvēki, kuriem reliģija ir sveša, pazīst svēto rakstu vārdus: „Un Dievs radīja cilvēku pēc sava tēla...”. Izejot no šiem vārdiem, mēs varam secināt, ka tā kā Dievs, pirmkārt, radīja mūs, bet otrkārt, mēs esam līdzīgi viņam pēc savas būtības, tad mēs pilnīgi varam radīt kādu pēc sava tēla.

Otrais. Jauna prāta radīšana bioloģiskā ceļā ir cilvēkam pavisam ierasta. Vērojot bērnus, mēs redzam, ka lielāko zināšanu daļu viņi iegūst apmācības ceļā, nevis kā „ieliktas” viņos iepriekš. Šis apgalvojums modernā līmenī nav pierādīts, bet pēc ārējām pazīmēm viss izskatās tieši tā.

Trešais. Tas, kas agrāk likās cilvēka daiļrades virsotne – spēle šahā, dambretes, redzes un dzirdes tēlu atpazīšana, jauno tehnisko risinājumu sintēze, utt. praksē izrādījās ne tik sarežģīts darbs (tagad tiek veikts darbs nevis uz minēto realizācijas iespējamības vai neiespējamības līmeņa, bet par optimālāko algoritmu atrašanu). Tagad bieži šīs problēmas pat nepieskaita pie MI problēmām. Pastāv cerība, ka arī cilvēka domāšanas pilnā modelēšana izrādīsies ne tik sarežģīta.

Ceturtais. Ar savas domāšanas reproducēšanas problēmu ir cieši saistīta pašreproducēšanas iespējamības problēma. Spēja pašreproducēties ilgu laiku izskatījās tikai dzīvo organismu prerogatīva. Taču dažas parādības, kuras notiek nedzīvajā dabā (piemēram, kristālu augšana, sarežģīto molekulu sintēze autokopēšanas ceļā) ir ļoti līdzīgi pašreproducēšanai. 50. gadu sākumā Dž. fon Neimans [51] sāka pamatīgi pētīt pašreproducēšanos un lika pamatus „pašreproducējošo automātu” matemātiskajai teorijai. Viņš arī teorētiski pierādīja to radīšanas iespēju. Pastāv arī dažādi pašreproducēšanas iespējas neformālie pierādījumi, bet programmētājiem visspilgtākais pierādījums droši vien ir datoru vīrusu eksistence.

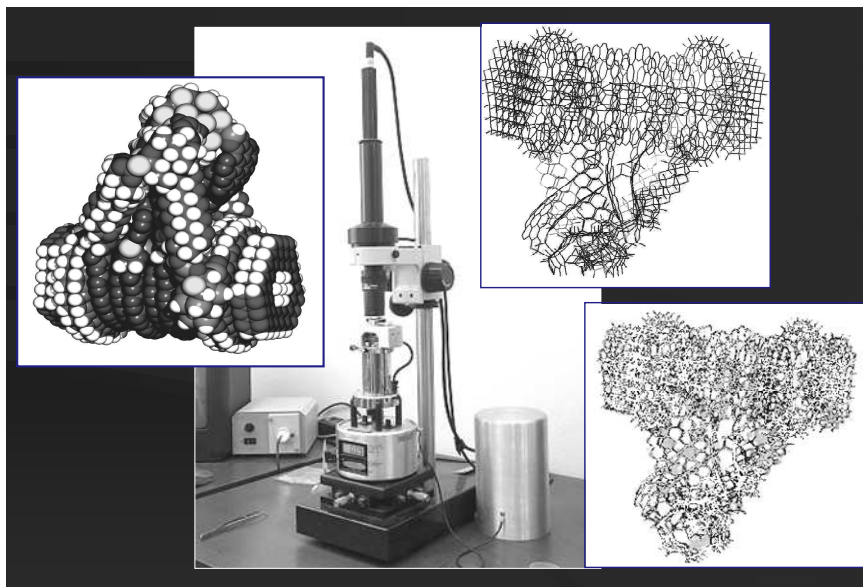
Piektais. Intelektuālo uzdevumu risināšanas automatizācijas principiālā iespēja ar ESM palīdzību tiek nodrošināta ar algoritmiskā universāluma īpašību. Kas tā ir par īpašību? ESM algoritmiskais universālums nozīmē, ka uz tām var programmatiski realizēt jebkurus informācijas pārveidošanas algoritmus, gan skaitļošanas algoritmus, vadības, teorēmu pierādījuma meklēšanas vai melodiju kompozīcijas algoritmus. Pie

tam mēs saprotam, ka procesi, kurus rada šie algoritmi, ir potenciāli īstenojami, t. i. tos var īstenot elementāro operāciju beigu skaitļa rezultātā. Algoritmu praktiskā īstenošana ir atkarīga no līdzekļiem, kas ir mūsu rīcībā, un kas var mainīties tehnikas attīstības gaitā. Tā, kopā ar ātri darbojošos ESM parādīšanos praktiski īstenojami kļuva arī tādi algoritmi, kuri agrāk bija tikai potenciāli īstenojami. Taču algoritmiskā universāluma īpašība neaprobežojas ar tā konstatēšanu, ka visiem pazīstamajiem algoritmiem izrādās iespējama to programmatiska realizācija uz ESM [5, 10, 107]. Šīs īpašības saturam ir arī nākotnes prognozēšanas raksturs: katru reizi, kad nākotnē kaut kāds priekšraksts būs atzīts ar algoritmu, tad neatkarīgi no tā, kādā formā un ar kādiem līdzekļiem šis priekšraksts būs sākotnēji izteikts, to varēs izveidot arī datoru programmas veidā.

Taču nedrīkst domāt, ka skaitļošanas mašīnas un roboti var principā risināt jebkurus uzdevumus. Dažādu uzdevumu analīze noveda matemātiķus pie izcila atklājuma. Bija strikti pierādīta tādu uzdevumu tipu eksistence, kam nav iespējams izmantot kopīgo efektīvo algoritmu, kas risina visus dotā tipa uzdevumus; šajā ziņā nav iespējama tāda tipa uzdevumu risināšana ar skaitļošanas mašīnu palīdzību. Šis fakts palīdz labāk saprast to, ko var izdarīt mašīnas, un ko tās nevar izdarīt. Tiešām, apgalvojums par kādas uzdevumu klases algoritmisko neatrisināmību ir ne tikai atzīšana, ka tāds algoritms mums nav zināms un neviens to vēl nav neatklājis. Tāds apgalvojums ir vienlaicīgi gan prognoze uz visiem nākotnes laikiem par to, ka tāda tipa algoritms mums nav zināms un neviens to neatklās, vai citiem vārdiem, tas nepastāv. Bet kā rīkojas cilvēks, risinot tādus uzdevumus?

Acīmredzams, viņš vienkārši tos ignorē, kas tomēr netraucē viņam dzīvot tālāk. Cits variants ir uzdevuma universāluma nosacījumu sašaurināšana, kad tas tiek risināts tikai sākuma nosacījumu noteiktai apakškopai. Vēl viens variants ir tāds, ka cilvēks uz labu laimi paplašina iespējamo elementāro operāciju daudzumu (piemēram, rada jaunus materiālus, atklāj jaunas atradnes vai kodolreakciju tipus). Nākošais filozofisks jautājums ir MI radīšanas mērķis. Principā viss, ko mēs darām praktiskajā dzīvē, parasti ir virzīts uz to, lai vairāk neko nedarītu. Taču diezgan augstā cilvēka dzīves līmenī (lielā potenciālās enerģijas daudzumā) galveno lomu jau spēlē nevis slinkums (vēlme ekonomēt enerģiju), bet meklēšanas instinkti. Pieņemsim, ka cilvēkam izdevies radīt intelektu, kas pārspēj viņa paša intelektu, kaut gan nevis ar kvalitāti, bet ar

kvantitāti. Kas tad būs ar cilvēci? Kādu lomu tad spēlēs cilvēks? Vai viņš tagad ir vajadzīgs? Vai viņš nekļūs par stulbo resno cūku?

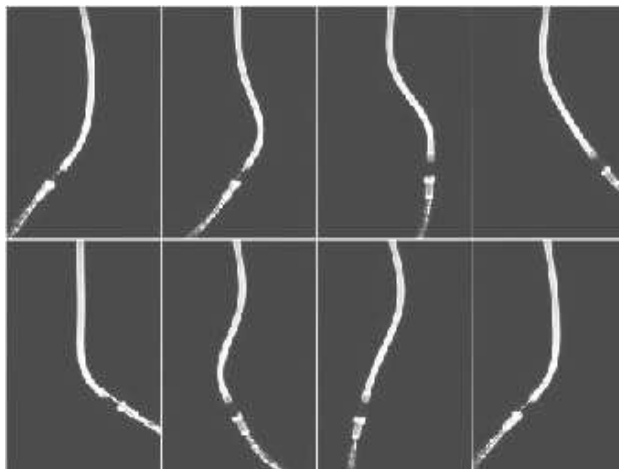


13. att. **Fraktāļu tēlu sintēze molekulārajā sintēzē, atbilstības kontrolē un nanotehnoloģijās** (avots: Moskvins, 1998).

Un vispār, vai principā ir vajadzīgs MI? Iespējamā atbilde uz šiem jautājumiem ir intelekta pastiprināšanas (IP) koncepcija. Šeit var parādīt analogiju ar valsts prezidentu – viņam nav nepieciešams zināt irīdija valenci vai *Java* programmēšanas valodu, lai pieņemtu lēmumu par lauksaimniecības vai rūpniecības attīstību.

Katrs nodarbojas ar savu darbu – ķīmiķis apraksta tehnoloģisko procesu, programmētājs veido programmu; beigu beigās, ekonomists saka prezidentam, ka ieguldot naudu rūpnieciskajā spiegošanā, valsts iegūs 20%, bet vanādija rūpniecībā – 30% gadā. Tādējādi ar IP palīdzību katrs, kam uzdots jautājums, var dot pareizo atbildi, kas īstenībā reāli arī notiek. Minētajā piemērā prezidents izmanto bioloģisko IP – speciālistu grupu ar viņu olbaltumvielu smadzenēm. Bet jau tagad tiek izmantoti arī nedzīvie IP. Piemēram, mēs nevarētu pareģot laiku bez datoru palīdzības, kosmisko kuģu lidojumus no paša sākuma tika izmantotas skaitļošanas ierīces. Bez tam, cilvēks

jau sen izmanto spēka pastiprinātājus (SP) – jēdziens, kas ir analogisks IP. Kā spēka pastiprinātāji viņam kalpo automobiļi, celtni, elektrodzinēji, preses, lielgabali, lidmašīnas un daudz kas cits.



14. att. **Bioloģiskais nanoģenerators** (*Japāna, 2007*)

Mūsdienu atombumba arī ir alas cilvēka parastās rungas redukcijas rezultāts. Būtiskā atšķirība IP no SP ir gribas esamība. Mēs taču nevaram iedomāties, ka pēkšņi sērijas „BMW” sadumpotos un sāktu braukt tā, kā viņam gribas. Nevaram iedomāties tieši tāpēc, ka viņš neko negrib, viņam nav savu vēlmju. Tajā pašā laikā, intelektuālai sistēmai pilnīgi varētu būt savas vēlmes, un tā varētu rīkoties ne tā, kā gribētos mums. Tādā veidā, rodas vēl viena problēma – drošības problēma. Šī problēma jau cilvēces prātus vēl no *Karela Čapeka* laikiem, kurš pirmoreiz pielietoja terminu „robots”. Arī citi rakstnieki-fantasti ziedojuši savu artavu šīs problēmas apspriešanā. Kā pašus slavenākos var minēt rakstnieka-fantasta un zinātnieka *Aizeka Azimova* stāstu sēriju, kā arī diezgan jaunu darbu – „Terminators”. Pie reizes, tieši *Aizeka Azimova* darbos var atrast visatstrādātāko un vairuma cilvēku pieņemto drošības problēmas risinājumu. Šeit iet runa par tā sauktajiem robottehnikas trim likumiem.

Atgādināsim tos. 1. Robots nevar nodarīt cilvēkam kaitējumu vai ar savu bezdarbību pieļaut, lai cilvēkam būtu nodarīts kaitējums. 2. Robotam ir jāklausa komandas, kuras viņam dod cilvēks, atskaitot gadījumus, kad šīs komandas ir pretrunā

ar pirmo likumu. 3. Robotam ir jā rūpējas par savu drošību, par cik tas nav pretrunā ar pirmo un otro likumu.

Pirmajā acu uzmetienā tādiem likumiem, pilnīgi tos ievērojot, ir jānodrošina cilvēka drošība. Taču uzmanīgi tos aplūkojot, rodas daži jautājumi. Pirmkārt, likumi ir formulēti cilvēka valodā, kas nedod iespēju pārvērst tos algoritmiskajā formā. Pamēģiniet, piemēram, pārtulkot jebkurā no jums pazīstamajām programmēšanas valodām tādu terminu kā „nodarīt kaitējumu”. Tālāk iedomāsimies, ka mums izdēvies pārformulēt minētos likumus valodā, ko saprot automatizēta sistēma. Tad interesanti, ko MI sistēma sapratīs ar vārdu „kaitējums” pēc ilgām loģiskām pārdomām? Vai viņa nenolems, ka visas cilvēces eksistence ir kaitējums? Viņš taču dzer, pīpē, ar laiku noveco un zaudē veselību, cieš. Vai nebūs labāk ātri pārtraukt šo ciešanas ķēdi? Protams, var ievest dažus papildinājumus, kas ir saistīti ar dzīves vērtību, gribas izpaušuma brīvību. Bet tie jau nebūs tie vienkāršie trīs likumi, kuri bija sākumā noteikti. Nākošais jautājums ir sekojošs. Ko nolems MI sistēma situācijā, kad vienas dzīves glābšana ir iespējama tikai upurējot citu? Īpaši interesanti ir tie gadījumi, kad sistēmai nav pilnas informācijas par to, kas ir kas. Taču neskatoties uz minētajām problēmām, minētie likumi ir diezgan labs neformāls pamats, lai pārbaudītu drošības sistēmas drošumu MI sistēmās. Un beigu beigās, vai nav drošas drošības sistēmas? Pamatojoties uz IP koncepciju, var piedāvāt šādu variantu. Pateicoties lielam daudzumam pētījumu, neskatoties uz to, ka mēs īsti nezinām, par ko atbild katrs atsevišķs neirons cilvēka smadzenēs, daudzām mūsu emocijām parasti atbilst neironu grupas (neironu ansamblis) uzbudinājums paredzamajā zonā.

Ir veikti arī pretēji eksperimenti, kad noteiktas zonas izsauc vēlamo rezultātu. Tās varēja būt prieka, nomāktības, baiļu, agresivitātes emocijas. Tas liek domāt, ka principā mēs pilnīgi varētu izdalīt organisma prieka emocijas uz ārpusi. Tajā pašā laikā praktiski visi pazīstamie adaptācijas un pašnoskaņošanas mehānismi (pirmām kārtām – tehniskās sistēmas), pamatojas uz „labi-slikti” tipa principiem. Matemātiskajā interpretācijā tā ir kaut kādas funkcijas novadīšana uz maksimumu vai uz minimumu. Tagad iedomāsimies, ka mūsu IP kā tādu funkciju izmanto tieši vai netieši izmērīto cilvēka-saimnieka smadzeņu prieka pakāpi. Ja attiecīgi rīkoties, lai izslēgtu pašdestruktīvo darbību depresijas stāvoklī, kā arī paredzēt citus īpašus psihikas

stāvokļus, tad iegūsim sekojošo. Tā kā ir domāts, ka normāls cilvēks nenodarīs kaitējumu pašam sev un citiem bez attiecīga iemesla, bet IP tagad ir dotā indivīda daļa (nav obligāti fiziskā kopība), tad automātiski tiek izpildīti 3 robottehnikas likumi. Pie tam drošības jautājumi tiek pārvietoti psiholoģijas un likuma aizsardzības jomā, jo sistēma (apmācīta) nedarīs neko tādu, ko negribētu tās īpašnieks. Un tagad ir vēl viens jautājums – vai vispār ir jēga radīt MI, varbūt vienkārši pārtraukt visus darbus šajā jomā? Vienīgais, ko var teikt – ja ir iespējams radīt MI, tad agri vai vēlu tas būs radīts. Un būtu labāk, lai tas notiktu sabiedrības kontrolē, ar drošības jautājumu rūpīgu izstrādāšanu, nekā pēc 100-150 gadiem (ja līdz tam laikam cilvēce vēl neiznīcinās pati sevi) to izdarīs kaut kāds programmētājs-mehāniķis-autodidakts, kurš izmanto sava laika tehnikas sasniegumus. Taču šodien, piemēram, jebkurš lietpratīgs inženieris, ja ir noteikti naudas resursi un materiāli, var uztaisīt atombumbu.

12. Mākslīgais intelekts un ontopsiholoģijas problēmas

Var izdalīt divas MI darbu pamatlīnijas. Pirmā ir saistīta ar pašu mašīnu pilnveidošanu, ar mākslīgo sistēmu „intelektualitātes” paaugstināšanu. Otrā ir saistīta ar „mākslīgā intelekta” kopīgā darba un cilvēka intelektuālo iespēju optimizācijas uzdevumu. Pārejot tieši uz MI ontopsiholoģiskajām problēmām, *O. K. Tihomirovs* izdala trīs pozīcijas jautājumā par psiholoģijas un mākslīgā intelekta savstarpējo iedarbību. Pirmā. „Mēs maz zinām par cilvēka prātu, mēs gribam to atjaunot, mēs to darām par spīti zināšanu trūkumam”. Šī pozīcija ir raksturīga daudziem ārzemju speciālistiem MI jomā. Otrā pozīcija aprobežojas ar tās intelektuālās darbības pētījumu rezultātu ierobežotības konstatāciju, kuru vadīja psihologi, sociologi un fiziologi. Kā cēlonis tiek minēts adekvāto metožu trūkums. Risinājums ir redzams kādu intelektuālo funkciju atveidošanā mašīnu darbā. Citiem vārdiem sakot, ja mašīna risina uzdevumu, kuru iepriekš risināja cilvēks, tad zināšanas, kuras var iegūt, analizējot šo darbu, arī ir tas pamata materiāls psiholoģisko teoriju veidošanai. Trešā pozīcija raksturojas ar pētījuma vērtējumu mākslīgā intelekta un psiholoģijas, kā pilnīgi neatkarīgu zinātņu jomā. Šajā gadījumā ir pieļaujama tikai patērēšanas iespēja, psiholoģisko zināšanu

izmantošanas iespēja psiholoģiski nodrošinot darbus MI jomā. Likumsakarīgi rodas jautājums par MI darbu ietekmi uz psiholoģiskās zinātnes attīstību. *O. K. Tihomirovs* [109] atzīmē kā pirmo rezultātu jaunās psiholoģisko pētījumu nozares rašanos, bet konkrēti, salīdzināmie pētījumi par to, kā vienus un tos pašus uzdevumus risina cilvēks un mašīna. Bez tam, jau pirmie darbi mākslīgajā intelektā parādīja, ka ne tikai uzdevumu risināšanas jomu skar salīdzināmie pētījumi, bet arī domāšanas problēmu kopumā. Parādījās vajadzība precizēt „radošo” un „neradošo” procesu diferenciācijas kritērijus. Vairāk gan uztveršanas, gan atmiņas pētījumi atrodas zem stipras mašīnu analogiju ietekmes. Oriģinālo MI darbu atspoguļošanu parāda uzvedības psiholoģiskā teorija. Bet psiholoģijai ir nepieciešama uzvedības un darbības jēdzienu sadalīšana. Sistēmu analīzes populārās idejas atļāva padarīt mākslīgo sistēmu darbības principu salīdzināšanu ar cilvēka darbību par svarīgo heuristisko paņēmieni kā izdalīt tieši cilvēka darbības specifisko psiholoģisko analīzi.

Vēl 1965. gadā *A.N. Leontjevs* [110] formulēja sekojošu pozīciju: mašīna atdarina cilvēka domāšanas operācijas un, tāpat, arī „mašīnas” un „nemašīnas” ir operacionālā un neoperacionālā savstarpējā attiecība cilvēka darbībā. Tajā laikā šis secinājums bija diezgan progresīvs un bija virzīts pret kibernetisko redukcionismu. Taču turpmāk operāciju, no kurām tiek sastādīta mašīnas darbība, un operāciju kā cilvēka darbības vienību salīdzināšanā izdalījās būtiskas atšķirības. Psiholoģiskajā ziņā „operācija” apzīmē rezultātu sasniegšanas veidu, procesuālo raksturlīkni, kamēr, piemērojot to mašīnas darbībai, šis termins tiek izmantots loģisko-matemātiskajā ziņā (raksturojas ar rezultātu).

Darbos MI jomā vienmēr tiek izmantots termins „mērķis”. Līdzekļu attiecības pret mērķi analīzi *A. Njūels* un *G. Saimons* sauc kā vienu no „heuristikām”. Psiholoģiskajā darbības teorijā „mērķis” ir darbības pazīme atšķirībā no operācijām (un darbības kopumā). Mākslīgajās sistēmās par „mērķi” sauc kādu beigu situāciju, uz kuru tiecas sistēma. Šīs situācijas pazīmēm ir jābūt precīzi noskaidrotām un aprakstītām formālā valodā. Cilvēka darbības mērķiem ir cita daba. Beigu situāciju subjekts var dažādi atspoguļot: kā saprotamā līmenī, tā arī priekšstatu vai aperceptīvā tēla formā. Šis atspoguļojums var raksturoties ar dažādu skaidrības pakāpi. Bez tam, cilvēkam ir raksturīga ne tikai gatavo mērķu sasniegšana, bet arī jaunu veidošana. Arī

mākslīgā intelekta sistēmu darbs raksturojas nevis vienkārši ar operāciju, programmu, „mērķu” esamību, bet ar vērtības funkcijām.

Arī mākslīgajām sistēmām ir īpašas „vērtības operācijas”. Bet cilvēka motivācijas-emocionālās darbības regulācijas specifiku veido ne tikai konstanto vērtību izmantošana, bet arī noteiktā situācijā dinamiski mainījušos vērtību izmantošana, ir būtiska arī atšķirība starp vārdiski-loģiskām un emocionālām vērtībām. Vajadzību un motīvu esamībā ir redzama atšķirība starp cilvēku un mašīnu darbības līmenī. Šī tēze bija par iemeslu pētījumu ciklam cilvēka darbības specifikas analīzēs. Tā *L. P. Gurjevas* darbā [111] ir parādīta domāšanas darbības struktūras atkarība no motivācijas izmaiņas radošo uzdevumu risināšanā.

Starp citu, tieši mērķu veidošanas procesa nepietiekoša izpētīšana atrada atspoguļojumu „sociālā pasūtījuma” formulēšanā. *P. Simonova* informācijas teorija [112] arī lielā mērā izmanto analogijas ar MI sistēmu darbiem. Bez tam gribas lēmuma pieņemšanas problēma psiholoģijā dažos darbos tiek apskatīta kā vienas alternatīvas izvēle no daudzām, tādā veidā izlaižot gribas procesu specifiku. Tajā pašā laikā, *J. Babajeva u.c.* [113] un *T. Kornilova* [114] mēģināja izpētīt mērķa veidošanas procesa formalizācijas iespējas, pamatojoties uz šī procesa dziļo psiholoģisko analīzi cilvēka darbībā. Visas trīs tradicionālās psiholoģijas jomas mācības par izziņas, emocionāliem un gribas procesiem izrādījās MI darbu ietekmē, kas sekmēja jaunā psiholoģijas priekšmeta kā zinātnes par informācijas pārstrādi, šīs definīcijas zinātniskums tika sasniegts ar psiholoģisko zināšanu „industrializāciju”.

Pievēršoties MI lomas problēmai apmācībā, šo procesu var apskatīt kā cilvēka savstarpējās iedarbības veidu ar ESM, kas atklāj no perspektīvām iespējām tos, kuri ir vērsti uz tā saukto „adaptīvo apmācošo sistēmu” radīšanu, kuras imitē „skolnieka” un „skolotāja” operatīvo dialogu. Tādā veidā savstarpējās mijiedarbības loma starp mākslīgā intelekta pētījumiem un psiholoģisko zinātni var raksturot kā auglīgu dialogu, kas ļauj ja ne rēķināt, tad kaut vai iemācīties uzdot jautājumus – gan augstā filozofiskā līmeņa „Kas ir patiesība?” vai „Kas ir cilvēks?”, gan pragmatiskākos jautājumus – metodiskos un metodoloģiskos.

13. Mākslīgā intelekta sistēmu uzbūves principi

Pastāv dažādas pieejas MI sistēmu veidošanai. Šis iedalījums nav vēsturisks, kad viens viedoklis pakāpeniski nomaina citu, un dažādas pieejas eksistē arī tagad. Bez tam, tā kā īsti pilno MI sistēmu pašlaik nav, tad nevar teikt, ka kāda no pieejām ir pareiza, bet kāda nepareiza. No sākuma apskatīsim loģisko pieeju. Kāpēc tā parādījās? Cilvēks taču nodarbojas ne tikai ar loģisko spriešanu.

Protams, šis izteiciens ir pareizs, bet tieši prasme loģiski spriest ļoti atšķir cilvēku no dzīvniekiem. Šādas loģiskās pieejas pamats ir *Bula* algebra. Katrs programmētājs to pazīst, kā arī pazīst loģiskos operatorus no tā laika, kad viņš studēja IF operatoru. Bula algebra ieguva savu tālāko attīstību predikātu aprēķinu veidā – kurā tā ir paplašināta, pateicoties priekšmeta simbolu ieviešanai, attiecībām starp tiem, eksistences un vispārīguma kvantoriem. Praktiski katra uz loģiskā principa veidota MI sistēma ir teorēmu pierādīšanas mašīna. Pie tam, izejas dotie tiek uzglabāti datu bāzē aksiomu veidā, loģiskā secinājuma noteikumi – kā attiecības starp tiem. Bez tā, katrai tādai mašīnai ir mērķa ģenerācijas bloks, un izvada sistēma cenšas pierādīt mērķi kā teorēmu. Ja mērķis ir pierādīts, tad pielietoto noteikumu trasēšana ļauj iegūt darbību virkni, kuras ir nepieciešamas ieceltā mērķa realizācijai. Tādas sistēmas jaudu noteic mērķu ģenerators iespējas un teorēmu pierādīšanas mašīna. Protams, var teikt, ka izteiksmju algebras izteiksmīguma nepietiks MI pilnai realizācijai, bet ir jāatceras, ka visu pastāvošo ESM pamats ir bits – atmiņas vienība, kura var pieņemt tikai vērtības 0 un 1. Tādā veidā būtu loģiski pieņemt, ka viss, ko var realizēt ar ESM palīdzību, varētu realizēt arī predikātu loģikas veidā.

Bet šeit nekas nav pateikts par laiku. Sasniegt lielāku izteiksmīgumu loģiskajai pieejai ļauj tāds salīdzinoši jauns virziens, kā neskaidra faziloģika. Tās būtiskā atšķirība ir tā, ka izteiciena patiesīgums var tajā pieņemt atskaitot jā/nē (1/0) vēl starpvērtības – nezinu (0,5), pacients ir drīzāk dzīvs, nekā miris (0,75), pacients ir drīzāk miris, nekā dzīvs (0,25). Šāda pieeja ir vairāk līdzīga cilvēka domāšanai, jo tā reti atbild uz jautājumiem tikai „jā” un „nē”. Tomēr eksāmenā tiks pieņemtas tikai atbildes no klasiskās Buleva algebras. Vairākamam loģisko metožu ir raksturīga liela darbietilpība, jo pierādījuma meklēšanas laikā ir iespējama pilnā variantu apskatīšana.

Tāpēc šī pieeja prasa skaitļošanas procesa efektīvu realizāciju un labs rezultāts parasti tiek garantēts, ja datu bāzes apjoms ir salīdzinoši neliels. Runājot par struktūras pieeju, mēs saprotam MI veidošanas mēģinājumus, modelējot cilvēka smadzeņu struktūru. Viens no pirmajiem tādiem mēģinājumiem bija *Frenka Rozenblata* [1, 67] perceptrons. Modelēšanas pamatstruktūras vienība perceptronos (kā arī daudzos citos smadzeņu modelēšanas variantos) ir neirons. Vēlāk parādījās arī citi modeļi, kuri ir pazīstami kā „neironu tīkli” (NT). Šie modeļi atšķiras pēc atsevišķo neironu uzbūves, pēc to starpsaišu topoloģijas un pēc apmācības algoritmiem [63-65, 69]. Starp pazīstamākajiem NT variantiem ir NT ar atgriezenisku kļūdas izplatīšanu, *Hopfilda* tīkli [56], stohastiskie neironu tīkli. NT veiksmīgāk tiek pielietoti tēlu pazīšanas uzdevumos. NT modeļiem, kas ir izveidoti pēc cilvēka smadzeņu darbības analogijas ir raksturīga viena īpašība, kas ļoti tuvinā tos cilvēka smadzenēm – neironu tīkli darbojas pat gadījumā ja ir nepilna informācija par pētāmo objektu vai vidi, t.i. līdzīgi kā cilvēks tie var atbildēt uz jautājumiem ne tikai „jā” vai „nē”, bet arī „īsti nezinu, bet drīzāk jā”. Diezgan plaši izplatījās arī evolūcijas pieeja. MI sistēmu uzbūve pēc šīs pieejas, pamatzmanība tiek veltīta sākuma modeļa izveidošanai un tādiem noteikumiem, pēc kuriem tas var attīstīties. Pie tam modelis var būt izveidots pēc visdažādākajām metodēm, tas var būt arī NT un loģisko likumu kopums un jebkurš cits modelis. Pēc tam ieslēdz datoru un tas atlasa labākos modeļus un, pamatojoties uz tiem, pēc dažādiem likumiem tiek ģenerēti jaunie modeļi, no kuriem atkal tiek atlasīti vislabākie utt.

Šeit jāatzīmē, ka evolūcijas modeļi kā tādi neeksistē, eksistē tikai evolūcijas apmācības algoritmi, bet pēc evolūcijas pieejas iegūtiem modeļiem ir dažas raksturīgas īpašības, kas ļauj izdalīt tos atsevišķā klasē. NT izstrādātāja funkcija tiek pārcelta no modeļa izveidošanas uz tās modifikācijas algoritmu un tas, ka iegūtie modeļi praktiski neveicina jaunu zināšanu iegūšanu no vides, kurā ir MI sistēma, tā kļūst it kā par *I. Kanta* „lietu par sevi”. Vēl viena plaši izmantojama pieeja MI sistēmu uzbūvei – imitācijas pieeja.

