



A gépi intelligencia társadalmi hatásai



A gépi intelligenciáról alkothatunk **pozitív, vagy negatív véleményt**, egy biztos: eredményei egyre jelentosebbek és egyre erosebb társadalmi hatást fejtenek ki.

A gépi intelligenciával kapcsolatos reményeket, elképzeléseket mindig **az aktuális kor technikai színvonala** futtötte. (Mechanikus robotok, elektronikus robotok, komputerizált robotok, ??)

A mesterséges intelligencia területén megfigyelhető jelenség, hogy a hétköznapi gyakorlatba bekerült eredményeket **egyre kevésbé tekintik mesterséges intelligencia eredménynek**. A mesterséges intelligencia produktumaival kapcsolatban felállított „**intelligenciaküszöb**” egyre emelkedik. Pl. a számítógépes sakk megtalálható nVidia demoprogramként, megtalálható játékboltok polcain, a szimbolikus számítást végző szoftverek az alkalmazott matematika hétköznapi szoftvereivé váltak, a beszédfeldolgozók bekerültek az operációs rendszerekbe: WARP3, Windows XP, a képfelismerő alkalmazások mindennaposak a palmtop számítógépekben, a szakértorendszer fel sem tunik a Windows súgójában, stb.

Megállapíthatjuk, hogy a gépi intellegencia eredményei **társadalmi hatással** bírnak.

† A gépi intelligencia társadalmi hatásai ..



A mesterséges intelligencia fejlődése, korszakai, interdiszciplinaritása

- **Kapcsolat az ókorral (i.e.-tol)**
A filozófia, a logika korai eredményei jelzik az emberi mentális működés törvényeinek megismerése iránti vágyat.
- **Alapvető kérdések a matematika homlokterében (1700-as évektől)**
David Hilbert Eldönthetőségi problémája: hatékony bizonyítási eljárásoknak létezik-e elvi korlátja? **Kurt Gödel:** elvi korlátok léteznek.
Nemteljességi tétel: bármely elegendően erős matematikai rendszerben lehet olyan formulákat megfogalmazni, amelyek igazsága, vagy hamissága nem bizonyítható a rendszeren belül. **Church-Turing tétel:** a Turing-gép képes minden kiszámítható függvényt kiszámítani.
Kezelhetetlen problémák, NP-teljes problémák: exponenciális méretfüggőség jellemzi őket.
Neumann János és Oskar Morgenstern játékelmélete megalapozza az ágens döntéseit.



A gépi intelligencia társadalmi hatásai ..



- **Az emberi mentális működés megismerése: a pszichológiai megalapozás (1800-as évektől)**
Introspekció, behaviorizmus, pszichoanalízis, kognitív pszichológia.
- **Az első lépések (1940 - 1956)**
Legelső MI eredmények: **McCulloch-Pitts neuron**, **Neumann: automata elmélet**, **Newell és Simon: Logic Theorist**, **McCarthy: „Artificial Intelligence”**.
- **A túlzó lekesedés korszaka (1957 – 1970)**
Newell-Shaw-Simon: General Problem Solver. Tételbizonyítók. Arthur Samuel: öntanuló dámajáték: a számítógép nemcsak arra képes, amire megtanították. (Tanítsuk meg tanulni.) **Lisp, Prolog. Mikrovilágok.**
- **Reális célkitűzések korszaka (1970 - 1980)**
Az 1960-as évek **optimizmusa**, mely az "általános problémamegoldó" (GPS) megalkotásához fűződött, a gyakorlati tapasztalatok eredményeképpen **lehiggadt**, és a kitűzött cél átalakult. Konkrét, egy szűk és jól definiált területen eredményesen alkalmazható rendszerek létrehozását tűzték ki a 70-es évek közepétől. Előtérbe került a **szimbolikus tudás-szemléltetés** és kidolgozták az első üzleti célú **szakértőrendszereket**. Dendral, Mycin, úttörők.



A gépi intelligencia társadalmi hatásai ..