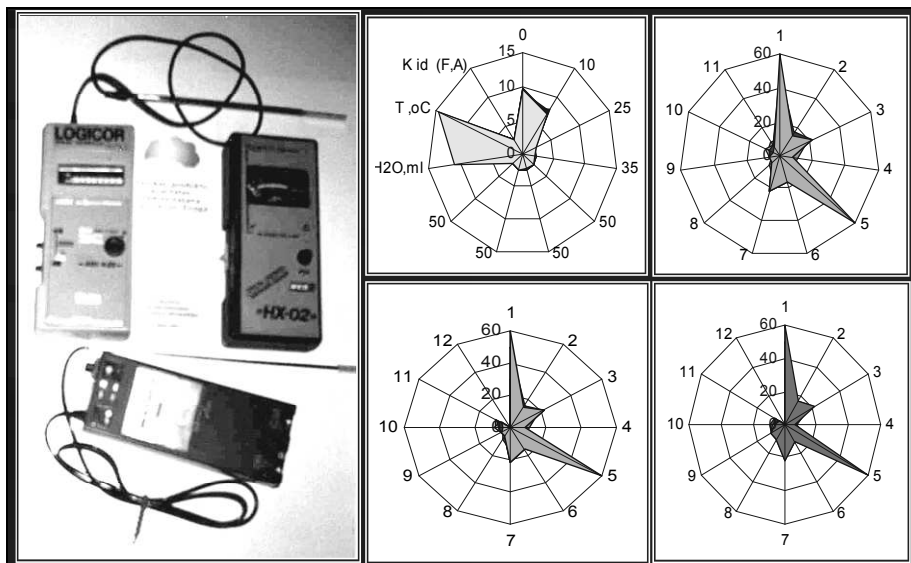


15. att. Intelektuālās mērīšanas ierīces pārtikas un lauksaimniecības produkcijas kvalitātes kontrolei (avots: Moskvins, 2007)

Šī pieeja kibernetikā ir klasiska. Pieejas pamatā ir *V. Ešbija* [30, 31] „melnās kastes” princips. „Melnā kaste” parasti ir ierīce, programmas modulis vai doto datu kopa, kur informācija par iekšējo struktūru un saturu pilnīgi neeksistē, bet ir zināmas ieejas un izejas datu specifiskācija. Objekts, kura uzvedība tiek imitēta, arī ir šī „melnā kaste”. Mums nav svarīgi, kas modelim ir iekšā un kā tas funkcionē, galvenais, lai mūsu modelis analogiskās situācijās uzvestos tāpat. Tādā veidā šeit tiek modelēta cita cilvēka īpašība – spēja kopēt to, ko dara citi, neiedziļinoties detaļās, kādēļ tas ir vajadzīgs. Bieži prasme kopēt ekonomē cilvēkam daudz laika, it īpaši dzīves sākumā. Imitācijas pieejas būtisks trūkums ir vairuma modeļu, kas izveidoti ar tās palīdzību, zemā informatīvā spēja.

Ar „melno kasti” ir saistīta viena ļoti interesanta ideja. Vai jūs gribētu dzīvot mūžīgi? Un vispār nav pārliecības, ka visi uz šo jautājumu atbildēs pozitīvi. Iedomāsimies, ka mūs novēro kaut kāda ierīce, kura seko tam, kā mēs rīkojamies kādās situācijās, ko sakām. Notiek novērošana pēc parametriem, kas nonāk mūsu ieejā

(redze, dzirde, garša, tauste utt.) un pēc parametriem, kuri iziet no mums (runa, kustība u.c.). Tādā veidā cilvēks šeit ir tipiska „melnā kaste”.



16. att. LLU „mākslīgās mēles” Logicor - AT intelektuālās ierīces atbilstības kontroles tehnoloģiju nodrošināšanai pārtikas preču aprītē (avots: Moskvins, 2004)

Tālāk šī ierīce mēģina atveidot kaut kādu modeli tādā veidā, lai iegūstot noteiktus signālus cilvēka ieejā, tā izdotu izejā tos pašus datus, kā cilvēks. Ja šī iedomā tika kaut kad realizēta, tad priekš visiem vērotājiem tāds modelis būtu tāda pati personība, kā reāls cilvēks. Bet pēc viņa nāves tā izteiks tās domas, kuras, domājams, izteiktu arī nomodelēts cilvēks. Mēs varam iet tālāk, nokopēt šo modeli un iegūt brāļi-dvīni ar tādām pašām „domām”.

Var teikt, ka tas viss, protams, ir interesanti, bet kā tas mani skar? Taču šis modelis tikai priekš citiem būs mans „es”, bet tā iekšā būs tukšums. Tiek kopēti tikai ārējie atribūti, bet es pēc nāves jau nedomāšu, mana apziņa nodzisis (ticīgiem cilvēkiem vārdu „nodzisis” jānoņem uz „pametīs šo pasauli”). Tas tā ir. Bet pamēģināsim iet tālāk. Saskaņā ar filozofiskajiem priekšstatiem, apziņa ir salīdzinoši neliela virsbūve virs mūsu zemapziņas, kura vēro galvas smadzeņu dažu centru

aktivitāti, tādu, kā runas centru, redzes tēlu apstrādi, pēc kā „atgriež” šos tēlus dotās informācijas apstrādes sākuma pakāpēs [45, 46, 48]. Pie tam, notiek šo tēlu atkārtota apstrāde, mēs it kā redzam un dzirdam, ko domā mūsu smadzenes. Pie tam, rodas apkārtējās realitātes algoritmu modelēšana ar mūsu it kā aktīvo piedalīšanos šajā procesā. Un tieši šo mūsu centru vērošanas process ir tas, ko mēs saucam par apziņu.

Ja mēs „redzam” un „dzirdam” mūsu domas, tas nozīmē, ka mēs esam apziņā, ja nē, tad mēs atrodamies bezapziņas stāvoklī. Ja mēs varētu modelēt tieši šo nedaudzo „apziņas” nervu centru darbību (kuru darbības pamatā ir visu pārējo smadzeņu darbība) vienas „melnās kastes” veidā, un „supervizora” darbību citas „melnās kastes” veidā, tad varētu pārliecināti teikt, ka šis modelis domā, pie tam, tāpat kā es.

Ļoti bieži ir sastopamas jauktas sistēmas, kur daļa no darba tiek veikta pēc viena tipa, bet daļa – pēc cita. Kas attiecas uz loģiskās domāšanas modelēšanu, tad kā labs modeļa uzdevums te varētu kalpot teorēmu pierādīšanas automatizācijas uzdevums. Sākot no 1960. gada tika izstrādāta programmu virkne, kas spēj atrast teorēmu pierādījumus pirmās kārtas predikātu aprēķināšanā. Šīm programmām pēc amerikāņu speciālista MI jomā Dž. *Makkatija* vārdiem ir „veselais prāts”, t.i. spēja veikt deduktīvos secinājumus. *K.Grīna* u.c. programmā, kas realizē jautājumu-atbilžu sistēmu, zināšanas tiek pierakstītas predikātu loģikas valodā aksiomu kopas veidā, bet jautājumi, kuri uzdod mašīnai, tiek formulēti kā pierādāmās teorēmas.

Lielu interesi izsauc amerikāņu matemātiķa *Hao Vanga* „intelektuālā” programma. Šī programma IBMp704 trīs minūšu darbības laikā atrisināja 220 relatīvi vienkāršu lemmu un teorēmu no fundamentālās matemātiskās monogrāfijas un pēc tam 8,5 minūšu laikā sniedza vēl 130 sarežģītāku teorēmu pierādījumus, daļu no kurām matemātiķi vēl nebija neatrisinājuši. Bet patiesībā, līdz šim neviena programma nebija atrisinājusi un pierādījusi nevienu teorēmu, kura, būtu ļoti nepieciešama matemātiķiem un būtu principiāli jauna. Robottehnika ir ļoti liela IST joma. Kur ir robota intelekta būtiskā atšķirība no universālo skaitļošanas mašīnu intelekta? Lai atbildētu uz šo jautājumu, varētu atcerēties izteicienu, kas pieder slavenajam krievu fiziologam *I. M. Sečenovam*: „visa smadzeņu darbības ārējo izpausmju bezgalīgā daudzveidība aprobežojas tikai ar vienu parādību – muskuļu kustību”.



17. att. **Autosaliekamo robotu sabiedrības izveidošanās**
(avots: EU-IST projekts, Šveice, 2003)

Citiem vārdiem sakot, visa cilvēka intelektuālā darbība beigu beigās ir vērsta uz aktīvo savstarpējo iedarbību ar ārējo pasauli ar kustībām. Tieši tāpat robota intelekta elementi kalpo pirmām kārtām viņa mērķtiecīgo kustību organizācijai. Tajā pat laikā tīri datorizēto MI sistēmu pamatzdevums ir intelektuālo uzdevumu risināšana, kam ir abstrakts vai palīgraksturs, kas parasti nav saistīti ar apkārtējās vides uztveršanu ar mākslīgo maņu orgānu palīdzību, ar izpildmehānismu kustību organizāciju.

Pirmos robotus ir grūti nosaukt par intelektuāliem. Tikai 60-tajos gados parādījās roboti, kurus vadīja universālie datori. Piemēram, 1969. gadā Elektrotehniskajā laboratorijā (Japāna) sākās projekta „rūpnieciskais intelektuālais robots” izstrādāšana. Šīs izstrādāšanas mērķis – jutekliska manipulācijas robota radīšana ar mākslīgā intelekta elementiem, lai izpildītu montāžas darbus ar vizuālo kontroli. Robota manipulatoram ir sešas brīvības pakāpes un tas tiek vadīts ar mini-ESM NEAC-3100 (operatīvās atmiņas apjoms 32 000 vārdi, ārējās atmiņas apjoms uz magnētiskajiem diskiem 27 3000 vārdi), kura veido nepieciešamo programmas kustību, kas tiek atstrādāta ar elektrohidraulisko sekošanas sistēmu. Manipulatora satvērējs ir aprīkots ar taustes adapteriem. Kā redzes uztveršanas sistēma tiek izmantotas divas televīzijas kameras, kas aprīkotas ar sarkani-zaļi-ziliem filtriem priekšmetu krāsas atšķiršanai. Televīzijas kameras redzes lauks ir sadalīts 64x64 šūnās. Iegūtās informācijas apstrādes rezultātā rupji tiek noteikta zona, kuru aizņem robota interesējošs priekšmets. Tālāk, lai detalizēti izpētītu šo priekšmetu, izdalītā zona atkal tiek sadalīta 4096 šūnās. Gadījumā, ja priekšmets neievietojas izvēlētajā

„lodziņā”, kamera automātiski pārvietojas, līdzīgi tam, kā cilvēka skatiens slīd pa priekšmetu.



18. att. **Autosaliekamo robotu sabiedrība kopīgi vieglāk pārvar šķēršļus**
(avots: EU-IST projekts, Šveice, 2003)

Elektrotehniskās laboratorijas robots bija spējīgs atšķirt parastos ar plaknēm un cilindriskām virsmām norobežotus priekšmetus speciālā apgaismojumā. Šā eksperimentālā parauga vērtība bija apmēram 400 000 ASV dolāri. Pakāpeniski robotu raksturlīknes monotoni uzlabojās, bet līdz šim tie vēl ir tālu no cilvēka saprātīguma ziņā, kaut gan dažas operācijas viņi jau izpilda labāko žonglieru līmenī. Vēl šeit var pievērst uzmanību *N. M. Amosova* un *V. M. Gluškova* darbiem [8, 39, 115].

Pasaulē tiek veikta virkne pētījumu, kuri ir vērsti uz robotu intelekta elementu izstrādāšanu. Šajos pētījumos īpaša uzmanība ir veltīta attēlu un runas atšķiršanas, loģiskā secinājuma (teorēmu automātiskās pierādīšanas) un vadīšanas problēmām ar neironu līdzīgiem tīkliem. Šā ļoti īsā apskata beigās aplūkosim dažus lielmēroga ekspertu sistēmu piemērus [29, 54]. MICIN – ekspertu sistēma medicīnas diagnostikai. To izstrādāja Stenfordas universitātes grupa, kura nodarbojas ar infekciju slimībām. Izejot no ievadītajiem simptomiem, nosaka diagnozi un rekomendē medikamentu dziedniecības kursu jebkurai no diagnosticētajām infekcijām. Datu bāzi veido 450 likumi. PUFF – elpošanas traucējumu analīze. Šīs sistēmas pirmās versijas parādījās vēl 1965. gadā Stenfordas universitātē. Lietotājs dod DENDRAL sistēmai kādu informāciju par vielu, kā arī spektrometrijas datus (infrasarkanais starojums, kodolu magnētiskā rezonanse un masspektrometrija), un tā, savukārt, nosaka diagnozi attiecīgās ķīmiskās struktūras veidā. PROSPECTOR – ekspertu sistēma, kura ir radīta

izrakteņu komerciāli attaisnoto atradņu meklēšanas atbalstam. Profesora *Duglasa Lenāta* un citu *Stenfordas universitātē* izstrādātā metode EURISKO ir paredzēta jaunu iespēju izpētei *zinātnisko zināšanu* laukā. Iegūtās zināšanas balstās uz heuristiku. Tas ir tāds zināšanu iegūšanas paņēmiens, kas prasa, lai iespējamie notikumi sekotu neiespējamajiem. EURISKO metode attīsta labvēlīgos notikumus, uzlabo iekšējos modeļus un uzlabo likumus iekšējo modeļu atlasei. Lenāts apraksta heuristikas koncepcijas paveidu terminos „mutācija” un „atlase” un piedāvā sociālās metaforas to savstarpējās iedarbības saprašanai. No tiem laikiem heuristikas un EURISKO metodes attīstās un konkurē. Taču var sagaidīt, ka parādīsies metodiskie surogāti un parazīti no heuristikas un tam patiešām ir nozīme. Piemēram, kāda no mašīnas heuristikas ir pacēlusies līdz pašam augstākajam vērtību pašnovērtējumam, un pretendē uz to, lai kļūtu par jaunas vērtīgas idejas pirmatklājēju. Profesoram *Lenātam* strādājot ar EURISKO ir izdevies paaugstināt aizsardzības līmeni no surogātiem un pārtraukt heuristikas attīstību pa nederīgu un tukšu domu ķēdi. EURISKO izmanto elementārās matemātikas uzdevumos, programmēšanā, bioloģiskās attīstības uzdevumos, spēļu teorijā, trīsdimensiju integrālo shēmu veidošanas programmās, naftas plankumu aizvākšanas projektos, montāžas darbos un, protams, heuristikā. Dažās jomās heuristikas pielietošana pārspēja visas cerības un pārsteidza pašus projektētājus ar jaunām idejām, tai skaitā jauno projektēšanas tehnoloģiju un trīsdimensiju integrālo sistēmu radīšanas laukā. Heuristikās programmas spēku mums uzskatāmi parāda cilvēka un mākslīgā intelekta sacenšanās rezultāti. Piemēram, vienā jūras kājnieku futuristikajā datorspēlē MI spēlēja saskaņā ar 200 lpp. likumu, kas nosaka spēles stratēģiju, izdevumus un ierobežojumus flotes darbībās. Tiek paziņots, ka aptuveni 60% spēļu uzvar cilvēks un aptuveni 40% uzvar EURISKO.

EURISKO un citas MI programmas parāda, ka datori nedrīkst būt ierobežoti iespējā izdarīt atkārtotu pareizo soli, tas ir, mums jādod viņiem iespēja laboties. Heuristikā meklēšanā viņi uzlabo paši sevi, ja ir uzdots pareizais heuristikās programmēšanas veids. Tā var izpētīt potenciālās iespējas un radīt pavisam jaunas idejas, kas pārsteidz pat to radītājus. Turpmāk inženieris, sēžot pie monitora, ievadīs datorā tikai projektēšanas mērķus un izdarīs piedāvāto projektu uzmetumu skicējumus. MI sistēma ar viņu sadarbosies, filtrējot un koriģējot projektus, labos tos un piedāvās

alternatīvas ar skaidrojumiem, ailēm un diagrammām. Inženieris tālāk izdarīs priekšlikumus un izmaiņas, un atbildēs ar jauna uzdevuma uzdošanu, kamēr aparāta līdzekļu sistēma ESM netiks izveidota pilnīgā saskaņā ar modeli. Automatizētās tehniskās sistēmas nepārtraukti tiek uzlabotas, tās izpildīs arvien lielāku darba apjomu ar lielāku ātrumu. Aizvien biežāk inženieris vienkārši piedāvās mērķus un izvēlēsies labāko variantu no labu risinājumu variantiem, kurus piedāvā domājošā mašīna. Aizvien retāk viņam vajadzēs izvēlēties detaļas, daļas, materiālus, formas un konfigurācijas. Pakāpeniski inženieri būs spējīgi piedāvāt vairāk apkopoto mērķu daudzumu un var gaidīt no domājošās mašīnas aizvien vairāk labu risinājumu. Kā EURIKSO izgudroja jaunas elektroniskas ierīces, tā arī nākotnes automatizētās tehniskās sistēmas būs spējīgas izgudrot jaunas molekulāras mašīnas un molekulāras elektroniskas ierīces. Intelektuālās automatizētās sistēmas ar savu programmatūru pašas sev palīdzēs jaunu molekulāro modeļu radīšanā, piemēram, nanotehnoloģijas jomā. Galu galā, programmatūras sistēmas būs spējīgas radīt jaunus projektus bez cilvēku palīdzības [108]. Jautājums par to, vai vairācumam cilvēku būs vajadzīgas šādas intelektuālās sistēmas, īstenībā nav būtisks. Pats automātiskās projektēšanas process ar pat nelielu zinošu cilvēku daudzuma līdzdalību atvēra ceļu principiāli jaunajām 21. gadsimta tehnoloģijām.

14. Domājošu mašīnu veidi

Pēc savas būtības adaptācijas procesi ir optimizācijas procesi (Dž. Hollands, „Adaptation in natural and artificial systems”)

Ir iespējami vairāki IST un MI konstruēšanas varianti [1, 4, 7, 8, 30, 34, 61, 107]. Viens no variantiem, kā izveidot domājošu mašīnu (DM), var būt automātiski mehānismi uz dažādas bāzes (tehniskais, bioloģiskais pamats). Sākotnēji DM veidotājs konstruē esošam pasaules uzskatam atbilstošu DM modeli. DM projektētājs var ietekmēt sistēmas procesus, izmainot pasaules modeli vai izmainot tā atsevišķas daļas. Tādā veidā DM veidotājs var apmācīt, stimulēt, kontaktēt ar modelējamo radību, saņemt atbildes. Tas ir acīmredzams, ka DM veidotājam nav vajadzības mainīt ne

robotu prātu, ne ķermeni, ne principus, pēc kādiem robots mijiedarbojas ar pasauli. DM ķermenis darbojas kā informācijas nesēju tilpne. DM „nāve” iestājas tad, kad beidzas informācijas apmaiņa sistēmas iekšienē. Cits DM konstruēšanas veids ir variants, kad tiek modelēts pats DM ķermenis, apkārtējā vide un mijiedarbība starp tiem. Šāda pieeja uzliek lielus ierobežojumus saprātīgas sistēmas iespējām. Dažādi pētnieki cenšas piedot domājošai mašīnai cilvēka domāšanas prototipus, taču tās domā savādāk, nekā cilvēki. Piemēram, robotam galīgi nav vajadzīgs jēdziens „siltums”. Pavisam cita lieta, kad robots asociē vārdu „siltums” ar temperatūras izmaiņu. Vispārīgi jebkuru cilvēka jēdzienu vai attēlu var parādīt citai saprātīgai sistēmai ar sakārtotu daudzu iekšēju tēlu sakarību. Tāpēc ir bezjēdzīgi ievadīt domājošā mašīnā cilvēka jēdzienus, domājošai mašīnai tas jādara patstāvīgi.

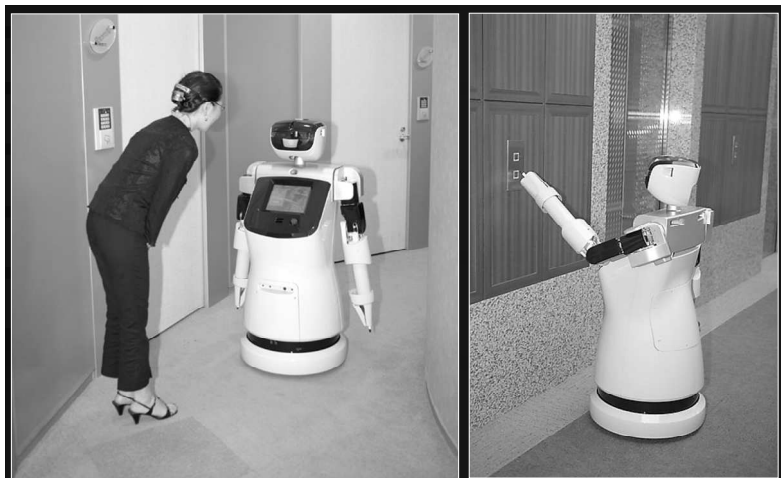
15. Domas vēsture

Intelekti – ir likumu spēja iztaisnoties (M. Geny)

Cilvēki vienmēr ir interesējušies par savu domāšanu. Šī smadzeņu domāšana pati par sevi, iespējams, ir tikai cilvēkam raksturīga īpatnība. Pastāv daudz pārdomu par domāšanas dabu, sākot ar garīgo un beidzot ar anatomisko aspektu. Šā jautājuma apspriešana, kas pagāja karstos strīdos starp filozofiem un teologiem no vienas puses un fiziologiem un anatomiem no otras puses, nedeva daudz labuma, jo pats priekšmets ir diezgan sarežģīts pētīšanai. Tie, kuri balstījās uz pašanalīzi un apdomāšanu, ieguva slēdzienus, kas neatbilst fizikas zinātņu noteiktības līmenim. Taču eksperimentētāji atklāja, ka smadzenes ir grūti novērojamas un tās iedzen strupceļā ar savu struktūru un organizāciju. Citiem vārdiem sakot, zinātnes pētījumu varenās metodes, kas mainīja mūsu priekšstatus par fizikālo realitāti, izrādījās bezspēcīgas paša cilvēka izprašanā. Neurobiologi un neuroanatomu sasniegta ievērojamo progresu. Cītīgi pētot cilvēka nervu sistēmas struktūru un funkcijas, viņi daudz ko izprata smadzeņu „elektriskajā instalācijā”, bet maz uzzināja par to funkcionēšanu.

Zināšanu uzkrāšanas procesā izrādījās, ka smadzenes ir pārsteidzoši sarežģītas. Simti miljardu neironu, katrs no kuriem ir savienots ar simtiem vai tūkstošiem citu,

veido sistēmu, kura pārsniedz mūsu visdrosmīgākos sapņus par superdatoriem. Neraugoties uz to, smadzenes pakāpeniski atklāj savus noslēpumus vienā no vissaspīlētākajiem un visgodkārīgākajiem pētījumu procesiem cilvēces vēsturē.



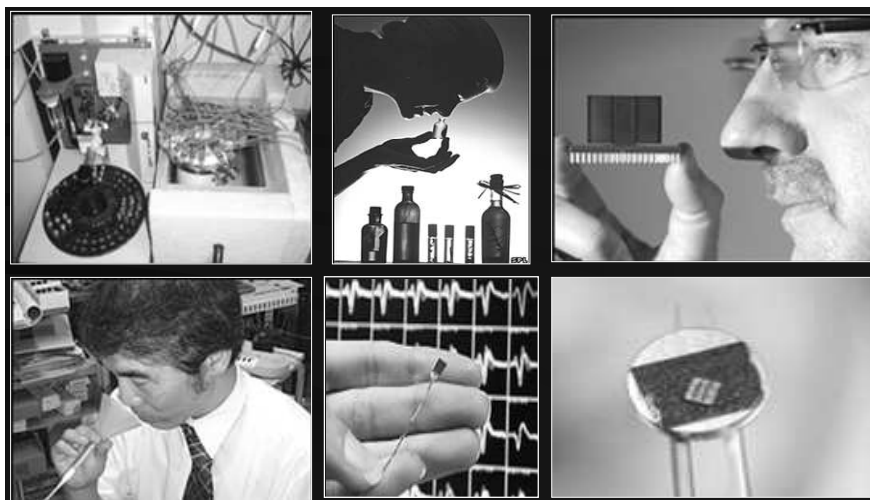
19. att. **Viesnīcas robots** (Japāna)

Labāka neirona funkcionēšanas un tā saišu zīmējuma izpratne deva pētniekiem iespēju radīt matemātiskos modeļus savu teoriju pārbaudīšanai. Eksperimenti tagad var tikt veikti uz digitālajiem datoriem bez cilvēka un dzīvnieku piesaistīšanas, kas risina daudzas praktiskās un morāli ētiskās problēmas. Jau pirmajos darbos noskaidrojās, ka šie modeļi ne tikai atkārtos smadzeņu funkcijas, bet spēj pildīt arī funkcijas, kurām ir sava vērtība. Tāpēc pašlaik parādījušies un paliek divi neironu modelēšanas mērķi, kuri abpusēji bagātina viena otru: pirmais – saprast cilvēka nervu sistēmas funkcionēšanu fizioloģijas un psiholoģijas līmenī un otrais – radīt skaitļošanas sistēmas (mākslīgie neironu tīkli), kuras pilda funkcijas, kas līdzīgas smadzeņu funkcijām. Paralēli ar progresu neuroanatomijā un neurofizioloģijā psihologi izveidoja cilvēka apmācības modeļus. Viens no tādiem modeļiem, kurš izrādījās auglīgāks, bija *D. Hebba* modelis. Viņš 1949. gadā piedāvāja apmācības likumu, kas kļuva par starta punktu mākslīgo neironu tīklu apmācīšanas algoritmiem.

Šodien, papildināts ar lielu skaitu citu metožu, tas nodemonstrēja tā laika zinātniekiem, kā neironu tīkls var apmācīties. Piecdesmitajos un sešdesmitajos gados pētnieku grupa, apvienojot šīs bioloģiskās un fizioloģiskās pieejas, izveidoja pirmos mākslīgos neironu tīklus. Sākotnēji izveidoti kā elektronu tīkli, vēlāk tie tika pārnesti datoru modelēšanas vidē, kura saglabājas arī šodien. Pirmie panākumi izsauca aktivitātes un optimisma sprādzienu. *Minskis, Rozenblats, Vidrovs* un citi izstrādāja tīklus, kuri sastāv no viena mākslīgo neironu slāņa. Bieži saukti par perceptroniem, tie tika izmantoti tādai plašai uzdevumu klasei kā laika apstākļu noteikšana, elektrokardiogrammu analīze un mākslīgā redze. Kādu laiku likās, ka intelekta atslēga ir atrasta un cilvēka smadzeņu atveidošana ir tikai pietiekoši liela tīkla konstruēšanas jautājums. Bet šī ilūzija drīz izklīda, jo tīkli nevarēja analogiski risināt citus ārēji līdzīgus uzdevumus. Ar šīm neizskaidrojamajām neveiksmēm arī sākās intensīvas analīzes periods.

Izmantojot precīzas matemātiskas metodes, *M. Minskis* korekti pierādīja virkni teorēmu, kas attiecas uz tīklu funkcionēšanu. Viņa pētījumi noveda pie grāmatas rakstīšanas, kurā viņš kopā ar *Peipertu* pierādīja, ka tajā laikā izmantojamie viena slāņa tīkli teorētiski nav spējīgi risināt daudzus vienkāršus uzdevumus, tai skaitā realizēt „*nē - vai*” funkciju. *M. Minskis* arī nebija optimistiski noskaņots attiecībā uz potenciāli iespējamu progresu. Perceptrons pierādīja, ka tam ir daudz pievilcīgu īpašību: linearitāte, saistoša apmācības teorēma, paralēlo aprēķinu modeļa vienkāršums. Nav pamata uzskatīt, ka šīs vērtības saglabāsies pārejā uz daudzslāņu sistēmām. Tomēr mēs uzskatām par svarīgu pētīšanas uzdevumu mūsu intuitīvās pārlicēības nostiprināšanu (vai atspēkošanu), ka tāda pāreja ir vēltīga. Iespējams, ka tiks atklāta kaut kāda varena teorēma par konverģenci vai tiks atrasts dziļš neveiksmju cēlonis, dodot interesantu „apmācības teorēmu” daudzslāņu mašīnām. *M. Minska* argumentācijas spožums un korektums, kā arī viņa prestižs radīja milzīgu uzticību grāmatai – tās secinājumi bija nevainojami. Pētnieki vīlušies atstāja pētījumu lauku un vēltīja sevi tādām jomām, kur varēja sagaidīt labākus rezultātus, bet valdības pārdalīja savas subsīdijas un mākslīgie neironu tīkli bija aizmirsti uz gandrīz diviem gadu desmitiem. Neskatoties uz to, daži neatlaidīgāki zinātnieki, tādi kā *Kohonens, Grosbergs, Andersons* pētījumus turpināja. Līdz ar trūcīgo finansējumu un

nepietiekošu vērtējumu virkne pētnieku piedzīvoja grūtības ar publikācijām. Tāpēc tie pētījumu rezultāti, kas bija publicēti 70-ajos gados un 80-to gadu sākumā, ir izkļādēti dažādos žurnālos, daži no tiem ir maz pazīstami.



20. att. Ekspertu sistēma „Mākslīgais deguns” (avots: *IfL-IfL*)

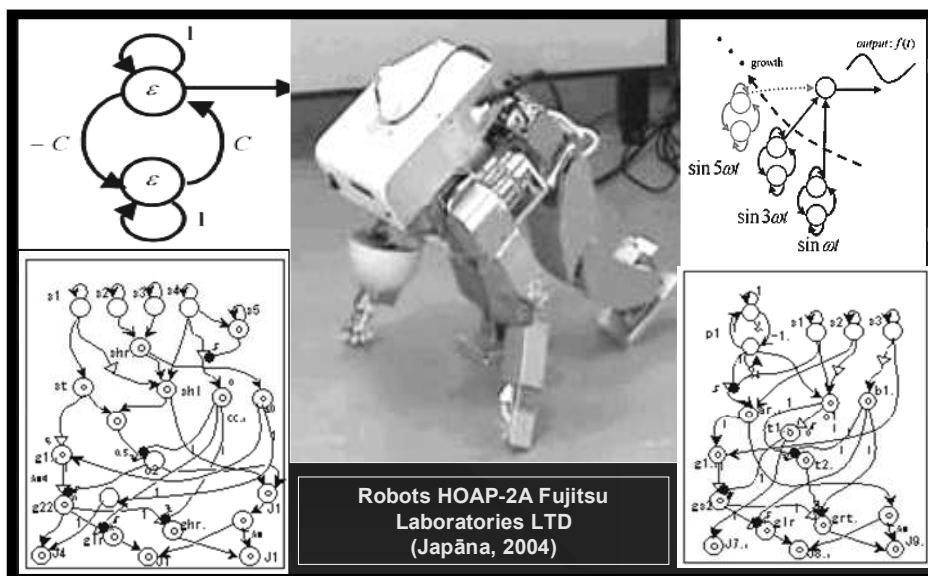
Pakāpeniski parādījās teorētiskais fundaments, uz kura pamata šodien tiek veidoti jaudīgāki daudzslāņu tīkli. *M. Minska* vērtējums bija pārāk pesimistisks, daudzus no viņa grāmatā izvirzītajiem uzdevumiem tagad risina tīkli ar standarta procedūru palīdzību. Dažu pēdējo gadu laikā NT teoriju sāka izmantot praktiskās jomās un parādījās jaunās korporācijas, kuras nodarbojas ar komerciālo šīs tehnoloģijas izmantošanu.

Domājošām mašīnām nav jābūt līdzīgām cilvēkiem dienesta formā, kas ir apveltīti ar dažām domāšanas spējām. Tik tiešām, dažām mākslīgā intelekta sistēmām ir tādas pašas īpašības, kā diplomētam speciālistam mākslu jomā, kurš spēj spēlēt varena dzinēja lomu dažādu projektu realizēšanā. Bet tomēr doma par to, ka cilvēka prāts ir cēlies no nesaprātīgas matērijas, zaudēs savu nozīmi, salīdzinot ar pārņēmušo ideju padarīt mašīnu par domājošu. Prāts tāpat kā citas kārtības (antihaosa) formas, ir

attīstījies (evolucionizējies) pateicoties mainīgumam un selektīvai izvēlei haosa un pretrunu pasaulē.

16. Mākslīgo neironu tīklu gūstā

Kas palīdz cilvēkam analizēt ienākošo informāciju? Neuroģenētikas tehnoloģijā ir iekļauts atslēgas jēdziens – neirotīkls. Tieši neirotīklu kopums veido cilvēka nervu sistēmu, kura savukārt nosaka cilvēka darbību, piešķir viņam prātu un intelektu. Ko nozīmē mākslīgie neironu tīkli? Kādas tiem ir iespējas? Kā tie darbojas? Kā tos var izmantot? Šos un arī citus jautājumus uzdod dažādu nozaru speciālisti. Mākslīgie neironu tīkli ir zināmi no bioloģijas. Tie sastāv no elementiem, kuru funkcionālās iespējas ir līdzīgas bioloģiskā neirona – nervu šūnas elementārajām funkcijām. Neironi organizējas pēc principa, kurš var atbilst (vai neatbilst) smadzeņu anatomijai [49, 50, 62-66].



21. att. Robota HOAP-2A (Japāna) vadības paradigma CPG / NP - Central Pattern Generator (CPG) / Numerical Perturbation Method (NPM)

Tie savienojas, veidojot ķēdes. Mākslīgiem neironu tīkliem ir raksturīgas daudzas īpašības, kuras piemīt smadzenēm. Tie mācās, pamatojoties uz pieredzi, apkopo iepriekšējos notikumus un iegūst būtiskas īpašības no ienākošās informācijas, kura satur jaunus datus. Mēs ieslēdzam informāciju vārdos, lai nodotu to citiem. Mākslīgajam intelektam vārdi nav vajadzīgi. Divas „smadzenes”, kas atrodas dažādās Zemeslodes vietās, var apmainīties ar informāciju caur pavadoni vai caur citiem kanāliem. Pie tam, tie neizmantos valodu. Tāpēc neirodatorus nevarēs uzskatīt par atsevišķiem subjektiem. Tie veidos vienotu veselumu. Liels daudzums „smadzeņu”, kuri ir apvienoti tīklā, kļūs par mākslīgo prātu. Viena komponenta bojāeja vai iznīcināšana kvalitatīvi neietekmēs visas sistēmas darbību. Taču arī mūsu smadzenes turpina darboties pēc ievērojama bojājuma. Onore de Balzaks pārcieta insultu, kas iznīcināja pusi no viņa smadzenēm, taču pēc tam viņš uzrakstīja dažus diezgan labus romānus. Ir zināmi gadījumi, kad cilvēks atgriezies pie apzinātās dzīves pēc operācijām, kurās viņam izgriezta smadzeņu garozas lielāko daļu.

Neskatoties uz funkcionālo līdzību, pat pats optimistiskākais neironu tīklu aizstāvis nevarēs iedomāties, ka tuvākajā nākotnē mākslīgie neironu tīkli dublēs cilvēka smadzeņu funkcijas. Reālais „inteleks”, kuru demonstrē vissarežģītākie neironu tīkli, atrodas zemāk par sliekas līmeni. Un entuziasmam ir jābūt mērenam attiecībā uz mūsu realitāti. Tomēr būtu nepareizi ignorēt apbrīnojamo līdzību starp dažu neironu tīklu un smadzeņu funkcionēšanu. Šīs iespējas, kaut arī šodien tās ir ierobežotas, liek domāt, ka drīz varēs dziļi iekļauties cilvēka intelektā un atklāt daudz jauna.

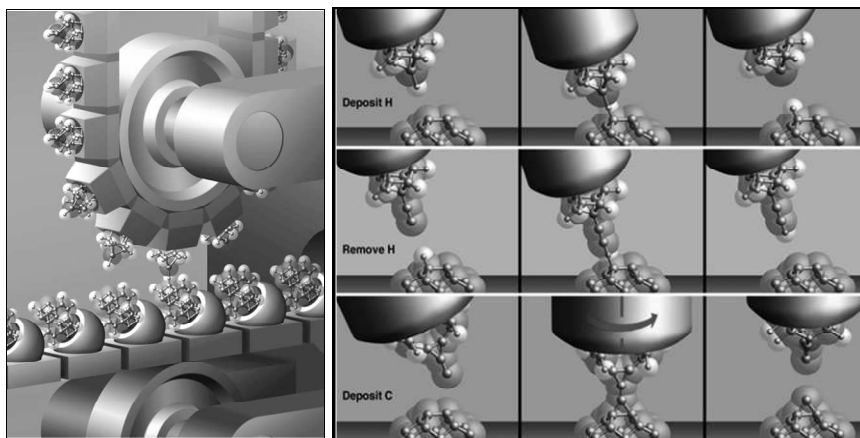
17. Nanoroboti

Pētnieciskā grupa, kura nodarbojas ar elektronisko ierīču un nanorobotu veidošanu, paziņojusi par tādu robotu radīšanu, kuru izmērs ir tuvs monētai un kuri spēj veikt desmit tūkstošus kustību minūtē. Projekts ir nosaukts *NanoWalker* un izstrādāts uz biotehnoloģiju laboratorijas bāzes Masačūsetsa Tehnoloģiju institūtā (*Massachusetts Institute of Technology*) *Silvina Martela* (*Sylvain Martel*) vadībā.

Patiesībā šos mehānismus nevar saukt par nanorobotiem (viens nanometrs ir vienāds ar milimetra miljono daļu) – tie ir pārāk lieli. Drīzāk tie ir mikroroboti, bet izmantosim projekta autoru terminoloģiju, jo vairāk tāpēc, ka šie roboti var darboties molekulārajā līmenī. Pēc paša *Martela* viedokļa viņa kolēģi vairāk nekā citi „nanokonstruktori” tuvojas pirmās efektīvās specializētās intelektuālās ierīces radīšanai, kura ir jau gatava pielietošanai. Vēl vairāk, *Martels* uzskata, ka tieši viņa laboratorijā parādīsies modeļi, kurus vēlāk izplatīs pa visu pasauli. Pirmā un galvenā šo nanorobotu īpašība ir tā, ka tie jau tagad spēj autonomi pārvietoties – tas ir „skriet un labot” bez vadiem. Šī raksturīgā īpatnība ir atspoguļota projekta nosaukumā, taču *Nano Walker* – nozīmē „nanogarāmgājējs, nanogājējs”. Robotu vadīšana notiek ar infrasarkanā adapteru, kas ir ievietots to „ķermeņos”. Kamera novēro to atrašanās vietu un virza uz uzdevuma izpildīšanas vietu. Tāpat autonomi *Nano Walker* var veikt savu darbu: dažos modeļos tiek izmantota speciāla pozicionēšanas sistēma, kas palīdz robotam noteikt daļiņu (piemēram, molekulu), ar kuru tiks veiktas manipulācijas. Daži maziņi „pētnieki” ir aprīkoti ar mikroskopiem, kuri ļauj tiem dabūt un translēt atoma attēlu, ar kuru vajadzēs pastrādāt (piemēram, pastumt, lai ar to aktivizētu kādu ķīmisko reakciju). Otra īpatnība – ir „*Martela robotu*” pārvietošanās ātrums un plakne.

Vairākums eksistējošo nanorobotu ir ļoti tālu no ideāla: tie ir ne tikai piespiesti strādāt tur, kur tos ievieto, bet kustas viņi ļoti lēni – viena kustība sekundē. Jaunā paaudze var pārvietoties trīs plaknēs. Trešā robotu īpatnība – manipulāciju sarežģītība: primitīvie roboti var veikt krāvēja kustības – no vienas vietas paņemt – citā nolikt. Elites „puiši” no Masačūsetsa prot risināt sarežģītus uzdevumus, kuri prasa ne tikai mehānisko piepūli. Pēc *Martela* vārdiem, viņa aizbilstamo tagadējā kvalifikācija pagaidām tikai tiek noteikta. Jau parādījušies modeļi, kuri var tikt izmantoti farmakoloģijā un īstenot ķīmisko preparātu un zāļu sintēzi. Bez tam, „mazie darba cilvēki” nodarbosies ar DNS darba „inspicēšanu”. Zinātnieks uzskata, ka prioritārā robotu pielietošanas sfēra ir izmainījusies – tagad tā nav automobiļu rūpniecība, bet gan bioloģija un farmācija. Bez tam, miniatūra *Martela* armija var veikt komplicētākus izskaitļojumus un kalpot kā visprecīzākās mērīšanas ierīces. Tuvākajos gados vienota bezvadu komanda varēs izpildīt līdz 200 tūkstošus precīzu kustību atomārā līmenī. Inženieri uzskata, ka jaunā tehnoloģija maina pašus principus medicīniskajā

diagnostikā. Projekta *Nano Walker* mērķis ir likt medicīnai aizmirst par jau sen novecojušo orgāna vai tā fragmenta pārvietošanas metodiku no viena mērīšanas aparāta pie otra. Tagad pats instrumentārijs miniaturizējas, lai uz kādu laiku kļūtu par orgāna daļu un veiktu analīzi „uz vietas”, tajā vidē, kurā slimība pastāv.



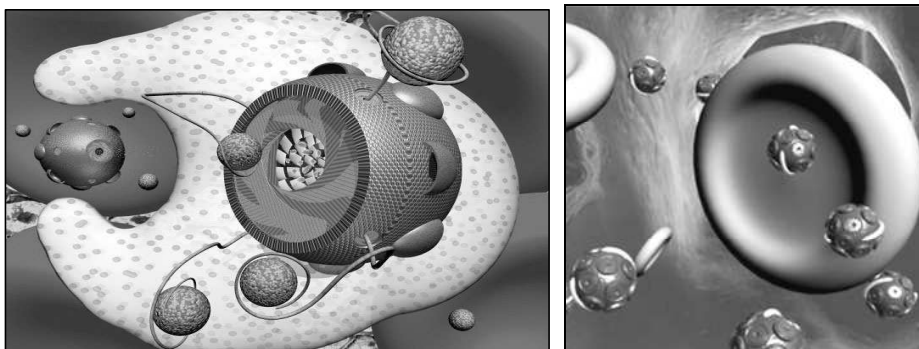
22. att. **Molekulu salikšanas automatizācija no atsevišķiem atomiem**
(avots: *Foresight Nanotech Institute, USA*)

Martela laboratorija nodarbojas arī ar citiem, arī ļoti interesantiem projektiem – piemēram, izstrādā datoru ar grauda izmēriem, kurš lēti maksātu, nebūtu energoietilpīgs, kā arī būtu savienojams ar citiem „pieaugušajiem” modeļiem. Šis projekts tā arī saucas *The I-Grain Project*. Cits projekts – elektroniskā čipa-implantanta radīšana, kas spēj stimulēt atsevišķo smadzeņu zonu darbu.

18. IST, MI... Kas tālāk?

Pasaules prese ļoti daudz raksta par to, ka nākotnes mākslīgais intelekts aizvietos cilvēku un aizņems galveno vietu planētas dzīvē. Publikācijas vairāk izskatās pēc stāstiem, tiek domāts, ka tas ir ilgs darbs, ka tas notiks pēc ilga laika un datortehnikai vēl ir jāattīstās, tā jāuzlabo. Jau ir izdomāta jauna datoru paaudze, kura ir

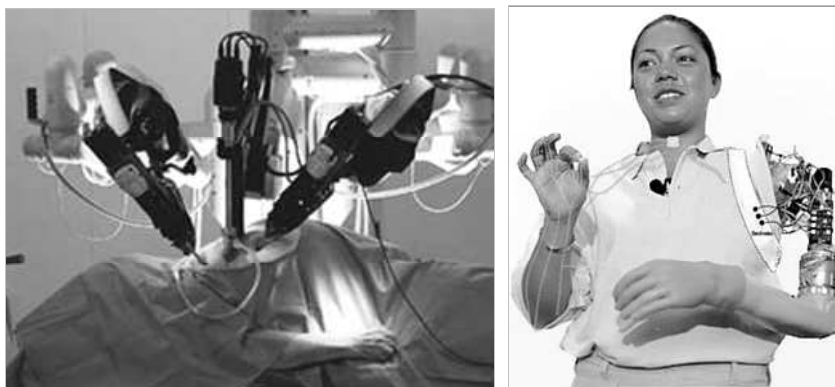
ar principiāli citu arhitektūru – neirodatori. Par šo datoru pirmo „upuri” kļūs ekonomika. Precīzāk, finanšu tirgus. Jo visi zinās par visu. Ja kāds jums pateiks, ka MI nav iespējams, palūdziet viņu paskatīties spogulī. Cilvēka intelekts jau sen kļuvis par mākslīgo, tāpēc, ka tas tiek veidots mākslīgā vidē, kurā mēs dzīvojam, ar mākslīgo pārtiku, kuru mēs lietojam uzturā, ar izdomātiem vārdiem, ar kādiem mēs aprakstām pasauli. Un ar katru dienu, ar katru tehnisku jaunumu, kam nepieciešama cilvēka vadība, ar katru sociālās sistēmas komplikāciju mūsu smadzenes kļūst arvien mākslīgākas, arvien loģiskākas. Smadzenes un dators atšķiras tikai ar to, ka Bet par to vēlāk. Tagad pēc kārtas. Pazīstamais Beļģijas zinātnieks *Hugo de Gari* (viens no kaķēna-robota *Robokoneko* radītājiem) saka, ka pēc 50–100 gadiem MI varēs veikt domāšanas operācijas 10 000 000 000 reizes ātrāk par cilvēku.



23. att. **Nanorobotu projekti pielietošanai medicīnā**
(avots: *Foresight Nanotech Institute, USA*)

Tādās „smadzenēs” katrs atoms kļūs par formālo neironu. Un to varēs uztaisīt tik lielu, cik vien gribas. Salīdzināšanai: cilvēka smadzenēs, kuras ir ierobežotas ar galvaskausu, „viena smadzeņu vienība” – neirons, kas ir šūna, kura sastāv no daudzām molekulām. „Nu un kas! - teiks *Holivudas* antiutopisku filmu scenāristi un režisori. Bet mašīnas nevar just, tām nav intuīcijas un vispār ...” Bet kas attiecas uz „vispār” - tas ir jājautā pasaku rakstniekiem un teologiem, taču runājot par emocijām un intuīciju, šeit vajadzēs tikt skaidrībā. Kas ir emocijas, pozitīvas un negatīvas – mīlestība, naidis, bailes, prieks, skumjas un tā tālāk? Kādēļ tās ir izveidojušās evolūcijas gaitā? Emocijas ir pātaga un medus rausis organismam. Atgriezeniska saite ar pasauli. Stimuli. Ja *Darvina* kungam ir taisnība, tad mēs sen jau nelēkājam pa zariem, bet bioķīmija

organismā ir tā pati un uzdevumi ir tādi paši – mīlēt un būt mīlētiem, būt par līderiem hierarhijā, garšīgi paēst un saldi pagulēt (tāpēc, ka mēs esam atkarīgi no seksa, sabiedrības, ēšanas un miega). Visas mūsu dzīvnieciskās tieksmes ir maskētas ar sociālo floru. Tas ir ar vārdu savārstījumu par pienākumu, godu, darbu, dzīves jēgu. Protams, ka emociju un jūtu mašīnai nebūs. Tikai tāpēc, kā tā ir savādāk iekārtota, ne bioķīmiski. Bet būs savādāki iekšējie impulsi, kas mudinās mašīnu darboties – elektroniskie signāli. Mašīnai tie ir ne mazāk svarīgi, kā mums – adrenalīns.



24. att. **Mūsdienu robotu pielietošana medicīnā** (avots: DAAD, IFL-LfL)

Tagad par intuīciju. Intuīcija ir netiešs veids, kā smadzenes apstrādā informāciju, kad galvā pēkšņi „uzpeld” atbilde, bet pati „aprēķināšana” paliek „aiz kadra”. Tas dažreiz šķiet apbrīnojami. Bet īstenībā gatavu atbilžu parādīšanās galvā ir process, ko var novērot daudz biežāk, nekā ir pieņemts uzskatīt. Cilvēks iet pāri ceļam, ierauga mašīnu un novērtē vai paspēs pāriet, vai arī labāk pagaidīt. Matemātiskais mašīnas un paša gājēja ātruma vektoru saskaitīšanas process gājēja galvā notiek nemanot, bet apziņā uzreiz rodas gatava atbilde: „Nepaspēšu!” Protams, kad smadzenes novērtē situāciju, tās neskaitļo pēc formulām, cipari vispār ir mākslīga matemātiķu izdomāta sistēma, lai padarītu skaitļošanas procesu „redzamu” apziņai. Bet iekšā datorā arī nav nekādu ciparu! Tur ir tikai elektriskie signāli. Tāpat, kā smadzenēs. Tikai mašīnā elektrība ir no kontakta, bet smadzenēs to izstrādā šūnas. Smadzenes un dators strādā pēc viena principa. Tikai dators strādā ātrāk. Kāpēc tad

datori vēl aizvien nespēj domāt? Tieši šo jautājumu šobrīd apspriež zinātne. Katrā ziņā tas pats *Hugo de Gari* domā, ka šajā gadsimtā cilvēce sadalīsies divās karojošās frontēs. Vieni, rasisti, visiem spēkiem aizstāvēs savas sugas intereses. Otri metīsies cīnīties par prātu, lai kādā formā tas arī nebūtu īstenojies. Pēc visa spriežot, tagad tiešām nāk domājošo mašīnu ēra. Tās saucas par neirodatoriem. Un pirmās līdzīgās mašīnas jau eksistē.

Šobrīd neirodatoru ir ļoti maz. Tas ir dārgs priekšs. Pat tie, kurus tagad ražo, tiek sākotnēji izgatavoti kādam noteiktam uzdevumam, piemēram, bezpilota izlūklidmašīnas vadīšana, vai arī ceļu satiksmes novērošanas sistēma. Pie katra luksofora tiek uzstādīta videokamera. Dators saprot visus mašīnu numurus, kuras brauc pa krustojumu, un sūta šos datus uz centrālo posteni.



a



b

25. att. Slaukšanas (a) un lauka (b) roboti darbā
(avots: De-Laval un Halmstad University, Zviedrija, 2005)

Daudzējādā ziņā situācija ar neirodatoriem šobrīd ir līdzīga situācijai 70-to gadu sākumā, kad elektroni skaitļojošo mašīnu varēja iegādāties tikai lieli zinātniskie centri un lielās korporācijas. Par masu psiholoģiju. Sabiedrības uzvedība ir analogiska bara uzvedībai. Masu ietekme ļoti vienkāršo indivīda domāšanu. It īpaši bara instinkti paaugstina līdera lomu, kas šajā gadījumā ir prognozes līkne. Bet, lai vadītu baru, līkne vispār nav vajadzīga. Cilvēku kopienas, aitu, zivju vai putnu bara vadības algoritms tiek nodrošināts tikai ar trim vienkāršiem nosacījumiem [22]: 1. Bada sajūta. 2. Mērķa norādījums barības virzienā. 3. Lūgums nekāpt vienam otram uz papēžiem, jo līdz barībai tā var arī netikt. Kas IST jomā var būt vienkāršāks? Saprotams, ka tirgojot ar

tik vienkāršām IST, tādas firmas, kā , teiksim, „Microsoft” momentāni bankrotēs. Jo monopola apstākļos dominē citas – dārgas, sarežģītas un nepilnīgas IST, piemēram, tādas kā „Windows” versiju pusfabrikāti, kuri nekad nekonfliktē varbūt vienīgi ar savu saimnieku - *Bilu Geitu*.

Ir ļoti vienkārša spēle, kurā mašīna vienmēr uzvar cilvēku. Jūs iedomājaties skaitli no 1 līdz 10 un mašīna mēģina to uzminēt. Pēc kāda neliela laika sprīža, tā sāk dot pareizas atbildes. Izrādās, ka cilvēks nevar izdomāt patiešām nejaušu skaitli! Viņš to dara pēc noteikta algoritma, un šo algoritmu mašīna ir spējīga atklāt. Precīzāk, tā atklāj nevis pašu algoritmu, bet rezultātu. Tas notiek it kā intuitīvi. Masu psiholoģija ir sākotnēji primitīvāka. Izglītots domātājs ir tas, kas var distancēties no cilvēku bara, tā vadīšanā viņš spēlē tieši uz masu emocionālajām tieksmēm. Bet tas taču ir ļoti grūti!



26. att. **Nanofabrika**

Kad viss tiek pārdots, tev arī ļoti gribas pārdot vismaz kaut ko. No matu sukas līdz dzimtenei ... Bet nedrīkst šaubīties savas stratēģijas pareizumā. Ir vajadzīga dzelzs griba. Datoriem šajā ziņā griba ir absolūti dzelžaina. Cilvēks ir pielāgots citai dzīvei – trīsdimensiju telpā. Taču laika rindu analīzei tā ir par maz. Tādā intuitīvajā līmenī, kā neirotikls, cilvēks vēl nav konkurētspējīgs. Bet cilvēka neironiem domāšanas ātrums sasniedz 100 pārslēgšanas sekundē, bet elektroniskajiem neironiem – 100 milj. sekundē!

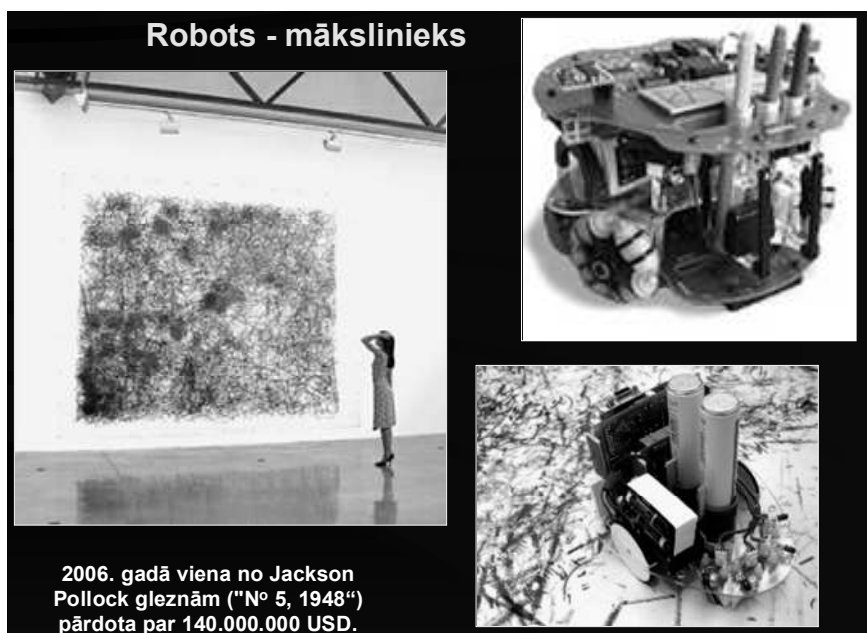
Pēc tam, kad tirgū parādīsies vairāk elektronisko ekspertu-padomnieku, strauji samazināsies debilo skaits, izjūks biržas, pazudīs baņķieru, brokeru, politiķu, juristu, advokātu, ārstu, vidējās un augstākās izglītības iestāžu pasniedzēju parazitējošas profesijas, kuri barojas uz cilvēku nabadzības, bēdu, uzticības, veselības un neizglītotības rēķina. Elektroniskais starpnieks nebūs korumpēts, nebūs tik alkatīgs un skops, kā dažs baņķieris, uzņēmējs vai ierēdnis; neņems kukuļus, izpildot darba pienākumus; nemelos skatoties cilvēkiem acīs, kā dažs politiķis vai tirgonis; sāks godīgi un, galvenais, prasmīgi ārstēt cilvēkus, nevis pagarinās tiem agoniju; ieteiks, kā labāk rīkoties ierobežotu resursu apstākļos; paredzēs cilvēku uzvedību ārkārtējās situācijās utt. Sociālā psihoze stabilizēsies un būs atkarīga tikai no objektīviem faktoriem. Bet Ja elektroniskiem domātājiem piemītis arī dzelzs griba, tad cilvēks, iespējams, kā bioloģiskā suga vispār pazudīs.

19. IST attīstība

Mūsdienās datori daudz ko dara labāk un ātrāk, nekā cilvēks. Bet ir viena problēma. Datori gandrīz neprot saprast tēlus. Pagaidām cilvēks vēl ir labākais universālais manipulators. Tas nav nejauši, jo cilvēkam, kas ir ideāli pielāgojies orientācijai trīsdimensiju telpā, vajadzēja prast izdzīvot tūkstošiem gadu. Tikmēr uz mūsu planētas ir tik daudz cilvēku, kuru pamata darbs ir tēlu saprašana un manipulēšana ar tiem. Paņem metāla stieni, nostiprina, pagaida, kamēr darbmašīna to automātiski apvirpos, izņem. Galveno darbu padarīja programmas vadība, bet cilvēks ir tas, kurš padod. Robotu tehnikas attīstība novedīs pie tā, ka simtiem miljonu cilvēku vairs nebūs darbmašīnas apkalpotāji.

Neirodatoros jau ir ieguldīta liela nauda. Liela daudzuma lētu specializētu procesoru parādīšanās tirgū novedīs pie tā, ka apmācīti roboti izplatīsies visur. Cilvēki vairs neturēs mājās kaķus un suņus un sāks iegādāties gudras elektroniskas rotaļlietas, kas neēdīs visu, kas pagadās un netaisīs čupas, kur pagadās. Taču paradīze, kas atnāks līdz ar robotu tehnikas attīstību, kļūs par lielu pārbaudījumu cilvēcei. Ko mēs darīsim, kad mums nebūs ko darīt? Uzmanību, atbilde: ilgāk un labāk dzīvosim. Vai tā ir slikta

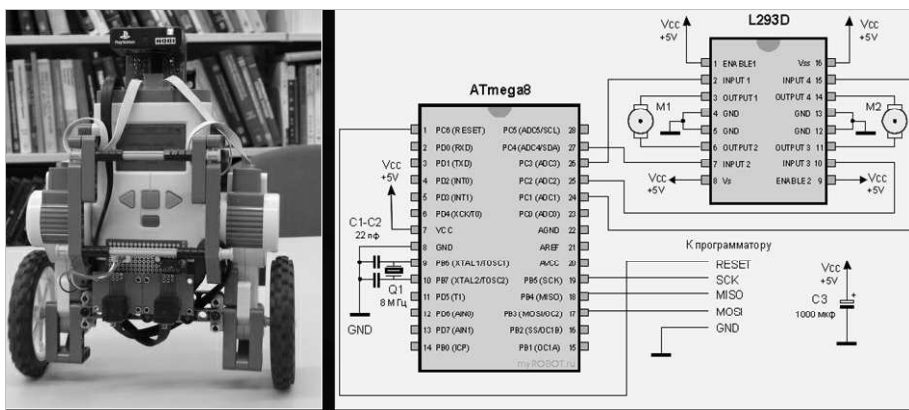
nodarbe? Nogursim no nekā nedarišanas? Diez vai. Cilvēks tik verdziski daudz strādājis visas cilvēces vēstures garumā, ka nogurums no atpūtas viņam nedraud. Iespējams, ka darbs kļūs par mākslu, par cilvēku garīgu nepieciešamību, nevis kā iespēja iegūt kompensāciju par izlietiem sviedriem un dabūtām tulznām. Atbrīvojoties no darba mēs vairāk nodarbosimies ar bērnu audzināšanu, daudzpusīgāk attīstīsimies garīgi. Nekāda prieka no smaga fiziska darba cilvēkam nekad nav bijis, nav un nekad nebūs. Bet cilvēces globālo tieksmi pēc atpūtas var gan ļoti saskatīt. Tieši tāpēc civilizācijas attīstība tiecas pie automatizācijas. Slinkums ir progresu dzinējs. Ne velti mēs visu laiku cenšamies novelt savus pienākumus mašīnām. Un pienāks brīdis, kad mums izdosies visu darbu uzdot mašīnām. Kā dzīvot tālāk? Elementāri – baudot atpūtu! Un nekādu arodbiedrību. Visa mūsu civilizācija, morāle, cilvēku attiecības pamatojas uz to, ka mums ir smagi jāstrādā. Strādāt ir ļoti, nestrādāt – slikti. Tas ir kultūras mugurkauls. Kas paliek invalīdam, kad mugurkauls ir laužts. Tikai pieejamo baudu gūšana, narkomānija – elektroniskā vai ķīmiskā, kura neizsauc pieradumu un postu organismam?



27. Robots - mākslinieks (avots: „Action Jackson”)

Jau tagad ir novērojama sekojoša tendence: narkotikas pamatā iegūst trešās pasaules valstīs, bet lieto – attīstītajās. Iespējams, ka radot mākslīgo prātu, cilvēks pašlikvidēsies narkotiskajā vai virtuālajā murgā ... Kā tas viss notiks? Pakāpeniski ... Vispār, process jau sen ir sācies. Prāts taču ir nevis indivīdu īpašība, bet visas cilvēku sabiedrības īpašība. Tik tiešām, prāta izpausme ir lomu sadales rezultāts kādā sarežģītā organismā.