- **Tudásfeldolgozás ipari méretekben (1980-tól)**
A szakértorendszerek fénykora. Ötödik generációs japán project.
Napjainkban szakértorendszerek százai vannak alkalmazásban (Nyugaton) a gazdaság, politika, gyógyászat, pénzügy, oktatás eszközeiként és még nagyon sok más területen alkalmazást nyernek.
- **A neurális hálók újra előtérben (1986-tól)**
Többretegu hálók, sikeres tanító algoritmusok, felügyelet nélküli, önszervező működés.
- **A jelenkor: a lopakodó gépi intelligencia (1990-tól)**
Előtérbe került a bizonytalanságkezelés: neurális hálók, fuzzy rendszerek. Globális keresők: genetikai algoritmus, szimulált kihűlés, tabu keresés. Fejlett beszédfeldolgozás. Humanoid robotok. Intelligens gyártórendszerek. Ágens megközelítés. Embervero sakkprogramok.



A gépi intelligencia társadalmi hatásai ..



Vámos Tibor akadémikus jóslata, 1988:

"Minden jel arra mutat, hogy a következő évtizedben olyan áttörés lesz, amely integrálja az eddigi eredményeket és erőfeszítéseket a számítógépes információs rendszerekben, jól használhatóan hozzáférhetővé és kezelhetővé teszi az emberi tudás minden ágát azok számára, akik a hozzáférhetőség eszközeivel, felkészültségével tudtak lépést tartani.

Ennek kapcsán rendkívüli újabb fordulat várható a gazdaságban, az emberi munka módszereiben, a foglalkoztatási struktúrákban, és a társadalmi viszonyokban. Az előző években a programozás alapfogalmaira vonatkoztatott számítógépes írni-olvasni tudás így nyer új tartalmat egy nemzetközi és nemzeti beszéd- és értelemkészség kialakulásában. A világpiacon az fogja megállni a helyét, aki ebben a beszéd- és értelemkészségben partnerré tud válni.,

(Gábor András: Szakérto rendszerek'88 SzámAlk, Budapest, 1988. p500)





A gépi intelligencia társadalmi hatásai ..



A (részben) mesterséges intelligencia kutatására indított nagyobb **projektek**:

<u>Projekt</u>	<u>ország</u>	<u>idoszak</u>	<u>összeg (millió\$)</u>
DARPA	USA	1984-88	600
Alvey	Anglia	1984-88	525
ESPRIT	EGK	1984-88	1330
5.gener.	Japán	1982-91	500



- A projektek hatására megnőtt a kereslet az MI termékek piacán is. Az MI termékek hasznosulásához elengedhetetlen **a társadalom fogadó-készségének** megteremtése, amely az emberek részéről újszerű gondolkodást és gépesített problémamegoldás alkalmazást jelent. A mesterséges intelligencia **áthatja** a gazdaságot, a kereskedelmet, az oktatást, a tudományt, a gyerekek számítógépes játékait egyaránt. Betör a sportba, megreformálja a gyógyítást.



A gépi intelligencia társadalmi hatásai ..



Az MI rendszerek beszerzésének és **alkalmazásának indukáló tényezói:**

- Növekvő **verseny**: üzleti szempontok, termelékenység növelése
- A modern társadalmakban erős az igény a **specializált szaktudásra**
- Erős az igény az **emberi problémáktól mentes** helyettesítőkre, mindezt könnyen kezelhető, megbízható szoftver-hardver rendszer alakjában képzelik el.