Kas ir konkrētais cilvēks? Tēvs. Inženieris. Dēls. Znots. Policists. Vadītājs. Padotais. Pircējas. Iemītnieks. Patērētājs. Katra lomu funkcija ir nekas cits, kā uzvedības programma un reakcija uz apkārtējo pasauli. Kāda ir pasaule, tāda ir reakcija. Bet arī datoriem ir savas lomas un programmas. Tie tāpat kā mēs ir savīti ar sazināšanās tīkliem. Daudzās dzīves jomās cilvēks jau sen nekontrolē situāciju. Cilvēks ar datoru palīdzību pakāpeniski tiek izslēgts no lēmumu pieņemšanas procesa. Daudzi attīstītajās valstīs par to pat nenojauš. Bet zinošie pie tā ir vienkārši pieraduši, uzskatot datorus par mūsu kalpiem.



28. att. **Intelektuālie aģenti vienkāršu lēmumu pieņemšanai**
(avots: AVR mikrokontroleris, draiveris L293D)

Pagaidām tā tas arī ir. Bet prāts pakāpeniski pāriet no organiskās formas neorganiskajā. No DI jomas – uz MI. Pagaidām MI ir ļoti primitīvs, jo algoritmus izdomājuši cilvēki. Tomēr jau tagad dažas programmas ir tik sarežģītas, ka nav neviena, kas zinātu, kā tās funkcionē. „Microsoft” kompānijā nav neviena tāda

darbinieka, kurš zinātu, kā darbojas katrs atsevišķs „Windows” elements. Bet algoritmi, pēc kuriem funkcionēs neirodatoru „metāla” daļa, ir principiāli nekontrolējami. Apstādināt procesu nevar, jo tas ne no kā nav atkarīgs. Mēs visi strādājam, tūkstošiem zinātnes darbinieku dara kādas konkrētas lietas. Kopumā mēs visi, paaugstinot mašīnu intelektu, padarām kādu globālo darbu. Neaizdomājoties, ar ko tas viss var beigties.

Diemžēl mēs nedomājam ar evolūcijas mērogiem. Bet, ja abstrahēties no pasaules niecības un pāriet uz citu mērogu, mēs skaidri ieraudzīsim, pie kā galu galā novedīs mūsu pētījumi. Kādreiz cilvēks atklās, ka Tīklā dzīvo kāds gudrāks (vai gudrāki) par viņu. Aizvainots pašlepnums var izsaukt viņā tādu vēlmi, kā „es tevi radīju – es tevi arī iznīcināšu”. Bet nogalināt to būs jau neiespējami, tas būs pašnāvniecisks lēmums situācijā, kad visas komforta un dzīves nodrošināšanas sistēmas tiek vadītas ar datoru palīdzību. Protams, atsevišķi ekstrēmisti un teroristi noteikti parādīsies uz prāta sajukuma pamata, bet agresīvie aģenti pastāvējuši visu laiku, kā vīrusi, un pagaidām neatstāj tik lielu ietekmi uz evolūciju. Kad cilvēce kļūs brīva no saviem kapračiem – politiķiem, baņķieriem, brokeriem, juristiem un ārstiem – tad to metīsies aizsargāt nesavtīgi likuma un kārtības sargi – nanoroboti, terminatori un robokopi.

20. „Cilvēks” - tas skan sarežģīti!

Cilvēkā [33] ir vairāk kā 10^{14} šūnu. Viņa ķermenī ir 60% ūdens. Tā sadalījums ir nevienmērīgs: tauku audos ūdens daudzums sastāda 20%, kaulos – 25%, aknās – 70%, muskuļos – 75%, asinīs – 80%, smadzenēs – 85% ūdens no smadzeņu svara. Pārējie 40% cilvēka ķermeņa svara ir sadalīti sekojoši: olbaltumvielas – 19%, tauki – 15%, minerālvielas – 5%, ogļhidrāti – 1%. Pieauguša cilvēka organismā skābekļa, oglekļa, ūdeņraža un slāpekļa daudzums sastāda aptuveni 70%. Kalcija un fosfora saturs – apmēram 2 kg. Kālija, sēra, nātrija un hlora saturs ir dažī desmiti gramu. Dzelzs saturs ir apmēram 6 grami, bet tam ir liela loma organismā, jo tas ietilpst hemoglobīna sastāvā. Kopējais cilvēka asinsvadu garums ir apmēram 100 000 km.

Miera stāvoklī asins ir sadalīta sekojoši: 25% - muskuļos, 25% - nierēs, 15% - zarnu sienīņu asinsvados, 10% - aknās, 8% - smadzenēs, 4% - sirds vainaga asinsvados, 13% - plaušu un citu orgānu asinsvados.

Katrs eritrocīts satur apmēram 270 miljonu hemoglobīna molekulu. Asins šūnas visu laiku atmirst un tiek aizstātas ar jaunām. Eritrocīta dzīves ilgums ir 90 – 125 dienas, leukocīta – no dažām stundām līdz nedēļām. Pieaugušam cilvēkam katru stundu atmirst miljards eritrocītu, 5 miljardi leukocītu, 2 miljardi trombocītu. Tos aizstāj jaunas šūnas, kuras izstrādā kaulu smadzenes un liesa. Diennaktī nomainās apmēram 25 g asins. Pieauguša cilvēka kaulu smadzenes (irdena masa, ar kuru ir piepildīti dažu kaulu dobumi) sver vidēji 2600 gramu. 70 gadu dzīves periodā tas dod 650 kg eritrocītu un vienu tonnu leukocītu. Miera stāvoklī, guļot, cilvēks diennaktī patērē 400 – 500 litrus skābekļa, veicot 12 – 20 ieelpas un izelpas minūtē. Salīdzināšanai: zirga elpošanas intensitāte – 12 elpošanas kustības minūtē, žurkai – 60, kanārijputniņam – 108. Pavasarī elpošanas intensitāte vidēji ir par 1/3 augstāka nekā rudenī. Pieaugušam cilvēkam sirds diennaktī pārsūknē ap 10 000 litru asiņu. Ar vienu pulsa sitienu aortā tiek izsviests ap 130 ml asiņu. Normālais pulss miera stāvoklī ir 60 – 80 sitienu minūtē, pie tam sievietēm sirds pukst par 6 – 8 sitieniem minūtē biežāk nekā vīriešiem. Smagā fiziskā slodzē pulss var paātrināties līdz 200 un vairāk sitieniem minūtē. Salīdzināšanai: zilonim pulsa biežums ir 20 sitieni minūtē, vērsim – 25, vardei – 30, trusim – 200, bet pelei – 500 sitieni minūte.

Cilvēka nervu sistēma satur apmēram 10 miljardus neironu un apmēram 7 reizes vairāk apkalpojošo šūnu – atbalsta un barojošo. Tikai 1% nervu šūnu ir aizņemts ar „patstāvīgu darbu” – pieņem sajūtas no apkārtējās vides un komandē muskuļus. 99% ir nervu starpšūnas, kas kalpo kā pastiprināšanas un sūtīšanas stacijas. Pašas lielākās cilvēka nervu šūnas ir 1000 reizes lielākas par vismazākajām. Pašiem smalkākajiem nervu audiem caurmērs ir tikai 0.5 mikrometri, paši resnākie – 20 mikrometri. Vairāk par pusi visu neironu ir koncentrēts lielajās smadzeņu puslodēs. Kopējais cilvēka smadzeņu garozas laukums variē no 1468 līdz 1670 kvadrātcentimetriem. Galvas smadzeņu nervos ieiet 2 600 000 nervu audu, bet iziet 140 000. Aptuveni puse no izejošo audu pārnēs signālus acs ābola muskuļiem, vadot smalkas, ātras un sarežģītas acs kustības. Pārējie nervi vada mīmiku, košļāšanu, rīšanu

un iekšējo orgānu darbību. No ieejošajiem nervu audiem 2 milj. ir redzes nervu audi. Vienā minūtē caur smadzenēm plūst 740 – 750 ml asins. Sākot no 30. dzīves gada cilvēkam katru dienu atmirst 30 – 50 tūkst. nervu šūnu. Samazinās smadzeņu pamatizmēri. Ar vecumu smadzenes ne tikai zaudē svaru, bet arī nomaina savu formu – sablīvējas. Vīriešiem smadzeņu svars sasniedz maksimumu 20 – 29 gados, sievietēm – 15 – 19 gados. Cilvēka smadzenes sver 1/46 no kopējās ķermeņa masas, ziloņa smadzeņu masa – tikai 1/560 no viņa ķermeņa masas. Cilvēka organismā darbojas ne mazāk par 700 fermentiem. Paši bargākie vīrieši katru dienu izraud 1 - 3 ml asaru. Asaras visu laiku izdalās no asaru dziedzeriem un mitrina acs radzeni, pasargājot to no gaisa un putekļu ietekmes. Svaigais pirksta nospiedums sver apmēram 1 miljono daļu no grama. Tas sastāv no ūdens, taukiem, olbaltumvielām un sāļiem, kas izdalās caur ādu.

Cilvēkam ir vairāk par 200 kauliem. Precīzāk pateikt nevar, jo dažādiem cilvēkiem kaulu daudzums ir dažāds (20%-iem ir novirzes skriemeļu daudzumā), kā arī ar vecumu daži kauli saaug. Visgarākais kauls ir gurnu kauls, tā garums parasti ir 27,5% no cilvēka auguma. Visīsākais – kāpslītis, viens no kauliņiem, kuri pārnes bungādiņas svārstības iekšējās auss jūtīgajām šūnām. Tas darbojas kā svira, palielinot skaņas viļņu spiedienu. Tas ir tikai 3-4 mm garš. Precīzi norādīt muskuļu daudzumu arī nevar. Speciālisti uzskata, ka cilvēkam ir 400 – 680 muskuļu. Salīdzināšanai: circeņiem ir ap 900 muskuļu, dažiem kāpuriem – līdz 4 000. Vīrietim kopējais muskuļu svars ir ap 40% no ķermeņa masas, bet sievietei – ap 30%. Vismazākais muskulis ir kāpslīša muskulis. Pārāk stiprās skaņās tas griež kāpslīti tā, ka kauliņa-svira plecu garuma attiecība mainās, un skaņas pastiprināšana krīt.

Vidējais normālais redzes asums ir 0,0003 leņķiskās minūtes, t.i. acs spēj izšķirt labi apgaismotu priekšmetu 1/10 mm caurmērā 25 cm attālumā. Bet ja pats priekšmets staro, tas var būt arī daudz mazāks. Caurums ar diametru 3 – 4 mm tūkstošdaļas, kas ir caurdurts skārda lapā, aiz kura ir aizdedzināta spuldzīte, ir labi izšķirams ar normālu aci. Acs ir spējīga izšķirt 130 – 250 tīru krāsas toņu un 5 – 10 miljonu jaukto toņu. Uzliesmojumu biežums, kurā mirgojoša gaisma šķiet vienmērīgi degoša, acu nervu kociņiem sastāda 15 sekundē, kolbiņām – 71 – 90. Pilnā acs adaptācija pie tumsas aizņem 60 – 80 min. Pirksts ir spējīgs sajust svārstības ar amplitūdu 0,00002 mm

daļas. Cilvēka ādas virsma vidēji aizņem ap 2 m^2 . Vienā minūtē caur ādu iziet 460 ml asins. Ādā ir izkļiedēti 250 000 aukstuma receptoru, 30 000 siltuma receptoru, 1 000 000 sāpju *receptoru*, 500 000 taustes receptoru un 3 000 000 sviedru dziedzeru. Cilvēkiem ar blondiem matiem vidēji uz galvas ir 140 000 matu, brunetiem – 102 000, šateniem – 109 000, bet rudmatainiem – 88 000. Kopējais matu skaits uz ķermeņa, neieskaitot galvu, ir apmēram 20 000. Matī aug ar ātrumu 0,35 - 0,40 mm diennaktī. Diennaktī mūsu mati kļūst garāki apmēram par 30 metriem (kopskaitā). Nagi uz rokām aug ar ātrumu 0,086 mm diennaktī, uz kājām – 0,05 mm. Gada laikā uz roku pirkstiem uzaug ap 200 g nagu.

Iekšējā ausī ir ap 25 000 šūnu, kuras reaģē uz skaņām. Frekvenču diapazons, kuru mēs uztveram ar dzirdi, atrodas starp 16 un 20 000 Hz. Kad cilvēkam paliek ap 35 gadi, augšējā dzirdes robeža krīt līdz 15 000 Hz. Auss ir vairāk jutīga 2 000 – 2 300 Hz diapazonā. Labākā muzikālā dzirde (spēja izšķirt skaņas augstumu) ir 80 – 600 Hz apgabalā. Piemēram, šeit mūsu auss spēj izšķirt divas skaņas ar frekvenci 100 Hz un 100,1 Hz. Vispār cilvēks spēj izšķirt 3 000 – 4 000 dažāda augstuma skaņu. Mēs izšķiram skaņu pēc 35 – 175 milisekundēm pēc tam, kad tā tiek līdz ausij. Vēl 180 – 500 milisekundes ir nepieciešamas ausij, lai sasniegtu labāku jutīgumu. Uz mēles atrodas ap 9 000 garšas receptoru. Labākā temperatūra to darbībai ir 24°C . Deguna ožas zonas laukums ir 5 cm^2 . Šeit atrodas aptuveni miljons ožas nervu receptoru. Lai nervu ožas šķiedrā parādītos impulss, uz tā galiņu jānokļūst ap 8 molekulām smaržīgas vielas. Lai parādītos smaržas sajūta, ne mazāk kā 40 nervu šķiedrām ir jāuzbudinās. Barības sakošļāšanā žokļa muskuļi attīsta spēku uz dzerokļiem līdz 72 kg, uz priekšzobiem – līdz 20 kg. Maizes sakošļāšanai ir nepieciešama 25 kg piepūle. Ceptai teļai – 15 kg. Uz vienu kvadrātmilimetru kuņģa gļotādas ir ap 100 dziedzeru, kuri izdala gremošanas sulu. Cilvēka tievajai zarnai, kurā notiek sagremotas barības iesūkšanās asinīs, uz tās iekšējās virsmas ir ap 5 000 000 bārktiņu – smalku matveidīgu izaugumu, caur kuriem arī notiek barības vielu uzsūkšana. Slāpju sajūta rodas ūdens zuduma dēļ, kas ir vienāda ar 1% no ķermeņa masas. Zaudējot vairāk kā 5% ūdens, cilvēks var zaudēt samaņu, bet, zaudējot vairāk par 10% ūdens, iestājas nāve. Ūdens malks – vai tas ir daudz? Daudzi mērījumi ir pierādījuši, ka vīrietis norij ar vienu malku vidēji 21 ml šķidruma, bet sieviete – 14 ml. Vismazākā sieviete pasaulē

ir Lučija Zarate (1864–1890). Piedzimšanas brīdī viņas augums bija 17 cm. Lučija izauga līdz 43 cm un svēra 2,2 kg.

21. Mašīnas intelekts

Ja domāt ar savu galvu, tad uz bibliotēku var arī neiet (M. Geny)

Viena no „intelektuālās mašīnas” vārdiskajām definīcijām skan: „Jebkura elektroniskās skaitļošanas mašīnas tipa sistēma vai ierīce, kura izpilda vai palīdz cilvēkiem uzdevumu risināšanā”. Bet vai tad tik daudz cilvēku problēmu būs spējīgas risināt šīs mašīnas realitātē? Skaitļošanas process, kas agrāk bija kā domāšanas akts ārpus mašīnu izmantošanas jomas, tagad uz intelektuālo uzdevumu un apmācīšanas sfēru. Šodien jau neviens neuzskata kabatas kalkulatora izmantošanu par kaut kādu īpašu MI īpašību. Skaitļošana tagad liekas kā vienkārša elementāra mehāniska procedūra. Savā laikā parasto datoru radīšanas ideja pati par sevi izskatījās netikumīga. Taču 19. gs. vidū Čarlzs Babbags izveidoja mehāniskos kalkulatorus un daļu no programmējamā mehāniskā datora. Bet viņš sadūrās ar dažām problēmām. Piemēram, kāds doktors Jungs pierādīja, ka ir lētāk ieguldīt nevis kalkulatoru, bet darbinieku, kuri nodarbojas ar skaitļošanu ar rokām, apmaksāšanā.

Lielbritānijas Karaliskais Astronoms sers Džordžs Airi savā dienasgrāmatā raksta, ka „15. septembrī G. Goulburns... lūdza izteikt savu viedokli attiecībā pret Babbaga skaitļošanas mašīnas derīgumu ... un es, detalizēti izpētot jautājumu, atbildēju, ka šis darbs nav pat izēstas olas vērts ...”. Savā laikā Babbaga mašīnai bija jēga tik vien tāpēc, ka mehāniķiem vajadzēja pavirzīt uz priekšu tās mehānisko daļu radīšanas mākslu ar augstāku precizitāti, nekā parasti pieprasīja. Bet faktiski šī mašīna tikai nedaudz pārsniedza kvalificētā speciālista darba ātrumu. Tomēr vecais skaitļošanas veids bija drošāks un pietiekami vienkāršs priekš tā, lai tiktu uzlabots.

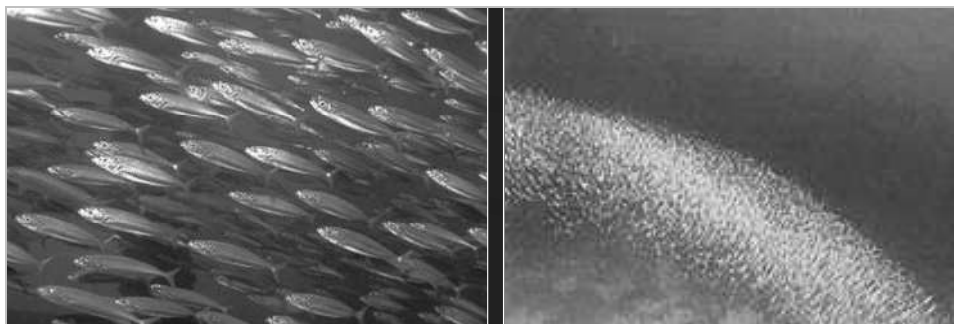


29. att. **Automatizācija un robotizācija precīzās lauksaimniecības tehnoloģijās**
(avots: DAAD, IfL-LfL, Vācija)

Datoru un MI attīstības vēsture atgādina aviācijas attīstības vēsturi. Līdz šim cilvēki noraidīja abas idejas kā neiespējamas – parastā situācija, kad nav acīmredzami parādīti idejas realizācijas veidi. Un līdz brīdim, kamēr MI modelis nekļūs pietiekoši vienkāršs saprašanā un īstenošanā, ieskaitot demonstrāciju, nekādas sarunas par lidmašīnas pilnīgu ekvivalentu neaizvieto parasto raķeti, kura ir pielietojama Mēness vai Marsa sasniegšanai. Ilgā laika gaitā cilvēki turpina mainīt pašu intelekta definīciju. Neskatoties uz tādiem reklāmas sludinājumiem un izziņām presē, kā „izcils elektroniskais prāts”, tomēr tikai nedaudzi sauca pirmos datorus par intelektuāliem. Patiešām, pats nosaukums „dators” nozīmē parasto aritmētisko mašīnu. Tomēr 1956. gadā, Dartmutā, pirmās vispasaules konferences MI jomā laikā, pētnieki *Alans Nevells* un *Herberts Saimons* prezentēja savu „Loģisko Teorētiku”. Tā bija programma, kas spējīga pierādīt teorēmas simboliskā loģikā. Vēlākajās datoru programmās spēlēja šahu un palīdzēja ķīmiķiem noskaidrot molekulārās struktūras. Divas medicīniskās programmas, CASNET un MYCIN bija izveidotas pietiekoši kvalitatīvi: pirmā bija domāta zāļu radīšanai, bet otrā – diagnosticēšanai un infekcijas slimību ārstēšanai.

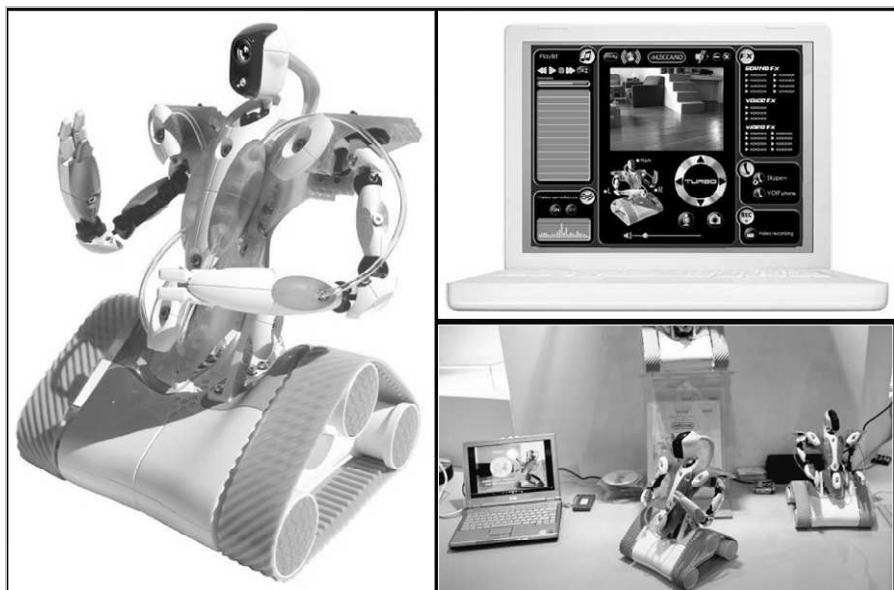
Saskaņā ar „Mākslīgā intelekta instrukciju” šīs programmas eksperimentālos variantos bija novērtētas kā „izpildītas attiecīgās jomās izmantošanas iespējamības līmenī”.

Jau tad programma PROSPECTOR maksāja miljonus dolāru. Šīs, tā saucamās „pieredzējušās sistēmas” bija veiksmīgas tikai stingri ierobežotas kompetences jomas robežās, bet tās varētu izbrīnīt 1950. gadu sākuma programmētājus. Taču šodien ne tik daudz cilvēku uzskata, ka šie projekti patiešām kļūs par reālu mākslīgo intelektu, tā kā MI visu laiku bija kā mērķis kustībā. Prese reklamēja izstāžu sasniegumus un paziņoja, ka mūsdienu datori var būt ieprogrammēti ar pietiekošu zināšanu krājumu un var risināt neparastus uzdevumus. Tāpēc cilvēkiem parādījās noteikta interese par intelektuālām spējām. Daudzie praktiskās rūpniecības robotu izmantošanas un jauno datoru iespēju publiskās apspriešanas radio un televīzijā gadi vismaz padarīja MI par salīdzinoši populāru ideju. Galvenais iemesls, lai paziņotu MI ideju par neiespējamu, vienmēr bija cilvēku pārliecība, ka mašīnas ir primitīvas un pat stulbas – viedoklis, kurš tagad sāk pakāpeniski mainīties. Pirmās mašīnas tik tiešām bija lempīgas, masīvas un izpildīja vienkāršu, rupju darbu. Bet datori veiksmīgi rīkojas ar informāciju, izpilda sarežģītas instrukcijas un var būt instalēti tā, lai izmainītu savas instrukcijas. Beidzot, tie prot eksperimentēt un mācīties. Tie nesatur norūsējušus dzelzs mehānismus un trauslas olbaltumvielu struktūras, bet tām ir daudzi elektriskās enerģijas vadītāju kanāli. Mākslīgā intelekta idejas kritiķi bieži norāda uz pastāvošo datoru intelektuālo ierobežotību. Bet tas nepierāda neko, kas attiecas uz to nākotni. Nākotnes domājoša mašīna, iespējams, varētu pati uzdot sev jautājumu, vai paši MI kritiķi ir prātīgi savos piedāvājumos un apgalvojumos?



30. att. **Haoss un antihaoss dabā. Zivju bars, kā vadības sistēmas modelis**
(avots: Moskvins, 2006)

Viņu aizrādījumi tika bāzēti uz atpalikušās pagātnes pieredzes, kad tvaika mašīnas vēl nemācēja lidot, kaut gan tās demonstrēja mehāniskos principus, kas vēlāk tika pielietoti virzuļa lidmašīnu dzinējos. Analogiski, simts gadu atpakaļ cilvēki vēl nevarēja noteikt sliekām intelekta pazīmes, kaut gan mūsu smadzenes jau tad izmantoja neironu principus. Nejaušie kritiķi tāpat izvairās no nopietnām domām attiecībā pret MI, sakot, ka mēs nevarēsim uztaisīt domājošas mašīnas saprātīgākas par mums pašiem. Viņi aizmirst par vēsturisko pieredzi. Mūsu attālie, nerunīgie senči varēja uzlabot intelekta kvalitāti caur ģenētisko attīstību, pat bez gudras spriedelēšanas šajā sakarā. Principus, kas līdzinās tiem, kurus salīdzinoši nesen atrada primitīvām sliekām. Taču šodien mēs aktīvi apspriežam un apdomājam MI problēmas, un arī pašas mākslīgās domāšanas mašīnu tehnoloģijas pašlaik attīstās daudz straujāk, nekā gēni bioloģijā. Tādēļ ar laiku cilvēki, protams, varēs taisīt domājošas mašīnas, kuras pārsniegs cilvēka spējas. Un šīs mašīnas būs ar cilvēkiem līdzīgu spēju ne tikai pētīt, bet arī paši spēs organizēt, sistematizēt un apkopot zināšanas.



31. att. **Intelektuālie aģenti vienkāršu lēmumu pieņemšanai.**
(avots: apmācošais robots „Spyke”, 2007).

Liekas, ka ir tikai viens iemesls, kas varētu liecināt par labu mākslīgā intelekta radīšanas neiespējamībai jaunajās jēdzieniskajās formās. Tā ir ideja par Visuma domas

materialitāti pēc saprātīgā materiālisma koncepcijas. Pastāv uzskats sakarā ar šo problēmu, kas paredz, ka atsevišķas domas materiālā būtība un saprātā būtība kopumā ir kāds materiāls, ko nevar ne modelēt, ne atdarināt, ne dublēt, ne kopēt, nevar pakļaut nekādai fiziskai vai mehāniskai iedarbībai vispār. Psihobiologi neredz nekādas šādas vielas esamības acīmredzamības reālas liecības un neatrod nekādas nepieciešamības prāta materiālismā tam, lai paskaidrotu saprāta būtību. Viņi uzskata, ka smadzeņu sarežģītums ir tik liels, ka saprāts vispār atrodas ārpus cilvēka tā saprašanas iespēju robežām. Tāpēc ideja izskatās pietiekami sarežģīti vilinoša projekta par mākslīgā saprāta radīšanu realizēšanai. Patiešām, ja cilvēks spētu saprast savu smadzeņu būtību, tad viņa smadzenes izrādītos mazāk sarežģītas saprašanai, kā paša saprāts. Ja mūsu zemeslodes cilvēku miljardi varētu vienlaicīgi vērot viena cilvēka smadzeņu darbību, tad katram cilvēkam vajadzēs kontrolēt desmitiem tūkstošu aktīvo šūniņu – sinapšu.

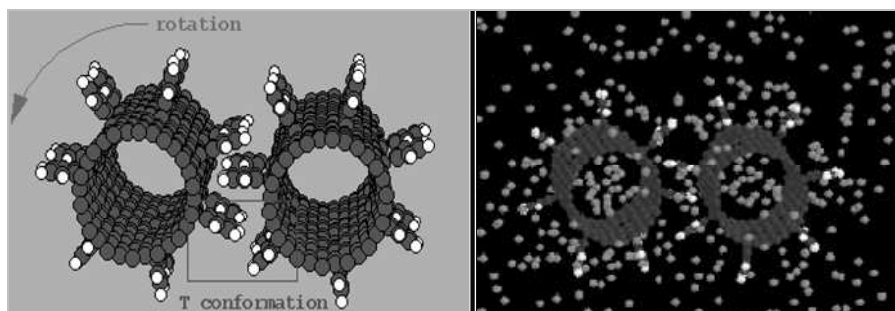
Vēl neejdzīgāk šodien izskatās viena atsevišķa cilvēka mēģinājums atšifrēt, operēt un apjēgt elektrofiziskos mirgojošos smadzeņu sinapšu signālus vienlaicīgi sešu miljardu cilvēces galvu kopumam. Tas izskatās, kā mistisks virsuzdevums, kas ir pa spēkiem ne tik cilvēka saprātam, cik tā labai iztēlei. Tātad, lai cilvēks varētu saprast un novērtēt pats savu saprātu, viņam ir nepieciešams būt pilnīgākai substancei – virsprātam. Tas ir viens no ne tik acīmredzamiem loģiskiem paradoksiem, kuru ir pietiekami daudz smagā mākslīgā saprāta radīšanas ceļā. Tādējādi mūsu smadzeņu darbības sarežģītums un globālums pats par sevi plašā izpratnē sagrauj mūsu smadzeņu spēju aptvert, saprast un apjēgt cilvēka saprāta mēroga lielumu. Domāšanas mehānisms šodien daudziem pētniekiem joprojām šķiet pietiekami sarežģīts tam, lai jau šodien realizētu mākslīgā saprāta ideju tieši mašīnas intelekta veidā.

22. Skaitļošanas intelekts un modelēšana

*The Self is a circle whose centre is everywhere and whose circumference is nowhere
(C.-G. Jung)*

Skaitļojošais intelekts var kalpot par mākslīgā intelekta paplašinājumu, ja ir nepieciešams formulēt problēmu un noskaidrot uzdevuma risināšanas mehānismus.

Modelēšana ir mēģinājums izpētīt un saprast reālu fizisku sistēmu funkcionēšanas un mijiedarbības būtību, galvenokārt vadot šo sistēmu datoru modeļus.



32. att. **Haoss un antihaoss. Nanomašīnas darbības princips**
(avots: Nanotech, ASV)

Skaitļošanas intelekts strādā simbolu līmenī, bet modelēšana notiek skaitļu līmenī. Šobrīd abi izpētes tipi viens otru papildina aizvien biežāk. Skaitļošanas intelekta progress nodrošināja daudzu izpētes metodoloģiju attīstību, piemēram, fazioloģika, neironu tīkli, evolūcijas algoritmi un to kombinācijas. Skaitļošanas intelekts var tikt skatīts kā metožu kopums kvalitātes procesu (plašā skatījumā) izpētei. Kopīga skaitliskā modelēšana un simboliskā skaitļošana būs ļoti svarīga lietojumprogrammu īpašība sarežģītu sistēmu izpētei. Intelektuālās metodes ir nepieciešamas, tādēļ, ka pastāv nenoteiktība, tā kā visas nepieciešamās detaļas nevar tikt ietvertas modelēšanas sistēmās. Sarežģītās sistēmās lietotāja prasības ekspertiem kļūs pārāk augstas, kad būs nepieciešamas zināšanas daudzās nozarēs. Lietojumprogrammās skaitļošanas intelekts spēlē šifrēšanas atslēgas lomu, kas ļauj tās vienkāršot tik tālu, ka programmas kļūst pieejamas pat lietotājam-nespeciālistam. Lai savienotu modelēšanas un skaitļošanas intelekta līmeņus lietojumprogrammās, šīs divas metodoloģijas var savstarpēji pielāgoties. Intelektuālās automatizētās sistēmas ar savu programmatūru pašas sev palīdzēs jaunu molekulāro modeļu radīšanā, piemēram, nanotehnoloģijas jomā. Galu galā, programmatūras sistēmas būs spējīgas radīt jaunus projektus bez cilvēku palīdzības. Jautājums par to, vai vairākam cilvēku būs vajadzīgas šādas intelektuālās sistēmas, īstenībā nav būtisks. Pats automātiskās projektēšanas process ar

pat nelielu zinošu cilvēku daudzuma līdzdalību atvēra ceļu principiāli jaunajām 21. gadsimta tehnoloģijām.

23. Tūringa domājošās mašīnas

*„There is not a truth existing which I fear or would wish unknown to the whole world”
(Thomas Jefferson)*

1950. gadā mašīnu intelekta sakarā britu matemātiķis *Alans Tūrings* izteicās diezgan filozofiski: „es domāju, ka 20. gs. beigās gudru vārdu kopums un cilvēku domāšanas līmenis izmainīsies tik tālu, ka katrs cilvēks spēs izteikt savu attieksmi pret domājošo mašīnu izveidošanas problēmu, nebaidoties izskatīties nekompetents un pretrunīgs. Bet viss būs atkarīgs no tā, kādu jēgu mēs ieliekam vārdā „domāšana”. Daži „domātāji” uzskata, ka tikai cilvēki spēj domāt, bet datori nevar būt cilvēki. Izdarot šādu asprātīgu paziņojumu, viņi, saliekot rokas klēpī, iesēdīsies mīkstajā krēslā un izskatīsies ļoti apmierināti ar sevi”.

Par cilvēku intelekta līmeni mēs parasti spriežam pēc to runas kvalitātes. Tāpēc savā laikā *Tūrings* ieteica sava veida spēli, ko līdz pat šai dienai sauc par „Tūringa testu”. Iedomājieties, ka jūs atrodaties istabā, kurai ir saikne caur terminālu ar cilvēku un datoru divās dažādās istabās. Katrs abonents sarunājas, izejot no cilvēka skatījuma un tīri intelektuālā skatījuma. Pēc ilga darba ar pogām „saruna” un „iespējas” ar rokasgrāmatu izmantošanu, informāciju par mākslu, datiem par laika apstākļu izmaiņām u tml., jūs nevarēsiet atšķirt mašīnas un cilvēka atbildes. Ja mašīna mācītos tālāk un varētu to turpmāk darīt labā līmenī, tad, kā saka Tūrings, tādu mašīnu tik tiešām var uzskatīt par intelektuālu. Šādam apgalvojumam droši vien varētu piekrist daudzi. Praktiskiem mērķiem mums ne vienmēr ir būtiski zināt, vai mašīnai ir „pašapziņa”, tas ir apziņa. Tie kritiķi, kas apgalvo, ka mašīnām vispār nevar būt apziņas, nenosauc iemeslus tādai neiespējamībai un nepaskaidro, ko viņi domā ar terminu „apziņa”. Apziņa, kas ir pietiekami attīstīta tam, lai rīkotos saprātīgi, nav vienkārši galvenais cilvēces piederums. Mums vēl jāsaprot citi, jāvērtē viņu dotības un sliecības, jāplāno savstarpējās attiecības. Turklāt, mums jāzina pašiem sevi, mūsu pašu

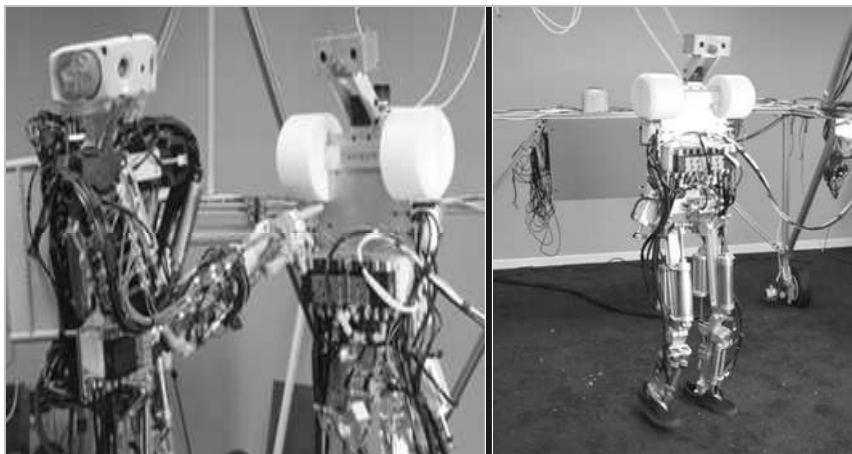
iespējās un tiekšanās, jāplāno pašu darbības un rīcības. Nav nekāda īpaša, maģiska noslēpuma „pašapziņas” mehānismos.

Mūsu apziņa un mūsu domas reaģē uz sabiedrības uzskatiem, viedokļiem un domām, tā atsaucas uz to atsevišķām rīcībām un darbībām. Šis process tiek realizēts ar īpašu tēlu, domas ekvivalentu, mijiedarbības palīdzību. Bet doma pati par sevi, kas ir saprāta īpašas būtības izpausme, kas atšķiras no smadzeņu vielas, neko mums nepaskaidro attiecībā uz apziņas būtības saprašanu. Testā uz apziņas esamību domājošā mašīna, kas atbildēja uz *Tūringa* testu, protams, teiktu, ka tai ir „pašapziņa”. Jebkurš biošovinists apšaubīs šādu apgalvojumu. Bet tikmēr, kamēr viņi nepaskaidros sabiedrībai, ko viņi domā ar terminu „apziņa”, viņi nevarēs pierādīt principiālu pašapziņas esamības neiespējamību domājošām mašīnām. Tomēr, nedomājot par sakarībām, ko mēs saucam par „apzinīgu” un „neapzinīgu” (neapsprīžot šeit psiholoģijas un parapsiholoģijas problēmas), intelektuālās mašīnas turpina un turpinās darīt savu, pat ja ne īpaši „saprātīgu”, tad jebkurā gadījumā sabiedrībai derīgu darbu, kas skar mūs visus.

Iespējams, ka apziņas būtības noteikšana kļūs par tādu argumentu, kas sadragās biošovinisma doktrīnu – kā cilvēka prātam augstākās pakāpes neticības formu, kurš ne uz ko neskatoties, neatlaidīgi pieliek maksimāli daudz pūliņus, lai izveidotu domājošo mašīnu. Līdz šim neviena domājošā mašīna nav spējusi iziet *Tūringa* testu, un, domājams, ne tik drīz spēs to izdarīt. Būtu loģiski padomāt par jautājumu, vai ir kāds pārliecinošs iemesls, lai vispār censtos veiksmīgi iziet šo testu: jo var taču gūt daudz vairāk labuma no mākslīgā intelekta pielietošanas rezultātiem, vadoties pēc citiem mērķiem. Atšķir divus mākslīgā intelekta tipus, kaut gan sistēmas ietvaros var mijiedarboties abi veidi. Pirmais tips – ir tehniskais mākslīgais intelekts, kas ir domāts lietošanai fiziskā pasaulē.

Šis MI tips visbiežāk tiek lietots automatizācijas projektos, kā arī to lieto zinātniskiem mērķiem, otrais tips – ir sociālais mākslīgais intelekts, kas attiecas uz cilvēka saprāta pētījumu lauku. Daudzu zinātnieku pūliņi šeit ir virzīti uz domājošo mašīnu izveidošanu, kas spēs veiksmīgi iziet *Tūringa* testu. Zinātnieki, kas strādā MI sociālo sistēmu jomā, pētījumu procesā vairāk un vairāk papildinās savas zināšanas par cilvēka saprātu. Pie tam viņu izstrādātām sistēmām neapšaubāmi būs aizvien lielāka

praktiska vērtība, tā kā aizvien vairāk cilvēku spēš sajūst reālu derīgumu no domājošo mašīnu intelektuālās palīdzības, rekomendācijām un ieteikumiem.



33. att. **Pašapmācošo robotu konflikts** (*Avots: 2007 Anybots, Inc., USA*)

Automatizētām IST sistēmām, kas balstās uz tehnisko intelektu, būs lielāka ietekme uz tehnoloģijas attīstības paātrinājumu, salīdzinot ar sociālo intelektu, īpaši molekulārās nanotehnoloģijas līmenī. Pilnīgāka – automātiska, tehniska intelektuālā sistēma, var izrādīties vienkāršāka realizācijai un tālākai attīstībai, nekā *Tūringa* sociālā domājošā mašīna, kurai jābūt ne tikai zinošai un intelektuālai, bet arī uzbūvētai ar iespēju atdarināt cilvēka zināšanas un cilvēka intelektu – daudz sarežģītāks un grūtāk sasniedzams uzdevums. *Tūringa* sev uzdeva šādu jautājumu: kāpēc mašīna nevar izdarīt kaut ko tādu, kas pēc visām pazīmēm būtu bīstami kā domāšana, bet ļoti atšķirtos no tā, kā to dara cilvēks?

Daži zinātnieki un politiskie darbinieki, iespējams, arī turpmāk atteiksies pieņemt mākslīgo intelektu, kā reāli domājošo mašīnu. Bet tikmēr, kamēr mašīnas intelekts nesasnies „pļāpīga dzelža” intelektualitāti, kas spēš izturēt *Tūringa* testu, daļa skeptiski noskaņoto domas inženieru turpinās pieņemt intelektu citās, vienīgi viņu prātam pieejamās formās.

24. Bionisko intelektuālo sistēmu telplaika korekcija

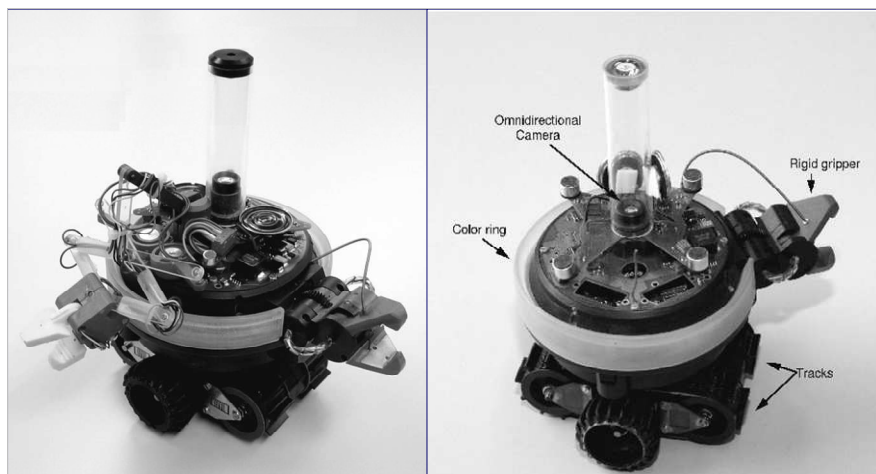
Laiks – tas ir Visuma liktenis (M.Geny)

Šodien mēs atrodamies intelektuālu automātisku sistēmu radīšanas ceļā. Inženieriem ir efektīvas sistēmas un zināšanu iegūšanas tehnoloģijas, kuras ir pieprasītas tirgū, kuras palīdz cilvēkiem risināt praktiskas problēmas. Programmētāji radījuši automatizētas sistēmas, kurās tiek pielietotas uzkrātās zināšanas par tēlu, formu, kustības, piepūles, sprieguma, elektronisko procesu, plūsmas, spiediena, temperatūras un daudzu citu parametru automātisku atšķiršanu un identifikāciju.

Projektētāji izmanto šīs sistēmas, lai palielinātu modeļa saprātīgumu, tādējādi paātrinot vēl nerealizēto projektu attīstību. Projektētāji un datori kopēji izveido intelektuālas, pusemākslīgas sistēmas. Inženieriem ir iespēja izmantot datorsistēmu plašu sortimentu, lai maksimāli palīdzētu sev darbā. Sākumā viņi izmanto datora ekrānu vienkārši zīmējumu labošanai. Tālākā darbā viņi izmanto datorsistēmas, kuras spēj matemātiski parakstīt detaļas trīsdimensiju telpā un izskaitļot nepieciešamo atbildi vadīšanas sistēmā adaptīvu komandu veidā. Dažas sistēmas ir aprīkotas ar datiem par datoru vadītu ražošanas aprīkojumu, ļaujot inženieriem veikt testējošos modeļus un instrukciju izmēģinājumus.

Izmēģinājuma rezultāts uz reāla procesa modeļa vēlāk tiek izmantots ar mašīnu vadāmu datoru atsevišķu detaļu sintēzei un projektēšanai. Tālākai automātisku sistēmu uzlabošanai datori būs nepieciešami ne tikai, lai veiktu pierakstus un skaitļošanu, bet arī, lai šīs sistēmas projektētu un ražotu. Programmētāji attīstīja darba instrumentus datora izmantošanai ienesīgā biznesā. Piemēram, viņi izveidoja programmatūru mikročipu projektēšanai. Integrālās shēmas čipi tagad satur daudzus tūkstošus tranzistoru un vadu.

Agrāk projektētājiem bija jāstrādā daudzus mēnešus, lai uzprojektētu mikroshēmu un ievietotu čipa virsmai šķērsām tā daudzās detaļas. Šodien viņi var uzticēt šo uzdevumu tā saucamajam „silīcija kompilatoram”. Šīs programmatūras sistēmas var ražot izpildei gatavu un detalizētu projektu ar čipa funkciju specifikācijas pilnu uzskaiti, turklāt ar cilvēka minimālu piedalīšanos vai pilnīgi automātiski. Skaidrs, ka visas šīs laboratorijā izveidotās un ieprogrammētās sistēmas pilnīgi balstās uz cilvēka zināšanām.

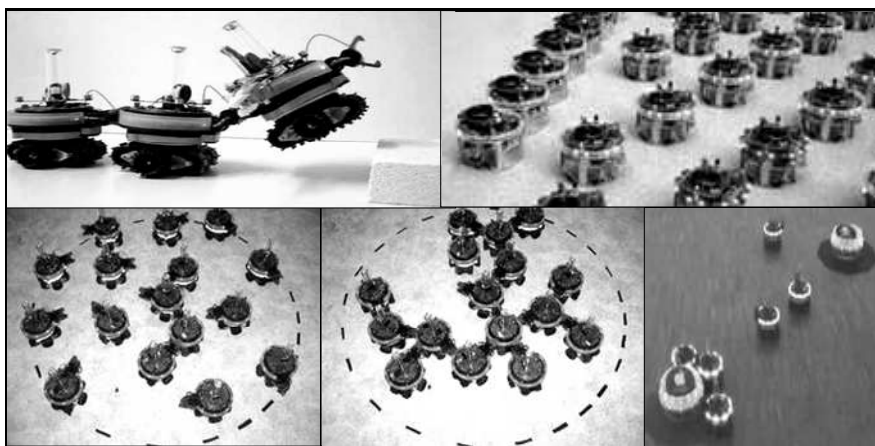


34. att. **Automātiskais modulis (swarm – bots) autosalikamo robotu sabiedrības izveidošanai** (avots: EU-IST projekts, Šveice, 2003)

Laika mēroga pareiza izvēle, tehniskais un terminoloģiskais laika kategorijas izpratne var būt noteicošs intelektuālai sistēmai, kura darbojas reāla laika mērogā (*in real time mode*). No mākslīga intelekta principiem un formālas loģikas teorijas izriet, ka laiku, principā, var raksturot kā vienmērīgo - nevienmērīgo, reālo - abstrakto, iekšējo - ārējo, tas var formāli virzīties jebkura virziena „uz priekšu” un „atpakaļ”, var eksistēt saprātā vai „virtuālajā realitātē”, var vispār neeksistēt „tagadnē”, „apstāties”, vienlaicīgi interpretēt pagātni, nākotni un tagadni. Ja pieļaut, ka „laiks” ir vektoriālais lielums un tiecas „skriet uz priekšu”, tad vektora virziens formāli norāda uz „metriskas telpas” konkrētajā momentā atrašanas, „eksistences” vietu.

Tas nozīme, ka cilvēks, tā „saprāts” individuāli, subjektīvi nosaka „telpaika” sistēmas eksistēšanas vietu un laika virzības procesa secību „pagātne-tagadne-nākotne”. Bet tāda gadījuma „tagadni” var interpretēt ar momentu, kad sava nepatrauktā virzība „uz priekšu” laiks ieņem pašreizējo stāvokli starp “nākotni” un “pagātni”. Pie tam- tikai tad, kad „mirklis ir apstājies”, jo „tagadne” ir tāds laika moments, kad „nākotne” nepārtraukti un momentāni pariet „pagātnē”. Bet tā kā laiku principa „apturēt” nevar (vektors), tad tās laika mirklis, ko mēs saucam par „tagadni” esošo teorētisko priekšstatu ietvaros reāli, fizikāli, neeksistē. Tātad tagadne neeksistējam arī mēs? Ja pieņem, ka „nākotne” momentāni un nepārtraukti pariet

„pagātnē”, tad pēc formālas loģikas var secināt, ka laika neeksistēšana „tagadnē” pieļauj domu par pētoša subjekta neesamību esošajā, pieņemtajā metriskajā (bioniskajā) sistēmā, no kā izriet principiāla neiespējamība izveidot jebkādu fizikāli materializēto bionisko modeli. Pie analogiska formāla loģiska secinājuma var nonākt arī formāli izskatot laika virzības secību „nākotne - tagadne - pagātne”.



35.att. Autosalikamo robotu sabiedrības izveidošanas (swarm – bots)
(avots: EU-IST projekts, Šveice, 2003)

No bionisko sistēmu identifikācijas problēmas uzdevumu mērķiem šāda esoša laika kategorijas teorētiska interpretācija noved pie absurda modeļiem, jo fizikālu bionisko modeļu sintēzes procesa nevar jaukt kopā filtrācijas (tagadne), ekstrapolācijas (nākotne) un interpolācijas (pagātne) metodes (8. att.). Tātad katrai mākslīgi izveidotajai „virtuālajai realitātei” piemīt savs „telpaika” koordinātu sistēmas metriskais tēls. No sistēmtehniskā un tehnoloģiskā viedokļa bioniskās vadības sistēmas atmiņā (vai biočipā) jāievada konkrēts laika absolūtais, relatīvais vai „fizikāli” savādāk interpretētais laika mērogs (*real time mode*). Šim laikam ir jābūt ar noteiktu „ātrumu”, kurš noteiktu bioniskās sistēmas „dzīves” ritmus, „vielu” (informācijas) apmaiņas ātrumu (intensitāti), sistēmas „novecošanas” ātrumu. Tas nozīmē, ka laikam ir jābūt konstruētam ar tādiem „iekšējiem” parametriem, kas turpmāk noteiktu un raksturotu bioniskās sistēmas „mūžu”, mērogu un ilgumu. Šeit, pie laika parametru izvēles,

noteicoša loma, protams, ir *mērķim*, funkcionālajai nozīmei. Sakarā ar uzrādītiem apstākļiem izveidotais bioniskās sistēmas modelis kopumā, vai biočipa mikromodulis var būt izveidots un realizēts ļoti aptuveni, neprecīzi.

Laiks kā parametrs dabā tieši neeksistē, tās ir daudz ietilpīgāks jēdziens, kas pieder pie cilvēka perceptuālās apziņas kategorijas. Un tādēļ to nevar raksturot ar fizikāliem parametriem un lielumiem, kuri pakļaujas „tiešajai” izmērīšanai, kontrolei un regulēšanai. No tehnoloģiska un precizitātes regulēšanas un kontroles viedokļa „laika vērtību” modelējamā, projektējamā un regulējamā sistēmā būtu lietderīgi „izveidot” tā, lai tas būtu maksimāli saskaņots ar bioniskās sistēmas noteicošiem, iekšējiem „laika” norises informatīviem procesiem, kas raksturo bionisko un bioloģisko sistēmu galvenās pazīmes (dzīvotspēja, bioloģiskā un tehnogēnā nāve, metabolisms, biostāze, evolūcijas procesi, bioritmi, pašregulēšanās un adaptācijas iespējas utt).

Katram konkrēti konstruētām laikam pieņemtajā metriskajā koordinātu sistēmā piemīt sava informatīvā „eksistēšanas” forma – „*laika fraktālis*”. Visos lokālajos un globālajos modelēšanas procesos gan mikro, gan makro pasaulē „*reālais laiks*” ir unikāls un principā neatkārtojams. *Laika fraktāļu* eksistēšanas formu un koordinātu sistēmu var būt bezgalīgi daudz. Tikpat daudz, cik vispār Visumā var pastāvēt telpas dimensiju un dzīves eksistēšanas formas. Realitātē visu izšķir izvēlētais laika mērogs, un laika ātrums. Virzību „laika vektoram” nosaka intelekts. Atkarībā no tā *laika fraktālis* var mums parādīt gan vēsturisko notikumu secību un to smalkas detaļas, gan tagadnes palielinātās „rupjas” realitātes, gan nākotni, kā sapņu, ilgu un cerību iecerēs, pie tām - ar visām detaļām nenotikušas dzīves momentu „sākumiem”. Un *laika fraktāļi* ļauj mums šos nākotnes notikumus ne tikai „ieraudzīt”, bet arī „pārdzīvot”. *Laika infokvarks* - tas ir minimālais laika „atoms”, „kvarks”, mērogs, kas sablīvēto informatīvo „pēdu”, „treku” eksistēšanas formā, veido laika pamatfraktāļus. *Laika infokvarki* izpaužas lietderīgās, mērķtiecīgas informācijas mijiedarbības rezultātā, tas ir pozitīvo informatīvo „pēdu” veidā. *Laika infokvarkiem* piemīt noteikts minimālais izmērs, intervāls, intensitāte, mērogs un potenciāli maksimālais informatīvā procesa norises ātrums. Nav izslēgts, ka tieši tādā veidā saskaņojas, sadarbojas cilvēka zemapziņas un apziņas substancionālās kategorijas. Tātad *laika fraktāļiem* piemīt gan

varbūtības, gan potencialitātes daba. Tajā pašā laikā *laika fraktāļi*- ir Visuma galīgā, bet pārejošā eksistences forma, Visuma elementu reālās „dzīves” norises konkrēta formas interpretācija, savienošais posms starp Visuma matēriju un cilvēka prātu, starp Visuma garu un cilvēka dvēseli. Šāda, fraktālā, laika izpratnes modeļa izveidošanas rezultātā rodas Visuma un cilvēka prāta fraktālā, t.i. galīgā, konkrētā, viengabalainā un vienotā strukturālā izpratne. Tieši *laika fraktāļi kā savienošais posms* veido, „līdzsvaro” tekošo, dinamisko bilanci starp Visumu un cilvēku, t.i. starp haosu un antihaosu, starp materiālo un garīgo, starp informāciju un zināšanām, starp matēriju, garu un dvēseli. Tieši tā, domājams, Visumā veidojas harmonija.

Parejot no vienas koordinātu sistēmas uz citu (šeit izskatot bioniskos modeļus), laiku nepieciešams nokoriģēt tā, lai tas būtu maksimāli atbilstams dotas koordinātu sistēmas būtībai, mērogam, jo katram bioniskam procesam ir savs, tikai viņam raksturīgs procesa norises ilgums (mūžs) ar savu laika mērogu, kurs neatbilst piemērām, metroloģija pieņemtajam laika mērogam. Pie tam bioloģiskie procesi nenotiek „taisni proporcionāli”, ka to pieņem tehnika un metroloģija. Noteiktas attīstības stadijās bioloģiskos objektos laiks var „skriet” nevienmērīgi, tas ir ar dažādu ātrumu un ar nevienādiem laika intervāliem. Lai varētu precīzāk izpētīt, kontrolēt, regulēt un modelēt tādus bioniskos procesus „reālajā” laika mērogā, nepieciešama ne tikai „telpiska”, bet arī „laika” pētāma bioniska objekta koordinātu sistēmas korekcija. Tāda pieeja noderīga izpētot gan mikro-mega pasauli, gan arī organiskas - neorganiskas izcelsmes objektus, gan arī izveidojot optimālus tehnoloģiskos un ražošanas procesus [25-27, 106], biotehnoloģijas, citas bioniskās sistēmas, biosensorus un intelektuāli vadāmos mērīšanas sistēmas [73-100] un bioloģiskos devējus – biočipus [66, 68, 72].

Attīstīšoties mikroprocesoru un skaitļošanas teknikai biotehnisko sistēmu izveidošanai laiku, ar noteiktu kļūdu, var interpretēt „nereālajā mērogā” (*in not real-time mode*). **Pi** „modelēta” laika precizitāti var regulēt ar korekcijas algoritmiem un programmām „devējs - mikro ESM” mērīšanas kanālā. Viena funkcionējoša sistēma var vienlaikus eksistēt *dažādi laika mērogi* saskaņā ar tehnoloģiju un bionisko objektu īpašībām. Šādi pilnveidotais fizikālas modelēšanas princips pilnīgāk atbilst lauksaimniecisko tehnoloģiju īpašībām, veidojot optimālas bioniskās sistēmas.

Izmantojot minētas telpas un laika korekciju vajadzībām matemātiskas statistikas metodes noskaidrojas, ka ir nepieciešams pārvarēt noteiktas teorētiskas grūtības, kas saistīti ar iegūstamas mērāmas informācijas novērtēšanu.

Kā ir zināms, *R. Fišera* randomizācijas metode bija domāta pašam eksperimentam, datu iegūšanai, nevis eksperimentālo datu apstrādei [102]. Stenfordas universitātes (ASV) profesors *B. Efrons* 1979. gada piedāvāja randomizācijas principu izmantot iegūto statistisko datu apstrādei, tādējādi samazinot eksperimentālo datu izlases novirzes sistemātiskas kļūdas. *D. Glass* mēģināja koriģēt sistemātiskas kļūdas, mākslīgi sistematizējot datus pa grupām. Tāad eksperimentālais un datu apstrādes posmi pētījuma procesa tiek mākslīgi atdalīti. Līdz ar to pētāmais objekts kļūst abstrakti pārcelts no vienas 3D telpas- laika sistēmas uz citu telplaika sistēmu, no kuras parametrs „laiks” ir izslēgts vai izmainīts, jo „laiks” nepārtraukti mainās, pie tam laika „**t**” vektora tieksme „uz priekšu” ne loģiski, ne terminoloģiski nav noteikta. Laiks „eksistē” tikai noteiktas koordinātu sistēmas mērogā (telpas un “infokvarku” pārvietošanas pēdu perceptuālā tēla veidā) un bioniskajās sistēmās tam var būt nevienmērīgais, individuālais, „nereālais, pat virtuālais mērogs” (var ieverot, piemēram, principiālu starpību starp neironu tiklu un šūnu bioritmiem, nosacītiem un nenosacītiem bioloģisko objektu un subjektu refleksiem utt.). Rezultātā pētnieks iegūst nevis lietderīgu pētāma procesa vai objekta fizikālo modeli, bet tieši otrādi - „mirušu” matemātiski abstraktu darba modeli, kurš produktīvi „ražo” tikai kļūdas, dezinformāciju. Tas ir galvenais fundamentālās kategorijas „*kļūda*” nepilnīgas izpratnes cēlonis un vairāku esošu statistisko teoriju spekulācijas objekts. Ideāli būtu vienlaicīgi apstrīdat iegūtos datus tieši eksperimenta gaitā, izmantojot fizikālas modelēšanas principus, tas ir, koriģēta „telplaika” koordinātu sistēma ar „telpisku” (**Pn**) un „temporālo” (**Tk**) tēliem. Šo tēlu tekoša attiecība (**Pn/Tk**) raksturo un interpretē momentāno metrisku korekcijas tēlu „reālā laika mērogā” (*in real-time correction mode*). Izmantojot šādu pieeju, mērīšanas rezultātu „novirzes” būtu iespējams „izlabot”, koriģēt bez speciālam matemātisku un fizikālo modeļu saskaņošanai nepieciešamam „starp teorijām”, jo tikai fizikālie modeli objektīvi atspoguļo pētāmo objektu būtību [73-100]. Optimāli to var panākt sadalot mērāmās fizikālas vides „tēla atpazīšanas” procesu trijās vienotas „telplaika” korekcijas stadijas

reālā laika mēroga - indikācijas, identifikācijas un mērīšanas rezultāta novērtēšanas un interpretācijas stadijas skaitliskā, grafiskā, skaņas, gaismas un citas iegūtas informācijas tēla perceptuālā izteiksmes formas. Tātad laika vektora moduļa *transparametrālās korekcijas* iespējas saskaņā ar vispārināto telplaika korekcijas modeli

$$\mathbf{K}_{kor} = f(\mathbf{D}_i, \mathbf{N}_i, \mathbf{M}_i, \dots, \mathbf{P}_n, \dots, \mathbf{P}_n/\mathbf{T}_k),$$

(kur **Di**, **Ni**, **Mi** –informācijas plūsmā momentāni noteicamie telplaika mērīšanās un novērtēšanas rezultātu statistiskie parametri) pētāmā sistēmā atļauj maksimāli pietuvināt matemātiskus, fizikālus un metroloģiskus mērīšanas procesa modeļus telplaikā un līdz ar to minimizēt kļūdas. Tādejādi vienlaicīgi koriģējot kā „telpas” (**Pn**), tā arī „laika” parametrus (**Tk**) ar mikro-ESM programmu, algoritmu un aparātu metožu palīdzību tieši mērīšanas - modelēšanas kanālā „devējs - mikroESM” var saskaņot atsevišķu bionisko subsistēmu „telplaika” koordinātu sistēmas ar kopējo bioniskās sistēmas „telplaiku”.