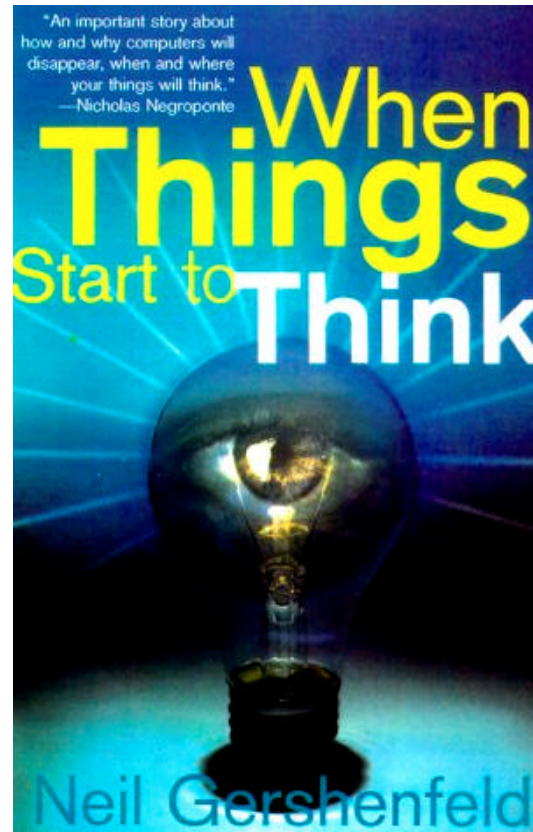
Terjedésének **gátló tényezói:**

- Idegenkedés az újtól
- Félelem a szellemi dominancia elvesztésétől
- Kísérletező jellegű empirikus tudomány, történelmi beágyazottsága gyenge.

† A gépi intelligencia társadalmi hatásai ..



- **Neil Gershenfeld** *When Things Start to Think* , Amikor a dolgok elkezdenek **gondolkozni** címu könyvében több futurisztikus téma, mint pl. beimplantált mikroáramkörök mellett foglalkozik a mesterséges intelligencia új korszakának emocionális és jogi oldalával is. Vázolja azt a kort és életet, amikor a dolgok mind gondolkodni fognak.

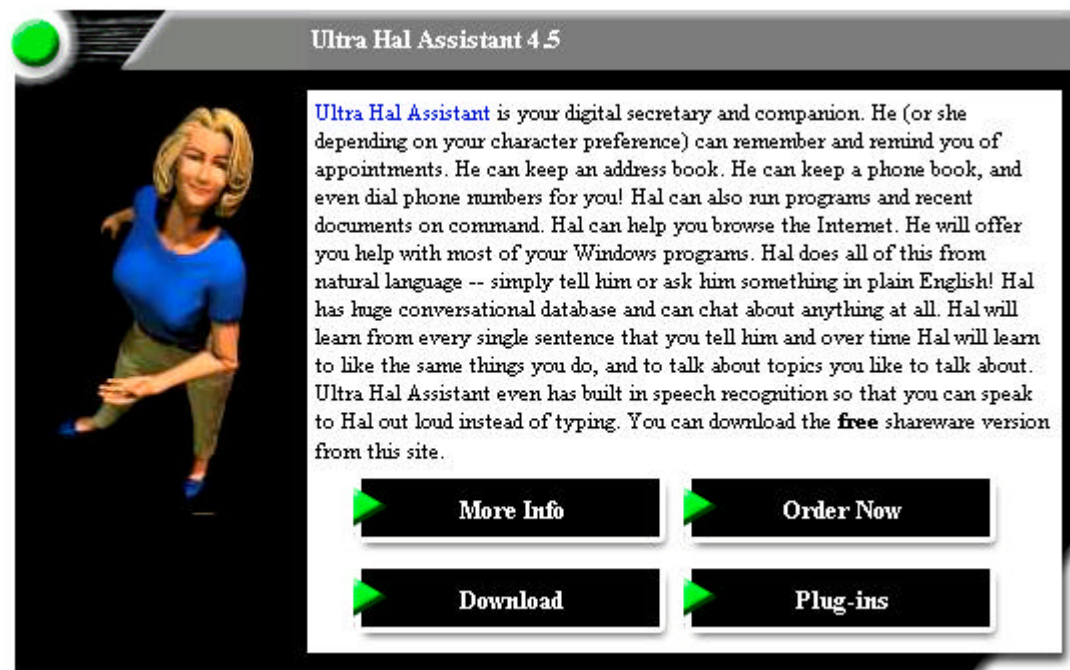




A gépi intelligencia társadalmi hatásai ..



- **ChatBot**-ok, beszélgetőrobotok sokasága található **Simon Laven** honlapján. Ezek egy része öntanuló, és személyi titkár szerepre is vállalkozik. Képviselőjük **Ultra Hal Assistant**, egy sokféle személyiség egyikét felölteni képes titkár. (Vesd össze Turing jóslatával.)