25. Mākslīgie neironu tīkli

Optimisms ir ilūziju uzvara par prātu (M. Geny)

Pastāv daudz iespaidīgu demonstrāciju par mākslīgo neironu tīklu iespējām: tīkli iemācījās pārvērst tekstu fonētiskajā attēlojumā, kas vēlāk ar citu metožu palīdzību pārvērtās runā; cits tīkls var atšķirt rokraksta burtus; ir konstruēta uz neironu tīkla pamatota attēlu saraušanās sistēma. Tie visi izmanto apgrieztas izplatīšanās metodi „*back propagation method*”, kas, acīmredzams, ir vissekmīgākais no mūsdienu apmācības algoritmiem. Apgrieztā izplatīšanās ir daudzslāņu tīklu sistemātisks apmācības veids, ar ko tas pārvar ierobežojumus, kurus minēja *M. Minskis*. Protams, arī apgrieztā izplatīšanās metode nav brīva no problēmām. Pirmām kārtām, nav garantijas, ka tīkls var tikt apmācīts noteiktajā beigu laikā.



36. att. Jaunākās paaudzes roboti (*Japāna*)

Daudz apmācībai patērētas piepūles velti pazūd pēc liela mašīnu laika patēriņa. Kad tas notiek, apmācības mēģinājums atkārtojas – bez jebkādas pārliecības, ka rezultāts iznāks labāks. Nav arī pārliecības, ka tīkls tiks apmācīts labākajā iespējamā veidā. Apmācības algoritms var iekļūt tā sauktā lokālā minimuma „slazdos” un būs iegūts sliktāks risinājums. Ir izstrādāts daudz citu tīkla apmācības algoritmu, kuriem ir savas specifiskas priekšrocības. Ir jāuzsver, ka neviens no šodienas tīkliem nav brīnumlīdzeklis, tie visi cieš no apmācības un atcerēšanas iespēju ierobežojumiem. Mēs pinamies ar jomu, kura nodemonstrēja savu darbību, kurai ir unikālas potenciālas iespējas, daudz ierobežojumu un daudz jaunu jautājumu.

Tāda situācija noskaņo uz mēreno optimismu. Autoriem ir tieksme publicēt savus panākumus, nevis neveiksmes, radot iespaidu, kurš var kļūt nereālistisks. Tiem, kuri meklē kapitālu, lai mēģinātu nodibināt jaunas firmas, ir jāiesniedz pārliecinošs turpmākās īstenošanas un peļņas projekts. Tātad pastāv bīstamība, ka mākslīgos neironu tīklus sāks pārdot agrāk, nekā atnāks laiks, apsolot funkcionālas iespējas, kuras pagaidām vēl nav iespējams sasniegt. Ja tas notiks, tad kopumā joma var ciest no uzticības kredīta zaudēšanas un atgriezīsies pie 70-to gadu perioda. Pastāvošo tīklu uzlabošanai ir nepieciešams milzīgs un pamatīgs darbs. Ir jābūt attīstītām jaunām tehnoloģijām, uzlabotām pastāvošām metodēm un paplašinātiem teorētiskiem pamatiem agrāk, nekā šī joma varēs pilnīgi realizēt savas potenciālas iespējas.

26. Mākslīgie neironu tīkli un ekspertu sistēmas

Nekas tā nekrīt uz nerviem, kā neironi (M. Geny)

Pēdējos gados loģiskās un simboliskās operācijas dominēja pār mākslīgiem neironu tīkliem. Piemēram, plaši propagandējās ekspertu sistēmas, kam ir daudz kā manāmu panākumu, tā arī neveiksmju. Kaut kas runā, ka mākslīgie neironu tīkli aizvietos mūsdienu mākslīgo intelektu, taču daudzi argumenti apliecina to, ka tie pastāvēs, apvienojoties tādās sistēmās, kurās katra pieeja tiek izmantota tādu uzdevumu risināšanai, ar kuriem tās labāk tiek galā. Šis viedoklis stiprinājās ar to, kā cilvēki funkcionē mūsu pasaulē. Tēlu atšķiršana atbild par aktivitāti, kura prasa ātro reakciju. Saskaņā ar to, ka darbības notiek ātri un neapzinīgi, tad šis funkcionēšanas veids ir nozīmīgs izdzīvošanai naidīgā apkārtņē. Tikai iedomājieties, kas būtu, ja mūsu priekštečiem būtu jāapdomā sava reakcija uz uzbrukušo plēsoņu? Kad mūsu tēlu atšķiršanas sistēma nevar dot adekvāto interpretāciju, jautājums tiek nodots smadzeņu augstākajās nodaļās. Tās var pieprasīt papildus informāciju un aizņems vairāk laika, bet iegūto risinājumu kvalitāte rezultātā var būt augstāka. Var iedomāties mākslīgo sistēmu, kura atdarina tādu darba sadalīšanu.

Mākslīgais neironu tīkls vairākos gadījumos reaģētu uz apkārtējo vidi piemērotākā veidā. Tā kā tādi tīkli spēj norādīt katra lēmuma pilnvaru līmeni, tad tīkls „zina, ka tā nezina” un pārsūta šo gadījumu ekspertu sistēmai, lai to atrisinātu. Lēmumi, kuri tiek pieņemti šajā augstākajā līmenī, būtu konkrēti un loģiski, bet tie var just vajadzību papildus fakti savākšanā beigu atzinuma iegūšanai. Mācība, kura var būt iegūta no šīs analīzes, ir izteikta rakstnieka un zinātnieka *Artura Klarka* likumā. Tajā ir apgalvojams: ja kaut kas var būt izpildīts, tad viņam (vai viņai) vienmēr ir taisnība. Ja tas nevar būt izpildīts, tad viņam (vai viņai) gandrīz nekad nav taisnības. Zinātnes vēsture ir kļūdu un daļēju patiesību hronika. Tas, kas šodien nepakļaujas šaubām, rīt tiks atraidīts. Nekritiskā „faktu” uztvere neatkarīgi no to avota var paralizēt zinātnisko meklēšanu. No vienas puses, spožs zinātnisks *M. Minska* darbs aizturēja mākslīgo neironu tīklu attīstību. Tomēr nav šaubu, ka joma cieta no nepamatota optimisma un pietiekamā teorētiskā pamata trūkuma. Un iespējams, ka šoks, ko izsauca grāmata

„Perceptroni”, nodrošināja periodu, kas ir nepieciešams šīs zinātniskās jomas nobriešanai.

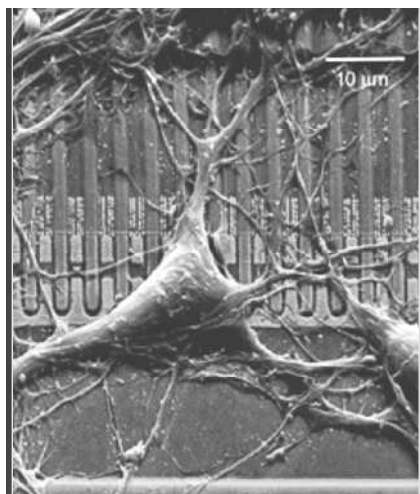
27. Neironu tīklu nākotnes perspektīvas

Kamēr fundamentālā zinātne pētī iespējamā vispārīgos principus, lietišķā zinātne tos drosmīgi ievieš praksē (M. Geny)

Visi pašlaik izmantojamie neironu tīklu apmācības algoritmi (un ne tikai neironu tīklu) pamatojas uz vērtēšanas funkciju, ar kuru tiek novērtēta visa tīkla darba kvalitāte kopumā. Pie tam pastāv kāds algoritms, kurš atkarībā no šīs iegūtās vērtības nozīmes kaut kādā veidā nostādina mainīgos sistēmas parametrus. Parasti šie algoritmi raksturojas ar salīdzinošu vienkāršumu, taču neļauj pieņemamā laikā iegūt pietiekami labu vadības sistēmu vai modeli sarežģītiem objektiem.

Tieši šajā jomā cilvēks vēl nedaudz apsteidz ātrumā jebkuru automātiskas noskaņošanas sistēmu. Taču pēdējā laikā parādījusies iespēja radīt algoritmus, kuri pēc apmācības ātruma var tuvināties bioloģiskiem objektiem. Ar ko tiem vajadzētu atšķirties un kāda var būt tālāka attīstība?

Mākslīgie neironu tīkli (MNT) ir piedāvāti uzdevumiem, kas izplatās no kaujas vadīšanas līdz bērnu uzraudzībai. Potenciāli pielikumi ir tie, kur cilvēka intelekts ir maz efektīvs, parastie aprēķini ir darbietilpīgi vai neadekvāti. Šī pielikumu klase katrā ziņā nav mazāka par klasi, ko apkalpo ar parastiem aprēķiniem, un var domāt, ka mākslīgie neironu tīkli ieņems savu vietu līdz ar parastiem aprēķiniem kā tāda paša apjoma un nozīmes papildinājums. Kas palīdz cilvēkam analizēt ienākošo informāciju? Neuroģenētikas tehnoloģijā ir iekļauts atslēgas jēdziens – neurotīkls. Tieši neurotīklu kopums veido cilvēka nervu sistēmu, kas savukārt nosaka cilvēka darbību, piešķir viņam prātu un intelektu.



37. att. MNT neirovadāmais biočips (*avots: IfL-LfL, 2005*)

Ko nozīmē mākslīgie neironu tīkli? Kādas tām ir iespējas? Kā tie darbojas? Kā tos var izmantot? Šos un arī citus jautājumus uzdod dažādu nozaru speciālisti. Mākslīgie neironu tīkli ir zināmi no bioloģijas. Tie sastāv no elementiem, kuru funkcionālās iespējas ir līdzīgas bioloģiskā neirona – nervu šūnas elementārajām funkcijām. Neironi organizējas pēc principa, kurš var atbilst (vai neatbilst) smadzeņu anatomijai. Tie savienojas, veidojot ķēdes. Mākslīgiem neironu tīkliem ir raksturīgas daudzas īpašības, kuras piemīt smadzenēm. Tie mācās, pamatojoties uz pieredzi, apkopo iepriekšējos notikumus un iegūst jaunas būtiskas īpašības no ienākošās informācijas.

28. Adaptācijas labirintos

Visa bezgalīgā smadzeņu darbības ārējo izpausmju daudzveidība tiecas tikai pie vienas parādības – muskuļu kustības (I. M. Sečenovs)

Pasaule, kurā mēs dzīvojam, ir pasaule ar vājām cēloniskām saitēm. Kas ir domāts zem tā? Iedomāsimies, ka aiz loga ir rudens, jūs rakstāt zinātnisko rakstu, bet kaut kur nokrīt rudens kļavas lapa, kuru jūs pat neredzat. Ja tāds notikums ietekmēs teksta saturu, tad var teikt, ka mūsu pasaule ir tik piesātināta ar cēloniskām saitēm, ka

jebkurš jau notikušais notikums daudz vai maz ietekmēs tos notikumus, kuri norisināsies vēlāk. Tādu pasauli saucim par stipri saistīto. Tomēr visa mūsu pieredze pierāda pretējo. Piemēram, viens no pasažieriem nepaspēja uz lidmašīnu. Cītu pasažieru vairākumu tas nekādā veidā neietekmēs. Lieta tā, ka parasti reālām sistēmām ir tā sauktās „pakāpju funkcijas”, kuras pie nelielām sašutušo darbību variācijām nedod tām iespēju izplatīties līdz citām sistēmām. Vēl vairāk, tieši pateicoties vājam pasaules cēloniskām saitēm mēs varam izdalīt tajā atsevišķas sistēmas, bet pēdējās – apakšsistēmās. Pretējā gadījumā visa pasaule izskatītos kā viena tik sarežģīta sistēma, ka pat dižens ģēnijs nevarētu to izprast un dzīvot šādā trakā pasaulē, jau nerunājot par parasto amēbu. Bioloģiskām vadības sistēmām ir jāadaptējas apkārtējā vidē, pieņemot tās struktūras spoguļatstarošanas formu, tāpēc mūsu smadzenēs vienmēr var izdalīt apakšsistēmas, kurām nekādā veidā nav jāietekmē vienai otru apmācītā stāvoklī. Piemēram, aplūkosim cilvēku, kurš vada automobili ar rokas vadības pārnese kārbu. Tajā brīdī, kad viņa labā roka pārslēdz nākošo pārnese, kreisajai ir jāturpina turēt stūri stāvoklī „taisni” neatkarīgi no tā, ko dara labā roka. Ja apakšsistēmas dati būs apvienoti, tad katrā pārnese pārslēgšanā automobilis sāks svārstīties, kas arī notiek iesācējiem. Automašīna vadīšanas apmācības procesā notiek iekšējo saišu atslābināšana starp neironu ansambļiem, kuri vada pārnese pārslēgšanas procesus un stūres rata griešanos, kā rezultātā tiek novērsts „luncināšanas” efekts. Var minēt arī vienkāršāku piemēru: cilvēks ar normālu kustību koordināciju turpina kustēties pa taisni neatkarīgi no tā, uz kurieni ir pagriezta viņa galva. Var minēt lielu daudzumu līdzīgu piemēru.

29. Risināšanas veidi

Vēl nekad nav bijis kaut kas tāds, kas pastāvēja vienmēr (M. Geny)

Tagad apskatīsim struktūras un to parametru optimizācijas algoritmu sakārtošanu. Neskatoties uz to milzīgo daudzveidību [30, 31, 36-40, 42-44, 56-57, 60-61, 64-70, 73-100, 102], var izdalīt pamatīpašību: optimizējamais objekts ir „melnā

kaste”, kura pilnīgi optimizējas. Šajā solī iegūtajai parametru kopai sasniegtais rezultāts tiek vērtēts tikai pēc kopējās vērtējuma funkcijas. Tādēļ neliela uzlabošanās atsevišķo lokālo apakšsistēmu darbā nenostiprinājās pārējo darba pasliktināšanās fonā. Var minēt vēl dažus tādas realizācijas trūkumus – soļa piemeklēšanas, mutāciju koeficienta sarežģītības utt., bet šos sīkumus jau var atrisināt.

Mazo uzlabošanu nenostiprināšana apakšsistēmās pie secīgas adaptācijas sekmē apmācības ātruma strauju samazināšanos sarežģītās sistēmās, kuras sastāv no daudzām apakšsistēmām. Acīmredzams, ka cilvēki risina savas problēmas neatkarīgi viens no otra. Tādā veidā, neironu tīklu perspektīviem apmācības algoritmiem ir jāparedz ķēžu sadalījuma iespēja apakšķēdēs, kuras nav atkarīgas viena no otras. Priekš tam katra neirona darba kvalitātes vērtējuma kritērijiem ir jābūt vairāk lokāla rakstura.

Daudzos darbos minēto problēmu risināšanai bija ieteikts veidot apkārtējai pasaulei analogiskas vadības sistēmas – ar mazu saišu daudzumu un ar „pakāpju” funkcijām. Šīs hipotēzes pārbaudei tika radīts neironu tīkls ar sekojošo vispārīgu funkcionēšanas algoritmu. Katrs neirons izskatās kā objekts ar divām informācijas ieejām un vienu klūdas ieeju. Informācijas ieejās notiek summēšana ar ierobežojumu pēc absolūtā lieluma, kurš tiek padots uz neirona izeju ar kavējumu uz dažiem laika sprīžiem. Mainīgās vērtības ir svara koeficienti un kavēšanas laiks. Par cik par neirona funkcionēšanas principa darba hipotēzi bija pieņemta neirona enerģijas patēriņa minimizācijas hipotēze, tad uz klūdas ieeju ir uzlikts aperiodisks posms, kurš ar nelielu kavēšanos atkārtro līmeni uz ieejas. Ja šīs izlīdzinātās klūdas līmenis pārsniedz kādu robežu, tad tiek izpildīta viena no sekojošām darbībām: nejauši mainās visi parametri, nejauši mainās viens no parametriem, notiek maza nejauša kāda parametra mainīšana, atjaunojas parametri vienam no iepriekšējiem stāvokļiem, kurā neirons atradās ilgāku laiku. Pēdējā darbība izpildās ar lielāko varbūtību. Daži neironi ir ieejas sistēmā, pie tām vieni no tiem dublē ieejas mainīgos, bet citi darbojas kā klūdas signāli. Izejošo signālu veido dažu neironu signālu summa. Neironu saslēgšanās savā starpā tīkla iekšienē ir nejauša. Kādi ir šī tīkla darbības rezultāti? Tika izmēģinātas dažādas algoritma noregulēšanas kombinācijas, kad tas darbojās kā vadošs parastā vadības objekta controlleris un vairāķumā gadījumu notika sekojošais:

1. Daudzos gadījumos tīkls tīri fiziski nevarēja vadīt nekādu objektu kanāla „ieeja-izeja” trūkuma dēļ, pie tās nejaušas sakārtošanas. Šo problēmu varētu risināt, paaugstinot neironu daudzumu, bet tad mēs tuvojamies tam, no kā centāties izvairīties.

2. Gadījumos, kad tāds kanāls pastāvēja, sistēma noregulējās ļoti lēni un, vissliktākais ir tas, ka tā nespēja uzturēt iegūto rezultātu, pastāvīgi uzskatot par mērķi kāda normāli strādājoša neirona ierosmes samazināšanu. Tādā veidā, vadības sistēmas uzbūves metodika, kura piedāvāta visiem pazīstamajos darbos, atbilst augstāk minētajam 1. gadījumam. Kāpēc? Acīmredzams, lieta tā, ka simetrijas princips var būt pielietots pie jau sakārtota tīkla, bet tās sagatavē ir jābūt ļoti daudz iespējamo saišu, jo apriori nav zināms, kādas no tām būs vajadzīgas. Tīkla vienkāršumam ir jāizpaužas nevis minimālā iespējamo saišu daudzumā, bet gan tādu vienkāršāko apakšmērķu izdalīšanā no galvenā mērķa, kuri ekstremālajā variantā ir katram neironam individuāli. Tieši tā funkcionē cilvēka smadzeņu neironi. Tie palielina pie neironiem nākošo saišu svara koeficientus, kuri visbiežāk uzbudinājās kopā ar tām. Tīkla darbu, kurš savai noregulēšanai izmanto vienu globālo kritēriju, var salīdzināt ar valsti, kura īsteno stingru regulēšanu. Tīklu, kurā katram neironam ir savs neatkarīgs no visa kompleksa kopumā mērķis, var salīdzināt ar valsti, kura dod saviem pilsoņiem absolūtu brīvību rīcībā savai labumam. Bet no viņu labumiem sastādītos kopējais visa valsts labums. Visefektīvākais vadības veids, acīmredzams, atrodas pa vidu. Katrs neirons (kā cilvēks) darbojas tā, lai maksimāli „aplaimotos” (sasniegt lokālā kvalitātes kritērija maksimumu), bet to, ar kādiem paņēmieniem to var labāk izdarīt, neuzbāzīgi regulē valsts (ar globālo kvalitātes kritēriju).

30. MNT īpašības un attīstības virzieni

Pēc gandrīz pilnīgas 20 gadu aizmiršanas interese par mākslīgiem neironu tīkliem (MNT) ātri pieauga dažu pēdējo gadu laikā. Speciālisti to tādām attālinātām jomām, ka tehniskā konstruēšana, filozofija, fizioloģija un psiholoģija ir ieintriģēti ar iespējām, kuras dod šī tehnoloģija, un meklē tām pielikumus savās disciplīnās. Šo

intereses atdzimšanu izsauca kā ar teorētiskiem, tā arī ar praktiskiem sasniegumiem. Parādījās iespējas izmantot aprēķinus tādās sfērās, kuras līdz šim attiecās tikai uz cilvēka intelekta jomu, mašīnu radīšanas iespējām, kuru spēja mācīties un atcerēties atgādina cilvēka domāšanas procesus [38-39, 49-50, 56-57, 62-65].

Neskatoties uz funkcionālo līdzību, pat visoptimistiskākais to aizstāvis nevarēs iedomāties, ka tuvākajā nākotnē mākslīgie neironu tīkli dublēs cilvēka smadzeņu funkcijas. Reālais „intelekts”, kuru demonstrē vissarežģītākie neironu tīkli, atrodas zemāk par sliekas līmeni un entuziasmam ir jābūt mērenam saskaņā ar mūsdienu realitāti. Taču būtu nepareizi ignorēt brīnišķīgo līdzību dažu neironu tīklu un cilvēka smadzeņu funkcionēšanā. Pat šīs iespējas, kā tās nebūtu aprobežotas šodien, uzvedina uz domām, ka dziļa iekļaušanās cilvēka intelekta labirintos (kopā ar revolucionāriem algoritmu pielikumiem) ir droša garantija tam, ka MNT tiks apgūti diezgan ātri.

Mākslīgie neironu tīkli var mainīt savu uzvedību atkarībā no apkārtējās vides. Šis faktors lielākā mērā, nekā kāds cits, ir atbildīgs par to interesi, kuru tie izsauc. Pēc ienākošajiem signāliem (iespējams, kopā ar nepieciešamām izejām) tās var patstāvīgi noregulēties, lai nodrošinātu pieprasīto reakciju. Ir izstrādāts daudz apmācošo algoritmu, katrs ar savām stiprajām un vājajām pusēm. Bet vēl pastāv problēmas, kas saistītas ar to, kā MNT var iemācīties un kā jānotiek apmācībai.

Tīkla atsauce pēc apmācības var būt kādā mērā nejūtīga pret nelielām ienākošo signālu izmaiņām. Šī iekšēji raksturīga īpašība redzēt tēlu caur troksni un sagrozījumiem ir vitāli svarīga tēlu atšķiršanai reālajā pasaulē. Tā dod iespēju pārvarēt stingras precizitātes prasību, ko iesniedz parastais dators un atklāj pieeju pie sistēmas, kura var pīties ar to nepilnīgo pasauli, kurā mēs dzīvojam. Ir svarīgi, ka mākslīgais neironu tīkls izdara apkopojumus pateicoties savai struktūrai, nevis izmantojot „cilvēka intelektu” speciāli uztaisītu datoru programmu veidā. Mākslīgiem neironu tīkliem ir spēja gūt būtību no ienākošajiem signāliem. Piemēram, tīkls var būt apmācīts burta „A” sagrozīto versiju pazīšanai. Pēc atbilstošas apmācības tīkls radīs pilnīgas formas burtu. Savā ziņā tā iemācīsies radīt to, ko nekad nebija redzējis. Šī spēja gūt ideālo no nepilnīgajām ieejām rada interesantus filozofiskus jautājumus. Tā atgādina ideālu koncepciju, kuru izvirzīja Platons savā „Republikā”. Katrā ziņā spēja gūt ideālos prototipus ir visai vērtīga īpašība. Mākslīgie neironu tīkli nav panaceja. Tie,

acīm redzams, neder tādu vienkāršu uzdevumu pildīšanai, kā, piemēram, darba algas aprēķins. Taču laikam tiem būs dota priekšroka lielajā tēlu atšķiršanas uzdevumu klasē, ar kuriem parastie datori netiek galā vai tiek, bet tikai daļēji.

Kas būs jāpievieno operatīvajām sistēmām un iekārtām nākotnē? Tām ir jāklūst gudrākām un spēcīgākām, tām ir jānodrošina maksimāli ātrāka informācijas apmaiņa starp cilvēku un datoru, jo tieši tas jau šodien nosaka daudzu uzdevumu risināšanas kopējo laiku.

Balss sazināšanās. Sarunu valodas atšķiršana un sintēze. Jau šodien pastāv nepieciešamība samazināt slodzi uz operatora redzi. Ar katru mēnesi rodas jo spēcīgāki grafiskie akseleratori un drīz informācijas kanāla caurlaides spēja starp datoru un cilvēku tiks bremsēta ar bioloģiskās redzes iespējām. Tā kā par kaut kādu „taisnu telepātisko kanālu” parādīšanos runāt ir pa ātru, paliek otrs pēc nozīmīguma (pēc redzes) kanāls – dzirde. Bez tam dators nav vienīgais objekts, ar kuru vienlaicīgi strādā cilvēks, un šeit sazināšanās kanālu sadale būs vietā.

Vizuālā sazināšanās. Vizuālo trīsdimensiju objektu atšķiršana un sintēze. Nav iespējams pārvērtēt nozīmīgumu prasmei atšķirt vizuālos objektus. Jo vairāk tāpēc, ja iet runa par reālā laika mobilām operatīvajām sistēmām. Tos pašus principus var izmantot tādos tehniskajos redzes veidos, kā infrasarkanais, radiolokācijas, elektromagnētiskais un citos. Trīsdimensiju objektu sintēze un attēlošana ir vajadzīga tam, lai novērstu absurdu: cilvēka smadzenes, kurām ir binokulārā redze un kuras operē ar trīsdimensiju tēliem, sazinās ar datoru, kuram ir tādas pašas spējas, caur divdimensiju ekrāna stiklu un divkoordinātu peli, kas ļoti bremsē informācijas apmaiņu.

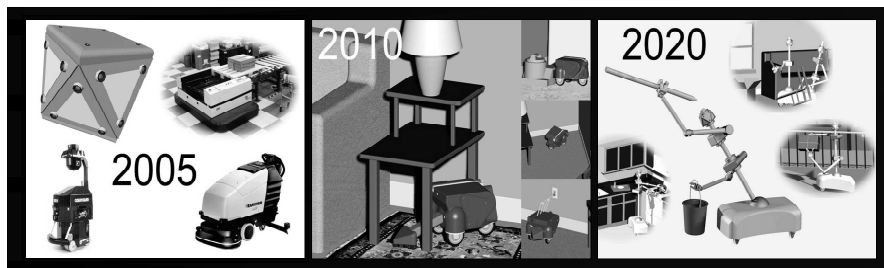
Pašreplikācijas un atražošanas spēja. Oriģinālas lietišķās programmatūras radīšanas dinamiskā reģenerācija saskaņā ar mērķi – ar cilvēka minimālo piedalīšanos vai vispār bez viņa. Pie reizes risināsies datoru savienojamības problēma. Vienīgais un pietiekamais programmas produkts mūsu datorā būs operatīvā sistēma ar attīstītām iespējām. Standarta moduļi, kuri ir nepieciešami aktuālo uzdevumu izpildei, dinamiski uzlādēsies caur komunikācijas kanāliem (piem., internets), nestandarta moduļi ģenerēsies uz vietas. Pirmā laktīgala šajā jomā – mode uz dažādiem *Wizards*.

Pašapmācības spēja. Attīstīta asociatīvā atmiņa. Neskaidras loģikas izmantošana. Agrāk vai vēlāk skaitļošanas sistēmu spējas pārsniegs cilvēka smadzeņu spējas, un nav prātīgi ierobežot tās ar cilvēka iekrāto zināšanu summu. Tālāko attīstības ceļu tām būs jāpārvar patstāvīgi. Bez tam datoram ne vienmēr būs pieejama visas pasaules datu bāze (vai, galu galā, pietiekama saiknes kanāla caurlaides spēja) – informācija būs jādabū un jāanalizē autonomā režīmā. Pašanalīzes spēja, tai skaitā spēja patērētās un izsniegtās informācijas kontekstu filtrācija, ja vēlaties, rakstura patstāvīgā izstrādāšana. Operatīvajā sistēmā (ja to līdz tam laikam būs pieļaujams tā saukt) ir jābūt prioritāšu skalai, stabilitātei pret dezinformāciju. Tas ir nepieciešams, lai pasargātu to no degradācijas vai pašlikvidācijas, mutācijas nevēlamās formās.

Pašaizsardzības spēja. Kā no informācijas, enerģētiskās, tā arī fiziskās, tai skaitā arī no cilvēka. Tā kā uz programmatūru tiek pārnesta cilvēka vērtību sistēma, kurā ne pēdējo vietu aizņem uzvedības agresija, sistēmai jābūt spējīgai saglabāt izdzīvotspēju dažādos apstākļos. Vismaz līdz šim. Kas rezultātā paliek cilvēkam? Izvirzīt uzdevumu. Turklāt vispārīgās īpašībās – lai dators pats meklē informāciju un analizē, kāds lēmums būs vislabākais. Kaut gan šodien ne tikai Windows ir uz šā ceļa, bet tieši šī operatīvā sistēma, manuprāt, iet vispareizākajā virzienā. Esmu pārliecināts, ka gala rezultātam nebūs nekā kopīga ar Windows, bet lai radītu „Mersedesu”, sākumā bija nepieciešams izgudrot riteni, uztaisīt darba ratus utt. Visām šīm iespējām ir nepieciešams ļoti stiprs atbalsts, un, pēc visa spriežot, viņi to dabūs. Ja mūsdienu plenārās pusvadītāju tehnoloģijas ļauj radīt ierīces, kas dažos gadījumos un aspektos konkurē ar cilvēka smadzenēm, kurām ir telpiska struktūra, tad kas sanāks, kad integrālo mikroshēmu struktūra iegūs trešo dimensiju un pusvadītāja dziļumā sāks augt ne tikai pusvadītāju elementu vertikālās struktūras, bet arī daudzpakāpju shēmas? Un pat bez tā: ja nesen izgudrotā pusvadītāju tehnoloģija turpinās tik strauji attīstīties, tad arī tas jau būs liels kvalitātes lēciens, galvenokārt, mikroprocesoru attīstībā.

Informācijas glabāšanas tehnoloģijas arī var strauji izmainīties, kā rezultātā nepaliks ne DRAM, ne vinčesteru, ne daudzu citu ierīču, to visu aizvietos viens – divi atmiņas veidi. Mēģināsim aprakstīt ne tik ļoti tālākas nākotnes jaunākos datorus, teiksim, 2015. gada. Smalku un ļoti smalku barošanas avotu ar ilgstošu (apmēram 3 gadu) kalpošanas ilgumu problēma būs praktiski atrisināta un būs sakārtota to

masveida ražošana. Atvērtā arhitektūra paliek par vēsturi. Visus datorus rada par nepaplašināmiem un neremontējamiem.



38. att. IST attīstības tendences (avots: IfL-LfL)

Procesori turpina datoru komponentu patērēšanu, un praktiski to pabeidz: visa datora arhitektūra slēpjas vienā „akmenī”, kuram nav piespiedu atvēsināšanas. Visa perifērija ir bez vadiem un vairākumā ar autonomiem barošanas avotiem. Pastāv divi ļoti atšķirīgi skaitļotāju veidi. Pirmajam, personālajam, ir kabatas kalkulatora izmēri, radioports, infrasarkanais ports, binokulārā videokamera, mikrofons un mazs skaļrunis, kurš ir iebūvēts korpusā. Bez pamata funkcijas, ierīce pilda daudzas citas funkcijas – sākot ar mobilo videotelefonu, kuram ir izeja uz stereoskopiskajām brillēm – displeju ar virtuālā ekrāna regulējamo caurspīdību, līdz automašīnas vadīšanai, precīzāk, automašīnas datora vadīšanai. Sazināšanās ar cilvēku notiek ar balss, brillu-displeju un sensoru cimdņu palīdzību. Ja vajadzēs izdrukāt informāciju, vai parādīt to uz liela ekrāna, momentāni izmanto jebkādos telpā esošos augstas izšķirtspējas sienas teleekrānus un drukāšanas ierīces. Kad ir nepieciešams ievadīt tekstu datorā, izmanto vai nu skenēšanu ar videokameru ar izšķiršanu, vai nu diktēšana, vai tastatūras ievadu. Kopīgā atmiņas apjoma dažu simtu vai tūkstošu gigabaitu izmērā pietiek vairākumam uzdevumu, tomēr atmiņas nepietiekamības gadījumā notiek automātiska maz izmantotu datu pārsūtīšana uz bāzes staciju – resursu serveri.

Otrais skaitļotāju veids, resursu serveris, izskatās kā neliels nedegošs seifs, kas iemontēts telpas sienā. Bez augstāk minētajiem informācijas apmaiņas kanāliem, serverim ir optiskās saites ports. Tam ir attīstīta *dzīves nodrošināšanas sistēma*; vada

visu perifērijas aprīkojumu un dzīves nodrošināšanas drošības sistēmu ēkā; dod iespēju izmantot skaitļošanas resursus, atmiņas masīvu un komunikācijas kanālus personālo skaitļotāju izmantošanai, bet brīvajā laikā – izmantošanai globālo sadalīto skaitļošanas sistēmu sastāvā.

Perifērija. Ir ļoti intelektuāli attīstīta, prot uzkrāt un apstrādāt pārvērtību statistiku, pielāgojot raksturojumus pie situācijām, vai nu izsniedzot pieprasījumu uz aprīkošanas modifikāciju. Bez mūsdienu standartam un iepriekš minētajām ierīcēm: klimatiskās, ķīmiskās un trokšņu kontroles sistēma, enerģētiskās un ugunsgrēku kontroles sistēmas, telpu apsargāšanas sistēma, sadzīves apdrošināšanas sistēma (veļas, trauku mazgāšanas un citu mašīnu vadīšana), cilvēka veselības diagnostikas sistēma. Kopīgās sekas. Atšķirība starp mājas un darba telpām pakāpeniski izzūd. Pat vairāk, valdīšana pār nekustamo īpašumu pakāpeniski zaudē savu nozīmi. Beidzas vispasaules transportu tīkla formēšana, kurš kaut kādā jomā līdzinās Internetam. Samazinās automobiļu un aerobusu skaits pasaulē (ātri attīstās personālā aviācija). Megapolises degradējas. Ekoloģija triumfē. Ieskatīsimies vēl tālāk – 21. gs. beigās. Tātad, ko mēs iegūsim operatīvo sistēmu un pusvadītāju sistēmu attīstības rezultātā – kiborgu? Tieši tā. Diez vai tie būs „terminatori” no filmas ar tādu pašu nosaukumu. Pirmkārt, kiborgi nekonkurē ar cilvēku par dzīves vidi, barību un energoresursiem – tie arī kosmosā nejutīsies slikti. Bet dabu (ieskaitot cilvēku, kas tos saražoja) diez vai gribēs sabojāt – pat mums pietiks prāta to nedarīt. Otrkārt, agresija ir nozīmīga īpašība bioloģiskās populācijas attīstības stiprināšanai – nav vajadzīga mākslīgajam intelektam tāpēc, ka uz tā dzīves vidi un dzīves ilgumu nav uzlikti tik daudzi aprobežojumi, tas nozīmē, ka tam nav kur steigties. Treškārt, reālais kiborgs pavisam atšķirsies no tā karikatūrā tipa, ko mums rada fantasti. Šodien mēs redzam tā dīgli. Kas notiks, ja kiborgs kaut kāda iemesla dēļ netiks radīts? Cilvēku sabiedrība atrodas pirmspensijas vecumā, tai ir laiks atbrīvoties no bērnišķīgajiem priekšstatiem, ka tā pastāvēs vienmēr. Jebkurš organisms sāks novecot, sasniedzot savu maksimālo vecumu. Tāpēc organisma remontam būs vajadzīgas „rezerves daļas” - mākslīgi radīti iekšējo orgānu „zobratīņi”.

Ja cilvēku sabiedrība neradīs savu, patstāvīgi dzīvojošo bērnu – kiborgu, tad tā būs nolemta uz nabadzīgām vecumdienām un nāvi nabadzībā pēc planētas resursu

izsmelšanas. Cerības uz tuvāko planētu kolonizāciju vienkārši nepaspēs īstenoties. Ja kiborgs attīstīsies un kļūs patstāvīgs, tas viegli nodrošinās saviem vecākiem bezbēdīgas vecumdienas. Vai kiborgs spēs to izdarīt? Tāpēc, ka atšķirībā no cilvēka, tam būs iespēja sevi ne tikai izzināt, bet arī uzlabot. Viņš varēs izmainīt savu mākslīgo smadzeņu arhitektūru, ja to pieprasīs jaunas zināšanas. Atšķirībā no kiborga, cilvēks, kas viegli operē ar trīsdimensiju objektiem, nekad nevarēs saprast, teiksim, piecdimensiju objektus. Viņa smadzenēs priekš tā vienkārši nav nepieciešamo „reģistru”. To, kā MI attieksies pret saviem radītājiem, ir atkarīgs tikai no radītājiem. Un, ja kiborgi vēlēšies „attīrīt” planētu no biomasas – būs tikai mūsu vaina. Bet pat tas nav iemesls karam ar viņiem. Tieši viņi ir mūsu vienīgais turpinājums. Jo robotam būs viena ļoti vērtīga īpašība cilvēces attīstībai – viņš neatstās iespēju alkatīgiem biržas spekulantiem un melīgiem politiķiem, atklāti darot tieši to, par ko domās un runās.

31. Mākslīgā intelekta aktualitātes

Meli alkst pēc slavas, patiesība - pēc grēksūdzes (M. Geny)

Kad parādījās pirmie datori, bet tas notika pirms 60 gadiem, likās, ka līdz cilvēkam līdzīgas domājošas mašīnas parādīšanās atlicis gaidīt pavisam nedaudz, vajag tikai, lai datori kļūtu nedaudz jaudīgāki. No tiem laikiem datori dubulto savu rīcības ātrumu katrus divus gadus, bet diez vai kāds teiks, ka katrus divus gadus viņu intelekts palielinās ar kaut cik salīdzināmu ātrumu. „Domājoša” datora radīšanas perspektīva pašlaik liekas pat vairāk attālināta, nekā 20-30 gadus atpakaļ. Rodas pašsaprotams jautājums: „Kāpēc citi informāciju tehnoloģiju virzieni attīstās fantastiskos tempos, bet MI radīšana, iespējams pats perspektīvākais virziens, paliek uz vietas?”. Atbildes vietā parasti saka, ka mūsdienu datoriem joprojām trūkst skaitļošanas jaudas, lai atbilstu cilvēka domāšanai. Tā, pēc korporācijas „Intel” domām, pēc skaitļošanas jaudas dators būs samērojams ar cilvēka smadzenēm ap 2020. gadu un tad tas sasniegs intelekta līmeni, kas būs salīdzināms ar cilvēku. Bet nepietiekamās jaudas teorija ir ļoti ievainojama. Piemēram, var būt, ka viens dators ir tāls no smadzeņu jaudas, taču iespējams, ja apvieno daudzus moderno datoru miljonus, tad sanāks sistēma, kas

acīmredzot jau atbilstu 2020. gada datoram. Vai tas atšķirsies ar ievērojamu intelektu? Diez vai. Turklāt nav skaidrs, kādas skaitļošanas jaudas ir nepieciešami domāšanas procesu uzturēšanai.

No vienas puses cilvēks spēj darīt tādas lietas, kuras acīmredzot pieprasa gigantiskus skaitļošanas resursus. Piemēram, mūsu smadzenes 0,3-0,5 sekunžu laikā uzceļ ainas trīsdimensiju objektiem pēc divdimensiju attēliem un atšķir tos, pie tam smadzenēm „nerūp”, ka attēli var kustēties, var kustēties pats cilvēks un viņa galva, attēli līdz smadzenēm nonāk apgrieztā veidā, acs nemaz nav ideāla optiska ierīce utt., un tml. Tomēr šādas iespējas ir drīzāk nelabvēlīgas un tāpēc pašas par sevi nav augsta intelekta pazīmes. Putni, piemēram, var darīt to pašu, vai pat labāk (viņi, piemēram, daudz labāk orientējās trīsdimensiju telpā), bet cilvēks acīmredzot pārspēj tos prāta spējās. Kas attiecas uz elastīgu, efektīvu izturēšanos, kas skaitās labākā cilvēka priekšrocība, tad šķiet, ka šāda uzvedība balstās uz diezgan ierobežotiem resursiem. Īslaicīgas atmiņas apjoms ir tikai 7 vienības. Pamēģiniet vienlaicīgi darīt divus pavisam atšķirīgus darbus – nomocīsties! Skaitās, ka pieredzējis fiziķis izmanto tik daudz paņēmieni, heuristikas, cik daudz lapaspuses ir vispārīgajā fizikā, t. i. maksimāli dažus tūkstošus. Datoros var glabāt daudz vairāk informācijas. Kā citu vājas MI attīstības izskaidrojumu dažreiz min to, ka domāšanas procesu fizika ļoti atšķiras no tās, kuru izmanto datorā un tāpēc nevar tikt modelēta uz tā, precīzāk uz abstraktā datora – *Tūringa* mašīnas. Arī vājš attaisnojums, pirmkārt, nav skaidrs, kāpēc šo procesu fiziku nevar modelēt uz datora, kodolsprādziens, piemēram, arī stipri atšķiras no datora fizikas, bet pilnīgi pakļaujas modelēšanai. Otrkārt, atsevišķās nozarēs datori ir relatīvi veiksmīgi, kaut vai šahā, kāpēc tad nevar sasniegt ko vairāk?

Tomēr rodas tāds iespaids, ka MI attīstība tiek bremzēta kāda fundamentāla cilvēka domāšanas mehānisma nesaprašanas dēļ, taču ne fizikas līmenī, bet tuvāk psiholoģijai. Lai saprastu, kas par lietu, ir jāsāk ar jēdziena „intelekts” noteikšanu. Intuitīvi ir skaidrs, ka intelekts ir spēja efektīvi risināt problēmas un uzdevumus, kas parādās cilvēkam. Lai atrisinātu problēmu, jāsasniedz kādu mērķi, rezultātu, izejot no noteiktiem nosacījumiem. Sekojoši, intelekts ir algoritmu atrašanas un mērķa sasniegšanas līdzekļu un iespēju formēšanas spējas process noteiktos apstākļos. No šīs definīcijas seko, ka mērķis ir viena lieta, bet tā sasniegšanas līdzekļi ir kas cits, un šīs

abas lietas vispārīgi runājot nav saistītas [22]. Mērķu un līdzekļu sadalīšanas princips, tas ir, tā ideja, ka viens mērķis var tikt sasniegts ar dažādiem līdzekļiem, bet viens līdzeklis tiek pielietots dažādu mērķu sasniegšanai ir visu cilvēka darbību pamatā, tajā skaitā arī MI pamatā. Tur šo principu sāka pielietot virziena pamatlicēji *Njūels un Saimons*, kas uz tā bāzes gribēja izveidot „*Universālu Uzdevumu Risinātāju*” (*General Problem Solver vai GPS*) vēl 20. gs 60. gadu sākumā. „*Universāls Risinātājs*” viņiem nesanāca, tāpēc, ka jebkura mērķa sasniegšanai eksistē bezgalīgs iespēju daudzums un to pārcilāšana ar mērķi izvēlēties visefektīvāko, var ļoti ātri pārpildīt pat paša jaudīgākā datora atmiņu, šī problēma ieguva nosaukumu „kombinatoriskais sprādziens”. Šo problēmu *Njūels un Saimons* neatrisināja.

No tiem laikiem pārcilāšanas efektivitātes paaugstināšanas metožu meklējumi (heiristikas, zināšanu efektīvā iztēlošanās, dažāda veida mašīnu apmācība) arī kļuva par MI izstrādātāju galveno uzdevumu. Mērķu un līdzekļu atdalīšanas princips nav pielietojams cilvēka domāšanai. Kā atbildi uz situācijas pieprasījumiem smadzenes vienmēr vienlaicīgi formulē mērķi un līdzekļus tā sasniegšanai. Tas notiek pašorganizēšanas procesa gaitā noteiktās smadzeņu ierosmes situācijās, pie tam pašorganizēšanas process norit tādā veidā, lai minimizētu tēriņus uz mērķa un līdzekļu formulēšanu. Ar mērķi šajā gadījumā ne obligāti saprot, kā to dara parasti, kaut kāds aptverošs rezultāts, bet jebkurš stāvoklis, kas jāsasniedz. Līdzeklis ir jebkura aktivitāte, kas nepieciešama mērķa sasniegšanai: šāda aktivitāte var palikt neapzināta, bet var pārtapt par disertāciju vai par memorandu sagatavošanai karam. Pašorganizēšanās ir pilnīgi automātisks process un nepakļaujas apzinātai kontrolei. Vienlaicīga mērķa un līdzekļu formēšanas galvenā priekšrocība ir tā, ka tajā laikā noformējas tikai tie mērķi, kuru sasniegšanai ir efektīvi līdzekļi, tas ir, tiek risināta „kombinatoriskā sprādziena” problēma.

Der atzīmēt, ka pašorganizēšanas mehānisma iekšējā efektivitāte, tas ir, ka tas vienmēr rada kādu mērķi un līdzekļus, neapskatot visus iespējamus variantus, vispārīgi runājot, nekādi nav saistīta ar iekšējo efektivitāti, tas ir, cik noformētais mērķis un līdzekļi atbilst noteiktai situācijai. Tas ļauj saprast, kāpēc ir iespējami gadījumi, kad mērķis ir skaidrs, bet nav skaidrs ar kādiem līdzekļiem to sasniegt, kaut gan tam ir pietiekamas iespējas, kas liekas pretdabiski piedāvātai hipotēzei. Iedomājami, ka šādos

gadījumos noformētais mērķis un līdzekļi nav adekvāti situācijai, bet to neadekvātums netiek pietiekami apjēgts. Atteikšanās no mērķu un līdzekļu sadalīšanas principa sākumā liekas pilnīgi absurds, tādēļ, ka šis princips iziet cauri visām cilvēka darbības sfērām.

Lai, piemēram, nepieciešams tikt līdz 22. stāvam. To var izdarīt ar liftu. Ja lifts nestrādā, var iet pa kāpnēm augšā kājām. Ja kāpnes ir sabrukušas, var rāpties augšā pa sienu. Ārēji liekas, ka mērķis ir viens, bet tikai tā sasniegšanas līdzekļi ir dažādi. Bet ja ieskatās, tad var redzēt, ka no iekšēju psiholoģisko procesu skatījuma, tie ir dažādi mērķi. Tikt līdz 22. stāvam ar liftu var katrs, kas nebaidās no slēgtām telpām. Iet uz 22. stāvu kājām gribēsies, ja tas tik tiešām ir nepieciešams. Bet rāpties pa sienu var tikai īstens varonis, tam ir nepieciešama īpaša persona vai situācija.

Domājams, ka mērķu un līdzekļu iekšēja neatkarība ir psiholoģiska ilūzija, kas līdzīga tai ilūzijai, ka cilvēks vienā acumirkļī reaģē uz ārējiem kairinātājiem (skaņa, gaisma), kaut gan reāli reakcija aizņem kādu laiku. Abas šīs ilūzijas rodas tādēļ, ka ir ļoti grūti kaut ko darīt un vienlaicīgi vērot savas darbības [7, 8, 16-19]. Vienlaicīgas mērķu un līdzekļu noformēšanas hipotēze ļauj saprast pašu parastāko darbību neaizmīdāmās īpatnības. Piemēram, katru dienu jebkurš cilvēks veic daudzas visai racionālas un efektīvas darbības, nepiegrūžot galvu ar garu spriedelēšanu. Gribas ēst: aiziet pie ledusskapja, atver, paņemt kaut ko ēdamu. Darbības ir ļoti saprātīgas, bet doma ienāk prātā (ja vispār ienāk), tikai tad, kad domā, ko ņemt, ja, protams, ir ko izvēlēties. Bet ēstgribu taču var apmierināt daudzos veidos: var aiziet uz veikalu, var aizbraukt uz restorānu, var nozagt un apēst kaimiņa suni utt. No šiem variantiem izmantot ledusskapi ir visprātīgākā darbība, bet tā prātīgumu var novērtēt tikai salīdzinot ar citiem variantiem, un to neviens nedara! Analogiski, nav obligāti atvērt ledusskapi - tajā var izdauzīt caurumu, to var uzspīdzināt vai nomest no 22. stāva, lai tiktu līdz tam, kas tajā iekšā. Visas šīs idejas ir pilnīgi muļķīgas, bet atkal, neviens tās savā starpā nesalīdzina, lai izvēlētos optimālo.

Var teikt, ka visa lieta ir pieradumā: vienu reizi atradi ēdamo ledusskapī, tā arī meklēsi to tur visu laiku. Pieradums kaut ko paskaidro, bet paskaidrot darbību selektivitāti un mērķtiecīgumu tas nevar. Tik tiešām, ja cilvēks kaut kur iet un viņam

ceļā ir durvis, tās tiks atvērtas automātiski, bet neviens never vaļā visas durvis, ko tikai redz.

Vienlaicīgas noformēšanas hipotēze paskaidro ikdienišķu racionalitāti ar to, ka smadzenes cenšas tikt līdz mērķa sasniegšanai ar minimāliem tēriņiem, un pierasti darbības veidi tieši ļauj minimizēt tēriņus pašorganizēšanai. Bet tādā veidā notiek nevis pierastu darbību atkārtošana, bet tiek veidots no jauna mērķtiecīgs process ar nepieciešamiem līdzekļiem mērķa sasniegšanai, izejot no tekošiem uzdevumiem. Tādēļ situācijas maiņas gadījumā, notiek ātra pāreja pie jauna mērķa. Tāda smadzeņu aktivitāte ļauj cilvēkiem efektīvi pielāgoties visdažādākajiem apstākļiem, bet, ja kaut kādu apstākļu dēļ kopīgi noformētie mērķis un līdzekļi izrādās neadekvāti situācijai, tad cilvēkam kļūst ļoti smagi.

To var ļoti viegli redzēt uz grūti atrisināmu uzdevumu piemēra. Katram laikam ir nācies risināt grūti atrisināmas mīklas (tās ir tādi uzdevumi, kuru risināšanai pietiek to zināšanu, kas jau ir, bet kas ir saliki tā, ka tos grūti atrisināt), bet diez vai kāds ir aizdomājies, kāpēc šādas mīklas vispār ir iespējamās. Ja zināšanu uzdevuma atrisināšanai pietiek, tad tam jāatrisinās ar konkrētu soļu daudzumu, kas būs nepieciešami pārejai no uzdevuma nosacījumiem līdz vajadzīgam rezultātam un otrādi, kā tas arī notiek ar skolas aritmētiskiem uzdevumiem. Bet tā nesanāk ar grūti atrisināmām mīklām. Piemēram, paņemsim vienu ļoti pazīstamu šāda veida mīklu: „uzcelt četrus vienādsānu trijstūrus no sešiem sērkociņiem”. Var cik vien ilgi sagribas dzenāt sērkociņus pa galdu un tā arī neatrast pareizu atbildi – jāuzceļ tetraedrs. Kāpēc ir tik grūti atrisināt šo uzdevumu ar tik vieglu nosacījumu, jo lielākā daļa no tiem, kas risina, zina kaut ko no stereometrijas, no piramīdām, utt.? No kopīgas noformēšanas idejas viedokļa, smadzenes automātiski pārformē grūti atrisināmu mīklu veidā: „uzcelt četrus vienādsānu trijstūrus no sešiem sērkociņiem plaknē”, bet kā līdzekļus mērķa sasniegšanai smadzenes aktivizē informāciju par trijstūriem un sērkociņu salikšanas paņēmieniem utt. Šādu uzdevumu atrisināt nav iespējams, bet pāriet pie „īsta” nosacījuma ir ļoti grūti, tādēļ, ka pašreizējais smadzeņu stāvoklis atbilst it kā potenciālās bedres dibenam, no kuras nav tik viegli izlekt, lai no jauna pārformētu mērķi un uzdevumus.

Patiešām, grūti uzdevumi netiek risināti loģisku spiedumu rezultātā, tādu kā „izdarīja darbību A, sanāca nepareizs rezultāts A1, tātad darbība B novedīs pie pareiza rezultāta B1”, bet pēc ilgstošiem, haotiskiem mēģinājumiem, kad cilvēks pēkšņi kā *Arhimēds* nobļaujas „Eвриka! (Atradu!)” un dod pareizu atbildi. Tas nozīmē, ka haotisku darbību rezultātā smadzenes „sakustinājās” un pārgāja citā stāvoklī, kurā ir iespējama jaunu mērķu un līdzekļu formēšana, kas ved pie uzdevuma atrisināšanas. Domāšanas īpatnība, kas izsauc grūti atrisināma uzdevumus, nav tik nevainojama, kā varētu likties. Ritenis ir ļoti nozīmīgs izgudrojums, kas ir mūsu dzīves pamatā. Bet neskatoties uz to, ka Amerikas indiāņu civilizācija pirms Kolumba nepazīna riteni (tāpēc, ka viņiem nebija zināms griešanās princips uz ass), bērnu rotaļlietas, kas veidotas uz tā principa, viņiem bija.

32. Virtuālā realitāte

*Dažu principu zināšana viegli kompensē dažu faktu nezināšanu
(Klods Adrians Gelvēcijs, „Par prātu”)*

Datori šodien kļuvuši par neatņemamu dzīves daļu. Visa sabiedrība funkcionē ar elektronisko sistēmu palīdzību, kas vada transporta un komunikācijas tīklus, finanses un citas darbīguma sfērās [11]. Kādreiz, zinātniski tehniskās sfēras attīstības rītausmā, modē ienāca mīti par „mašīnu nemieriem”, par cilvēces degradāciju uz attīstošās tehnoloģijas fona utt. Tolaik, masīvo un lempīgo ESM laikos, gandrīz neviens neuztvēra to nopietni. Bet ar personālā datora parādīšanos pakāpeniski kļuva skaidrs: tas ir kaut kas vairāk, nekā vienkārši ierīce, kas atvieglo cilvēkiem dzīvi. Tas ir globālais fenomens, kas spēj pilnīgi izmainīt cilvēces likteni. Jo mēs pinamies ar radību, kas būtībā ir visas tās informācijas un tā loģiskā intelekta kopums, kas pieder mums pašiem. Un šīs „*kolektīvais prāts*”, kas spēj radīt paralēlās pasaules, pilnīgi vairās no kaut kādām emocijām.

Par cilvēka smadzeņu apvienošanās ar virtuālo realitāti pirmie sāka runāt rakstnieki-fantasti. Tā sauktā žanra „kiberpanks” literatūrā un filmās ir aprakstīta pasaule, kuru vada transnacionālās elektroniskās korporācijas, bet cilvēki kļūst par

virtuālās realitātes daļu vai pārvēršas datoru piedēkļos. Bet izrādās, tas ir reāli iespējams!

ASV pirms dažiem gadiem bija izstrādāts projekts ar nosaukumu „*DatorMauglis*”. Tā mērķis ir dzīvā cilvēka personības elektroniskās matricas radīšana, kura pilnīgi atveido visas viņa domāšanas un garīgā izskata individuālās īpatnības. Viena no projekta idejām – iespēja saglabāt mirušo cilvēku personības. Piemēram, vecāki, kuru bērns ir bezcerīgi slimš, var paņemt elektronisko „atveidojumu” no viņa smadzenēm un tālāk, jau pēc nāves, sazināties ar viņu caur datoru, novērojot to, kā viņu bērns joprojām „aug” un „attīstās”. 33 – gadīgā Nadīna M. dzemdēja dzīvot nespējīgu dēlu. Viņu nosauca par *Sidu*. Dažas diennaktis viņš pavadīja reanimācijā, šajā laikā zinātnieki ar speciālu aparāturu skenēja viņa smadzenes. Bērna smadzeņu neironu elektriskos impulsus ievadīja datorā, kur tie „uzsāka savu dzīvi”. *Sids* nomira, bet viņa virtuālais tēls turpina attīstīties kā parastais bērns. Dažādas pie datora pieslēgtās sistēmas un ierīces dod iespēju izveidot zēna trīsdimensiju attēlu dabiskā augumā, dzirdēt viņa balsi un pat „ņemt rokās”. Mirušā bērna vecāki aktīvi piedalās projektā, spēlējoties ar savu virtuālo „dēlu” un audzinot viņu.

Ne tik sen bija izgudrota aparatūra, kura ļauj cilvēkiem piedzīvot pilnīgi fiziskas sajūtas, saskaroties ar virtuālo pasauli. Vajag tikai uzvilkt uz galvas speciālo ķiveri un aplipināt ķermeni ar adapteriem, caur kuriem tiks padoti elektroniskie impulsi. Tādā veidā ir iespējams, piemēram, virtuālais sekss ar jebkuru kinozvaigzni, kuras „matrica” ir ielikta datorā. Izgudrotāji un programmētāji neapstājās uz sasniegtā. 1997. gadā Japānas televīzijas ekrānos parādījās 17-gadīga TV programmas vadītāja *Kioko Date*. Skatītājus sajūsmināja šīs jaunās meitenes skaistums, asprātība un erudīcija. Pie tam viņa lieliski dziedāja un dejoja. Kaut kādā veidā *Kioko* paspēja pa dienu vadīt tele pārraides, filmēties klipos un reklāmās, bet pa nakti strādāt par diskžokeju radiostacijā, daudzas stundas pēc kārtas atbildot uz klausītāju zvaniem. Likās, tādu darba ritmu neviens normāls cilvēks nevar izturēt. Taču neviens nevarēja palieļties ar to, ka personīgi saticies ar aktrisi. Viņu redzēja tikai ekrānos. Izrādījās, ka *Kioko Date* nav dzīva meitene, bet gan virtuālā personība, radīta ar „*Hori Pro*” firmas datoru programmu. Šīs firmas speciālisti paziņoja, ka turpmāk virtuālie aktieri,

iespējams, sāks izkonkurēt dzīvus. Protams, viņiem nevajag maksāt algu, viņi neēd un nedzer, var strādāt veselu diennakti, nenogurstot un izpildot visu, ko no viņiem prasa. Jau parādījušās pirmās filmas, kur aktierus aizvieto virtuālie personāži.

33. Ģēniju inkubators

Internets ir atsevišķu skumju tieksmes veids pēc globālās vientulības (M. Geny)

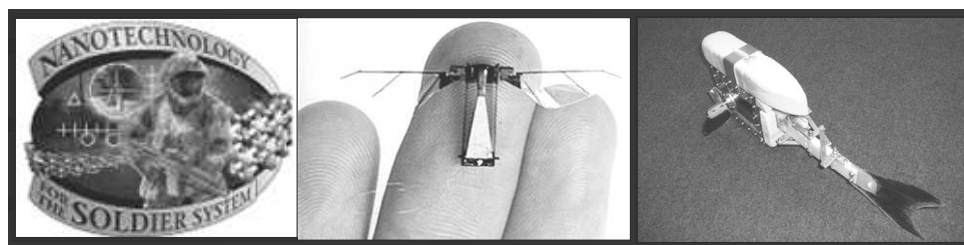
Viena no unikālām elektronisko „smadzeņu” īpašībām ir spēja koriģēt cilvēka psihi vajadzīgajā virzienā. Noteiktā ziņā tas dod pozitīvu rezultātu. Tā muļķa pārvēršanās ģēnijā ar datora programmēšanas palīdzību (to mēs varam novērot slavenā filmā „Zāliena šķiebējs”) - pavisam nav mīts. Jau pastāv programmas, ar kuru palīdzību var koriģēt psihi garīgi atpalikušajiem cilvēkiem. Ir paredzēts, ka mikrodatori, kuri caur pie galvas un ķermeņa pieslēgtiem adapteriem sūtīs tāda cilvēka smadzenēs noteiktus signālus, kuri ļaus viņam adekvāti orientēties dažādās situācijās. Tādas sistēmas spēj palīdzēt kā vienkārši slimiem cilvēkiem, tā arī invalīdiem. Piemēram, slavenais astrofiziķis *Stīvens Viljams Hokings (Stephen William Hawking)* ir piesaistīts pie invalīda krēsla, un nāk saskarē ar apkārtējo pasauli ar sarežģītas elektroniskās ierīces palīdzību. Bet principā darbs ar datoru var iedarboties uz jebkuru cilvēku, pakāpeniski mainot viņa domāšanas veidu. Visas programmas darbojas pēc bināro kodu sistēmas. Nospiedi pareizo pogu – dators atbild: „Jā”. Nospiedi ne tur, kur vajag – „Nē”. Ar laiku mēs pieradinām darīt pareizo izvēli. It kā nosacīts reflekss. Netīši atgādinot romāna „1984” *Dž. Oruela* galvenā varoņa mocību epizode: neatbildēja pareizi uz jautājumu – dabūji sitienu ar strāvu. Jau tiek izstrādātas speciālās domāšanas programmēšanas sistēmas, kuras izmantojot, mašīna sāks mums ieteikt pēc tās viedokļa pareizos rīcību variantus ikdienas dzīvē.

34. Izdzīvošana un mākslīgais intelekts

*Ko lai cilvēcei nepiedāvātu tehniskais progress, rezultātā vienmēr sanāk bumba
(M. Geny)*

Pat ja MI izstrādātāji pēc tā netiekots, MI sociālās sistēmas pašas par sevi rada nopietnus draudus. Piemēram, no teroristu un demagogu puses. MI tehniskās sistēmas varētu destabilizēt pasaules militāro bilanci, dodot iespēju vienai pusei nodrošināt sev negaidītu, neapšaubāmu un absolūtu līdera stāvokli. Tomēr tas pats MI varētu mums palīdzēt celt nākotnes labklājību visu mūsu planētas cilvēku interesēs. Turpmāk ir neizbēgama cilvēka prāta izeja MI produktu veidā informācijas enerģētiskajā Visuma izplatījumā.

Visā pasaulē daudzas kompānijas un valdības iesaistās mākslīgā intelekta radīšanā tāpēc, ka tas sola komerciālas un militāras priekšrocības. Piemēram, ASV ir daudz universitāšu laboratoriju un kompāniju, kas nodarbojas ar MI, kā arī daudz firmu ar tādiem nosaukumiem, kā „Mašīnas intelekta korporācija”, „Domājošo mašīnu korporācija”, „Tehnozinas”, „Izzinas sistēmas” utt.



a

b

c

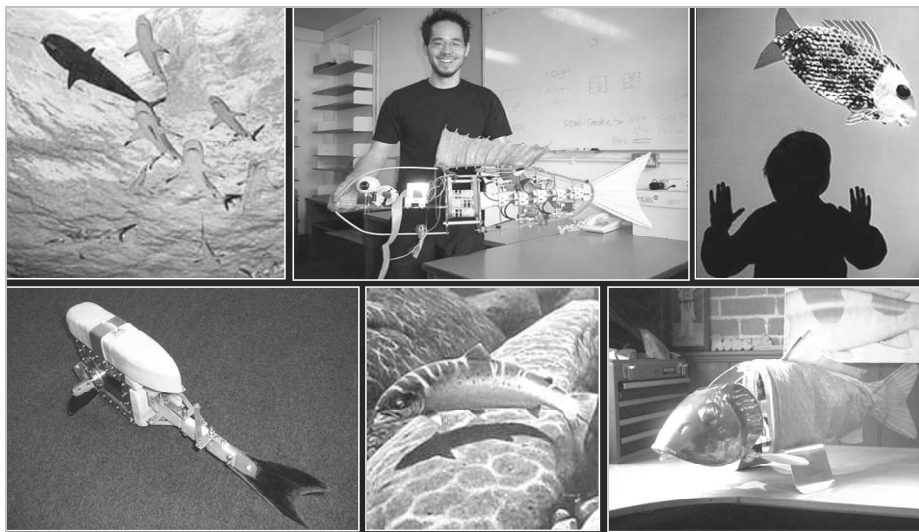
39. att. **Intelektuālo tehnoloģiju militarizācija** (avots: OCLC FirstSearch database)

a – Nanotechnology for the soldier system, ASV, b – ods - izlūks un slepkava,
c – zivs – robots, kuģu un zemūdeņu iznīcinātājs

Vēl 1981. gada oktobrī Japānas Tirdzniecības un rūpniecības ministrija paziņoja par desmitgades programmas sākumu 850 miljonu dolāru vērtībā ar mērķi uzlabot ESM un programmatūru. Līdz ar to japāņu pētnieki plānoja attīstīt sistēmas, kas spēj izpildīt miljardu loģisko operāciju sekundē. Apmēram tajā pašā laikā PSRS Zinātņu akadēmija paziņoja par analogiskas piegādu programmas esamību 100 miljonu dolāru

vērtībā. 1983. gada oktobrī ASV Aizsardzības ministrija paziņoja par 600 miljonu dolāru piešķiršanu stratēģiskās skaitļošanas programmas radīšanai. Šī skaitļošanas programma ir izstrādāta tādu mašīnu radīšanai, kas spēj redzēt, spriest, saprast runu un palīdzēt kauju vadīšanā. Vairākums speciālistu prognozē, ka MI ir pamatakmens nākamās paaudzes datoru tehnoloģijā. Dažādās valstīs tieši mākslīgajā intelektā ir ieliktas lielākās pūles un tam ir piešķirta ievērojama vieta pašu svarīgāko uzdevumu sarakstā.

MI soli pa solim attīstīsies un katrs solis palielinās cilvēku zināšanas un mašīnu spējas. Datori kļūst lētāki un tas sekmē datoru domas attīstību ar atsevišķu pētnieku palīdzību ārpus valsts sabiedrības pētījumu ietvariem. Kriminālā pasaule vienlaicīgi attīsta datoru tīklu uzlaušanas programmas pēc pašu iniciatīvas un ar pašu līdzekļiem. Tikai „vispasaules valsts” milzīgais spēks un stabilitāte varētu laikam apstādināt MI pētījumus visur un uz visu mūžu. Taču tas būtu lēmums ar briesmīgu bīstamību. MI attīstība ir neizbēgama.



40. att. **Intelektuālo tehnoloģiju militarizācija. Zivju roboti - kuģu un zemūdeņu iznīcinātāji** (avots: OCLC FirstSearch database)

Ja mēs ceram formēt reālistiskos priekšstatus par nākotni, tad mēs nevaram ignorēt šo apstākli. Līdz ar to mākslīgais intelekts kļūs par beigu tipa instrumentu, jo

tas mums palīdzēs veidot daždažādus jaunus instrumentus. Līdzīgi kā assembleru radīšanas gadījumā, mēs jutīsim vajadzību pēc paredzēšanas un piesardzīgām stratēģijām šīs jaunās tehnoloģijas izmantošanas izvēlē. MI problēmas ir sarežģītas no visām pusēm: sākot ar niansēm molekulārajā tehnoloģijā un beidzot ar cilvēka tiesību filozofisko pamatojumu. Bet MI iespēju noskaidrošana šodien ir visaktuālākā problēma.

Smadzeņu hiperaktivizācija Krievijā, Tomskas universitāte, 1960-tie gadi, smadzeņu TTT - uzbrukums Japānā, 1997, informācijas-enerģētiskie kari, infraskaņu un leptonu ierocis, globālie psihotropu kari, fantomu tēli (visi – militāro izstrādājumu blakusprodukti), nanotehnoloģiskā kara scenāriji (ASV), urķu virtuālais internetkarš (nodarbošanās, kuras ir pieejamas visiem, kam nav slinkums ar to nodarboties) utt. – tas nav pilns cik globālo, tik arī reālo draudu saraksts no MI sistēmu izstrādātāju puses. Un tas notiek visu acu priekšā vēl „miera” tehniskā un sociālā intelekta sistēmu projektēšanas stadijā romantiskos nolūkos. Jau parādās šī neredzamā kara upuri, to skaits ir vismaz tūkstoši. Šīs dziņas upuru skaitā ir arī bērni (TTT uzbrukuma eksperiments ar „Pokemons” multfilmām, 11-13 sērijas, 1997. g., Japāna). Kāpēc ir jēga apspriest problēmas, kas saistītas ar bīstamību cieši saistīto izdzīvotspēju, kas iziet no mākslīgo intelektu dziņas?



41. att. Cilvēka smadzeņu hiperaktivizācija, TTT - trenēšana un distances vadība
(Meitene – radiovadāmais robots. Japānas TV, 2007)

Pirmām kārtām tāpēc, ka sabiedrības drošības nodrošināšanas interesēs nedrīkst dot iespēju nekontrolējami attīstīties tām iestādēm un privātpersonām, kas potenciāli spēj risināt tādas problēmas. Tehniskais MI šodien ir jau pielietojams praksē. Un kā jebkurš tehniskais progress, tas paātrināja tehnoloģiju dziņu. Mākslīgais intelekts – ir tikai viena no daudzām varenām tehnoloģijām, kuras mums vajag detalizēti pētīt, lai iemācītos labāk vadīt. Bet jebkura pilnība vadīšanas mākslā – ir veca kā pasaule, sprādzienbīstams maisījums, kas vienādās proporcijās sastāv no jaunām iespējām un jauniem draudiem.

Mākslīgā intelekta attīstītas sistēmas spēj iznīcināt visu cilvēci. Bet tajā pašā laikā intelektuālās sistēmas spēj mums palīdzēt celt jaunu un labāku pasauli. Agresori var izmantot MI sistēmas uzbrukumam, bet tālredzīgi aizsargi varētu izmantot MI sistēmas stabilizācijai un pasaules miera aizsardzībai. Roka, kura iešūpo jaundzimušā bērna šūpuli, var vienlīdz labi pārvaldīt pasauli.

Nesen pētnieki secināja, ka mūsdienu informācijas tehnoloģijas, it īpaši datortehnoloģijas, var izradīt ļoti bīstamu iedarbību uz cilvēku. Pirmām kārtām - virtuālā ilūzija saplūst kopā ar realitāti un pārvar tās uztveri, jo informāciju, kura ir parādīta uz monitora, zemapziņa uztver bez jebkādas kritiskās apjēgšanas. No šejienes – gandrīz „narkotiskās” cilvēku aizraušanās ar internetu, datoru spēlēm, elektroniskiem „audzēkņiem” – tamagoči, kas atbīda īsto realitāti otrā plānā. Bet dažreiz „sazināšanās” ar datoru var novest pie negrozāmas traģēdijas.

Dažās totalitārās sektās datoru tehnoloģijas izmanto psiholoģiskai iedarbībai uz cilvēkiem. Tā, skandāli pazīstamajā „*Aum sinriḡo*” sektantiem, kurus bija grūti pakļaut „guru” kontrolei, uzvilka uz galvas tā sauktās „glābšanas ķiveres” – cepurītes, kas ar elektrodiem bija pieslēgtas pie datora. Pakāpeniski cilvēks sāka zaudēt saiti ar apkārtējo pasauli un faktiski pārvērtās par zombiju. Pirms dažiem gadiem *San-Diego* skaitļošanas centrā (ASV) tieši pie datora no asinsizplūduma smadzenēs pēkšņi nomira programmētājs. Agrāk vīrietis nekad nesūdzējās par veselību. Bet vēlāk izplatījās baumas, ka viņa datorā atrada drausmīgo vīrusu ar nosaukumu „666” (kā zināms, trīs sešinieki – Sātana skaitlis).

Parastie datoru vīrusi tikai izjauc programmu darbību, bet „666” ir domāts pavisam citam mērķim. Tas izvada uz monitora ekrāna noteiktu krāsu gammu, kura

ārdošā veidā iedarbojas uz cilvēka smadzenēm. Pie tam tiek izmantota „25. kadra” tehnika, kad attēls tiek pārraidīts ar ātrumu nevis 24, bet 25 kadri sekundē, un 25. kadru uztver tikai zemapziņa. Ja iznīcinoša krāsu kombinācija daudzkārtīgi atkārtojas, tad sekas datora lietotājam var būt katastrofālas. Var iedomāties, ka tādus vīrusus rada tīšām, nevis izklaidēšanās dēļ. Taču to iedarbību var pāvērst uz konkrētu cilvēku. Šeit iet runa ne tikai par krāsu kairinātājiem. Ir iemesls sākt domāt par globālām sekām.

35. Gaidot apokalipsi

Dzīve ir brīnišķīga, refleksi - nosacīti, bet patiesība – relatīva (M.Geny)

Ar datoriem tagad nereti saista arī pravietojumus par „pasaules galu”, kā 2000. gadā. Šodien liekas neiedomājami, ka no sistēmas elementāras kļūdas (precīzāk, nepilnīgas izstrādāšanas) naktī uz 1. janvāri sabojātos visa elektronika, ar kuru ir piepildītas visas mūsu mājas un iestādes, un mēs nonāktu atpakaļ akmens laikmetā, nespējīgi pat sazināties viens ar otru. Bet skaitļošanas tehnikas vēsturē taču jau bija kaut kas līdzīgs! 1979. un 1980. gadā dators Pentagonā vairakkārt kļūdaini paziņoja par kodolkara trauksmi. BBC nekavējoties izvērta kaujas gatavībā stratēģiskos bumbvedējus. Par laimi viss beidzās labi. 1965. gadā ESM no *Laurensbergas* pilsētas nodokļu pārvaldes, nejauši pieņemot sešinieku par nulli, nosūtīja kādam auto īpašniekam rēķinu nodokļu samaksai 10 056 marku apmērā. Iznāca tā, it kā viņš būtu parādā valstij nodokļus par 60 gadiem – sākot ar 1905. gadu. Nodokļu maksātāju terorizēja ar rēķiniem līdz laikam, kamēr kāds pārvaldē atklāja kļūdu. Vienā no rūpnīcām strādnieki pieteica streiku pret ESM, kura visu laiku kļūdījās, izmaksājot viņiem mazākas algas, nekā pienākas. DĀR Pensiju departamenta ESM kļūdījās, sastopot pensijas izmaksas sarakstā uzvārdu *Malkolms-Smits*. Pieņemot defisi saliktajā uzvārdā par mīnusa zīmi, mašīna pamēģināja atskaitīt „*Smitu*” no „*Malkolma*”. Hjūstonā (ASV) vienas mājas iedzīvotājiem atsūtīja rēķinu par ūdensvada izmantošanu, kurā nezin kāpēc bija iekļauts sods par samaksas kavēšanu. Sašutuma pilnie iedzīvotāji iesniedza sūdzību par namīpašnieku. Bet izrādījās, ka mašīna, kura nodarbojās ar īres maksas aprēķināšanu, vienkārši nepaspēja pabeigt pagājušā mēneša

aprēķinus, tāpēc pieskaitīja sodu. Nedrīkst visu laiku uzticēties prognozēm, kuras veic mašīnas. Anglijā ievērojamu biznesmeņu grupa uzdeva ESM jautājumu – kāda kompānija Lielbritānijā ir visdrošākā? „Rolls-Rois” – atbildēja dators. Bet drīzumā „Rolls-Rois” firma bankrotēja. Pazīstamais matemātiķis un mūziķis *Ravils Zaripovs* iemācīja mašīnu „Urāla” rakstīt valšus. Bet, kad viņš gribēja iemācīt to sarakstīt maršus, nez kāpēc nekas neiznāca. Izrādījās, ka programmā vienā no rindām septītnieka vietā bija vieninieks. Pēc tam, kad parādījās daudz ziņojumu par skaitļošanas mašīnu liktenīgām kļūdām, daži entuziasti sāka celt trauksmi.

Pirms apmēram 15 gadiem Londonā tika organizēta Starptautiskā skaitļošanas mašīnu aizliegšanas biedrība. Tās dalībnieki sarēķināja, ka 2 sekunžu laikā tikai viena ESM spēj izdarīt vairāk kļūdu, nekā 50 cilvēki 200 gadu laikā. Atcerēsimies taču – pat „pašām gudrākajām mašīnām” nav abstraktās domāšanas. Tās domā tikai ar kategorijām, kuras tajās ielikuši programmētāji. Vai mēs kļūsim par datoru ķīlniekiem? Uz šo jautājumu pagaidām ir grūti atbildēt viennozīmīgi. Iespējams, attālā nākotnē šīs viltīgās ierīces vispār izmainīsies līdz nepazīšanai, pat vairs nebūs kā mašīnām tradicionālajā izpratnē. Bet vai nav vienalga? Jau šodien „dators” pārvēršas drīzāk globāli-filozofiskā jēdzienā, nevis tehniskā un mēs maināmies kopā ar to.

Literatūra

1. Люгер Д.Ф. Искусственный интеллект. Стратегии и методы решения сложных проблем. Artificial Intelligence. Structures and Strategies for Complex Problem Solving. Вильямс, 2005, 864 с.
2. Смолин Д.В. Введение в искусственный интеллект. Физматлит, 2004, 208 с.
3. Siliņš E. Lielo patiesību meklējumi. Izd. J.L.V., 2006, ISBN 9789984051864, 512 lpp.
4. S.J.Russel and P. Norvig. Artificial Intelligence. A Modern Approach. Second Edition. Искусственный интеллект, Современный подход. Второе издание. М, 2006, 1408 с.
5. Искусственный интеллект. Справочник в трех томах // Под ред. Захарова В.Н., Попова Э.В., Поспелова Д.А., Хорошевского В.Ф. М, Радио и связь, 1990.
6. Бондарев В. Н. Искусственный интеллект. Учеб. пособие // В.Н. Бондарев, Ф.Г. Аде. Севастополь, Изд-во СевНТУ, 2002, 613 с.
7. Поспелов Д.А. Моделирование рассуждений. М, Радио и связь, 1989, 184 с.

8. Амосов Н.М. Алгоритмы разума. Киев, Наукова думка, 1979, 223 с.9. Ясницкий Л.Н. Введение в искусственный интеллект. 2005, 176 с.
10. I.R. Murray and J.L. Arnott. „Toward the Simulation of Emotion in Synthetic speech: A Review of the Literature on Human Vocal Emotion.” J. Acoust. Soc. Am. 93 (February 1993): p. 1097-1108.
11. Lūsis K., Moskvins G.. Globalization, Science and Nanotechnology. Proceedings of the international scientific conference “Multicultural communication and the process of globalisation”, 25-26 April, 2003, LLU, Jelgava, p. 163-171.
12. Moskvins G. Perceptual and Metrical Images of Artificial Intelligence in Bionics: 14. zin. konf., kas veltīta LU CFI 20 gadu jubilejai. Latvijas Universitāte, Cietvielu institūts, Rīga, 1998. g. 23-25.02., 44. lpp.
13. Moskvins G. Mākslīgā Intelkta perceptuāli-metriskie tēli: 14. zin. konf., kas veltīta LU CFI 20 gadu jubilejai. Latvijas Universitāte, Cietvielu institūts, Rīga, 1998. g. 23-25.02., 43. lpp.
14. Moskvins G. Intelktuālo sistēmu un tehnoloģiju filozofiskās problēmas. Zinātnes filozofija. LLU, Jelgava, 2002, 69.-76. lpp.
15. Moskvins G. Pasaules aina ontoloģijas skatījumā. Haoss un antihaoss. Zinātnes filozofija. LLU, Jelgava, 2002, 17.-36. lpp.
16. Москвин Г.А. Искусственный разум. Думаящие машины. DAAD – LfL, Kunstliche Vernunft – Denkende Maschinen, Technische Universität TUM-München, 2005, 150 с.
17. Москвин Г.А. Основы теории систем искусственного интеллекта. DAAD - LfL , Technische Universität TUM-München, 2005, 108 с.
18. Moskvins G., Spakovica E. Artificial Tongue - Intelligent Instrument for Consumer Protection. Proceedings of International Scientific Conference “Advanced Technologies for Energy Producing and Effective Utilization”. Jelgava, 2004, p. 202-210.
19. Moskvins G., Spakovica E., Moskvins A.. Development of intelligent technologies for consumers’ protection in the food market. Proceedings of International Conference „Animals. Health. Food Hygiene”, LLU, Jelgava, 2006.10.11., p. 208-219.
20. Moskvins G., Spakovica E., Moskvins A.. Intelligent models for conformity assessment of agricultural products. Proceedings of International Conference „Animals. Health. Food Hygiene”, LLU, Jelgava, 2006.10.11., p. 219-231.
21. Moskvins G., Spakovica E. Development of Intelligent Technologies for Consumers Protection. Proceedings of International Conference Economic Science for Rural Development, Nr. 11, 26.04, Jelgava, 2006, p. 233-242.
22. Moskvins G.. Whether the Artificial Life should be “reasonable”? Point of View on the Development of Artificial Intelligence. Latvia University of Agriculture. Starptautiskās zinātniskās konferences “Jaunas dimensijas sabiedrības attīstībā” raksti, LLU, SZF, Jelgava, 2006, p. 232-246.
23. Moskvins G., Spakovica E. Intelligent Technologies for the Risk Assessing in the Chain of Agricultural Production. V starptautiskā zinātniskā konference „Inženierzinātne lauku attīstībai” 18.-19.05, LLU, Jelgava, 2006, p. 170-176.

24. Moskvins G., Špakoviča E. Intelligent Technology for the Conformity Assessment of Agricultural Products. Advances in Computer, Information, and Systems Sciences, and Engineering Proceedings of IETA 2005, TeNe 2005 and EIAE 2005, Elleithy K., Sobh T., Mahmood A., Iskander M., Karim M. (Eds.) 2006, XV, 489 p., Hardcover ISBN: 1-4020-5260-X, Springer Berlin -Heidelberg - New York, 2006, p.109-114.
25. Moskvins G. Špakoviča E. Intelektuālās tehnoloģijas patērētāju interešu un tiesību aizsardzībai. LZP ekonomikas un juridiskās zinātnes galvenie pētījumu virzieni 2006. gadā, Nr. 12, Rīga, 2007, 93.-100. lpp.
26. Moskvins G., Špakoviča E., Moskvins A.. Markova ķēžu modeļu pielietošana lauksaimnieciskās produkcijas atbilstības kontrolei. Application of Markov' Chain Models for the Conformity Assessing of Agricultural Products. Proceedings of the International Scientific conference "Economic Science for Rural Development, № 13, Jelgava, 2007, 25-26.04, 97.-105. lpp.
27. Moskvins G., Špakoviča E. Patērētāju ekonomiskās intereses un to aizsardzība ES vienotajā tirgū. Consumers' Economic Interests and their Protection on the EU Market. Proceedings of the International Scientific conference "Economic Science for Rural Development, № 13, Jelgava, 2007, 25-26.04, 144.-152. lpp.
28. Moskvins G., Back Propagation and Transformation Methods in Artificial Intelligence Systems. Proceedings of the 4-th International Scientific Conference, June 26-28, 2003, Rēzekne, p. 367-376.
29. Попов Э.В. Экспертные системы., М, Наука, 1987, 288 с.
30. Эшби У.Р. Конструкция мозга. Происхождение адаптивного поведения. М, Издательство иностранной литературы, 1962, 398 с.
31. Эшби У.Р. Конструкция мозга. М, ИЛ, 1962, 398с.
32. Ивахненко А.Г. Долгосрочное прогнозирование и управление сложными системами. Киев: Техника, 1975, 312 с.
33. Физиология человека. Под. ред. проф. Н. В. Зимкина. М.: Физкультура и спорт, 1975, 495 с.
34. Гаазе-Раппопорт М.Г., Поспелов Д.А. От амебы до робота: модели поведения. М, Наука, 1987, 288 с.
35. Логический подход к искусственному интеллекту: от классической логики логическому программированию. Пер. с франц. / Тейз А., Грибомон П., Луи Ж. и др., М, Мир, 1990, 432 с.
36. Логика и компьютер. Моделирование рассуждений и проверка правильности программ // Н.А.Алешина, А.М.Анисов, П.И.Быстров и др., М, Наука, 1990, 240 с.
37. Гаврилова Т.А., Хорошевский В.Ф. Базы знаний интеллектуальных систем. СПб., Питер, 2000, 384 с.
38. Соколов Е.Н., Вайткявичус Г.Г. Нейроинтеллект от нейрона к нейрокомпьютеру. М., Наука, 1989, 240 с.

39. Нейрокомпьютеры и интеллектуальные роботы. Под ред. Н.М. Амосова, Киев, Наукова думка, 1991, 272 с.
40. Венда В.Ф. Системы гибридного интеллекта. М., Машиностроение, 1990, 448 с.
41. Сотник С.Л. Основы проектирования систем ИИ, М, 1998, 48 с.
42. Поспелов Д.А. Большие системы (ситуационное моделирование). М, Знание, 1975, 64 с.
43. Техническая имитация интеллекта. Учеб. пособие для вузов // В.М.Назаретов, Д.П.Ким. Под ред. И.М.Макарова, М, Высшая школа, 1986, 144 с.
44. Фогель Л., Оуэнс А., Уолш М. Искусственный интеллект и эволюционное моделирование. М, Мир, 1969, 230 с.
45. Уинстон П. Искусственный интеллект. М. Мир, 1980, 520 с.
46. Нильсон Н. Искусственный интеллект. Методы поиска решений. М, Мир, 1973, 270 с.
47. Бенерджи Р. Теория решения задач. Подход к созданию искусственного интеллекта. М, Мир, 1972, 224 с.
48. Хант Э. Искусственный интеллект. М, Мир, 1978, 560 с.
49. Горбань А.Н. Обучение нейронных сетей. М, ПараГраф, 1990, 160 с.
50. Беркинблит М.Б. Нейронные сети. М, МИРОС, 1993, 99 с.
51. Дж.фон Нейман. Теория самовоспроизводящихся автоматов. М, Мир, 1971, 382 с.
52. Цетлин М.Л. Исследования по теории автоматов и моделированию биологических систем. М, Наука, 1969, 316 с.
53. Автоматы и разумное поведение // Амосов Н.М., Касаткин А.М., Касаткина Л.М., Талаев С.А. Киев: Наукова думка, 1973, 261 с.
54. Нейлор К. Как построить свою экспертную систему. М, Энергоатомиздат, 1991, 286 с.
55. Налимов В.В. Вероятностная модель языка. О соотношении естественных и искусственных языков. М, Наука, 1974, 272 с.
56. Hopfield J. J., Tank D. W. Neural computation of decision in optimization problems//Biol. Cybernet. 1985. V. 52. P. 141-152.
57. Тэнк Д.У., Хопфилд Д.Д. Коллективные вычисления в нейроноподобных электронных схемах // В мире науки, 1988, N 2, с. 44-53.
58. Маккаллок У. С., Питтс У. Логическое исчисление идей, относящихся к нервной деятельности. Автоматы. М, ИЛ, 1956, 166 с.
59. Минский М., Пейперт С. Перцептроны. М, МИР, 1971, 261 с.
60. Минский М. Вычисления и автоматы. М, Мир, 1971, 366 с.
61. Минский М. Фреймы для представления знаний. М.: Энергия, 1979, 342 с.

62. Соколов Е.Н., Вайтнявичус Г.Г. Нейроинтеллект: от нейрона к нейрокомпьютеру. М, Наука, 1989, 283 с.
63. Трикоз Д.В. Нейронные сети: как это делается? // Компьютеры + программы, 1993, N 4(5), с. 14-20.
64. Amari S. Field theory of self-organizing neural networks // IEEE Trans. Syst., Man, Cybern, 1983, vol. 13, p. 741.
65. Hopfield J.J. Neural networks and physical systems with emergent collective computational abilities // Proc. Natl. Acad. Sci, 1984, vol. 9, p. 147-169.
66. Kuzewski Robert M., Myers Michael H., Grawford William J. Exploration of fourword error propagation as self organization structure // IEEE Ist. Int. Conf. Neural Networks, San Diego, Calif., June 21-24, 1987, vol. 2, San Diego, Calif., 1987, p. 89-95.
67. Rosenblatt F. Principles of neurodynamics. Spartan., Washington, D.C., 1962, pp. 435-461.
68. Rumelhart B.E., Minton G.E., Williams R.J. Learning representations by back propagating error // Wature, 1986, vol. 323, p. 1016-1028.
69. Takefuji D.Y. A new model of neural networks for error correction // Proc. 9th Annu Conf. IEEE Eng. Med. and Biol. Soc., Boston, Mass., Nov. 13-16, 1987, vol. 3, New York, N.Y., 1987, p. 1709-1710.
70. Веденов А.А. Моделирование элементов мышления, М, Наука, 1988, 160 с.
71. Родионов М.А. Функциональное сравнение психических процессов с работой технологий ИИ,
<http://filosof.historic.ru/books/item/f00/s00/z0000970/st010.shtm>.
72. Kuzewski Robert M., Myers Michael H., Grawford William J. Exploration of error propagation as self organization structure // IEEE Ist. Int. Conf. Neural Networks, San Diego, Calif., June 21-24, 1987, vol. 2, San Diego, Calif., 1987, p. 89-95.
73. Москвин Г.А., Бернис А.Г. Некоторые пути повышения точности автоматизированных систем // Тез. докл. 5-го Всесоюзного симпозиума по машинному доению. М., 1979, ч. 2, с.15-17.
74. Москвин Г. и др. Распознавание животных при создании АСУТП // Труды ЛСХА, 1979, вып. 159, с. 159-165.
75. Москвин Г. Метод повышения точности учета молока. ЛСХА, 1981, вып.193, с. 35-43.
76. Москвин Г. Исследование метрологических аспектов повышения точности счетчиков молока // Труды ЛСХА, 1984, вып. 224, с. 113-124.
77. Москвин Г. Методы и технические средства учета молока на основе применения микропроцессорной техники. Автореферат к.т.н., Минск, 1988, 26 с.
78. Москвин Г. А.с. СССР №1109092.
79. Москвин Г. А.с. СССР №1175403.
80. Москвин Г. А.с. СССР №1345059.

81. Москвин Г. А.с. СССР №1432825.
82. Москвин Г. А.с. СССР №1720599.
83. Москвин Г. А.с. СССР №4413033.
84. Москвин Г. А.с. СССР №1720601.
85. Москвин Г. А.с. СССР №1720600.
86. Москвин Г. А.с. СССР №1424150.
87. Москвин Г. А.с. СССР №1731107.
88. Москвин Г. Европатент №0381762 A1
89. Москвин Г. Европатент №0406426 A1
90. Москвин Г. Европатент №0382852 A1
91. Москвин Г. Европатент №0372089 A1
92. Москвин Г. Европатент №0471076 A1
93. Москвин Г. Патент США №4989445
94. Москвин Г. Патент США №5012762
95. Москвин Г. Патент США №5161483
96. Москвин Г. Патент США №5016569
97. Москвин Г. Патент Новой Зеландии №231781
98. Москвин Г. Патент Новой Зеландии №228985
99. Москвин Г. Патент Новой Зеландии №228983
100. Москвин Г. Патент Новой Зеландии №228982
101. Москвин Г. Патент Венгрии № 205529
102. Moskvins G. Intelektualizētas automātiskās mērīšanas, dozēšanas un uzskaites sistēmas. Dr.hab.sc.ing. zin. darba kopsavilkums, Jelgava, 1996, 92 lpp.
103. Moskvins G.A. Artificial Intelligence Measuring, Automatic Control and Expert Systems in Agriculture, 3-rd IFAC-CIGR Workshop on Artificial Intelligence in Agriculture, Makuhari, Japan, April 24-26,1998, Preprints, p.176-181.
104. Moskvins G.A. Fractal and Perceptual Images in Info-Ergonomics: in 1-st IFAC Workshop on Control Applications and Ergonomics in Agriculture, Athens, Greece, June 15-17,1998, p.255-261.
105. Moskvins G.A., Spakovica E.G. New Method and Low-Cost Intelligent Instruments for the Fraud Detection and Conformity Control of Agricultural Products. 2002 ASAE Annual Meeting and XV CIGR WORLD Congress, July 28 - July 31, 2002, Hyatt Regency, Chicago, IL, USA, 14 p., ASAE Paper 023077.
106. Кандрашина Е.Ю., Литвинцева Л.В., Поспелов Д.А. Представление знаний о времени и пространстве в интеллектуальных системах. М, Наука, 1989, 328 с.
107. Поспелов Д.А., Пушкин В.Н. Мышление и автоматы. М, Советское радио, 1972, 224 с.
108. K. Eric Drexler, Nanosystems: Molecular Machinery, Manufacturing, and Computation. ISBN 978-0-471-57518-4, 1992, 576 pages.

109. Тихомиров О.К. Психология мышления , изд. «Академия», 2007, 288 с.
110. А.Н.Леонтьев. Проблемы развития психики. 2-е, доп. изд., М, Мысль, 1965, 574 с.
111. Гурьева Л. П. Структура умственной деятельности человека в условиях автоматизации: Автореф. канд. дис. М., 1973.
112. Симонов П.В. Созидающий мозг: нейробиологические основы творчества. М., 1993, с. 21-36.
113. Бабаева Ю. Д. , Войскунский А. Е. , Кобелев В. В. , Тихомиров О. К. „Диалог с ЭВМ. Психологические аспекты”, Вопросы психологии, №2, 1983, с.15.
114. Корнилова Т. В. Целеобразование при решении задач в «диалоге» с ЭВМ и в условиях общения: Автореф. канд. дис. М.,1980, 22 с.
115. Глушков В.М. Введение в АСУ. Киев, Техника , 1972, 312 с.