<http://www.zabaware.com/>



A gépi intelligencia társadalmi hatásai ..



- Egyik képviselőjük, Júlia elérhető az Interneten is:
<http://www.simonlaven.com/>



- Jelenlegi képességeikről ad információt a következő táblázat:
- Ezen virtuális személyiségek jövőbeni elterjedése, intelligenciájuk növekedése várható. A jövőben a számítógépből kilépve, 3D hologram alakjuk és Internetes kapcsolatuk a központi nagy intelligencia- és problémamegoldó-szerverhálózattal várható.

	Free Version	Advanced Version
Can chat with you	YES	YES
Can assist you	YES	YES
Can speak out loud	YES	YES
Ability to learn	YES	YES
Plug-in support	YES - Limited	YES
Scripting support	YES	YES
Built in Voices	10 voices	27 voices
Built in Characters	1 character	20 characters (shown below)
Personalities	1 personality	3 personalities
Built in Skins	2 Skins	5 Skins
Price	FREE	\$34.95
Speech Recognition - You can speak out loud instead of type	NO - You must type to Hal	YES
File Size	11 megabytes	80 megabytes



Vélemények az MI mellett és ellen



A mesterséges intelligencia kutatását végigkísérik a filozófikus, vagy gyakorlatias megközelítések a végső cél: az emberrel összemérhető gondolkodási képességu gép kifejlesztésének lehetőségéről, ill. lehetetlenségéről. Az ideológiai alapállástól sem független vélekedések, meggyozodések többféle szintet képviselnek:

- **Nem** lehet intelligensnek tuno, vagy ténylegesen intelligens gépeket létrehozni.
- **Gyenge MI nézet:** Lehet a gépek cselekvéseit úgy kialakítani, *mintha* intelligensek lennének.
- **Eros MI nézet:** Az intelligensen cselekvo gépeknek **valódi tudatosságuk** lehet.

A tárgy témájából eredoen az 1. pontot nem kommentáljuk, bár néhány ilyen irányú véleményt bemutatunk.



Vélemények az MI mellett és ellen ..



Ellenérvek a gyenge MI-vel szemben, még az a szint sem elérhető:

- Vannak dolgok, amiket a számítógép nem lesz képes megtenni, függetlenül attól, hogyan programozzák be
- Az intelligens programok fejlesztésének bizonyos módszerei hosszú távon megbuknak
- A megfelelő programok szerkesztése, mint feladat, nem valósítható meg.

A fenti ellenérvek legjobb cáfolata a lehetetlennek tartott muködések megvalósítása.

Hihet valaki az Eros MI-ben anélkül, hogy a Gyenge MI-ben hinne: intelligensen viselkedo gépeket nem lehet létrehozni, de ha valaha mégis sikerülne, akkor azok bírhatnak tudattal.



Vélemények az MI mellett és ellen ..



Alan Turing a **COMPUTING MACHINERY AND INTELLIGENCE** címu cikkében bevezetett „imitation game”, imitációs játék gondolkísérletével felvetette, hogy nincs több jogunk intelligensnek tartani egy válaszaiban magát nonek beállító férfit, mint egy „magát” embernek beállító gépet. Ezt a kísérletet Turing tesztként ismerjük és alkalmasnak tartják az Eros MI kritériumnak való megfelelés eldöntésére.

- A kutatóknak a gép szándékosságáról, tudatosságáról, eros MI-re való alkalmasságáról álljon itt **néhány vélemény**: .

*"Hiba lenne azt gondolni, hogy a gépek egy napon elérik intelligenciafokunkat és itt megállnak, vagy azt feltételezni, hogy értelmünk mindenkor megállja helyét a gépekkel szemben. Ahhoz, hogy megőrizzük ellenőrző szerepünket a gépek felett, az lenne szükséges, hogy minden tettünket és cselekedetünket a jelenleg a földön lévő legmagasabb intelligenciaszint befolyásolja tartósan." (Marvin L. **Minsky**)*



Vélemények az MI mellett és ellen ..



Karl Steinbuch: " Nincs nyomós érv, ami cáfolná azt az állítást, hogy a számítógépek intelligensebbek lehetnek az embernél. Ez ugyanolyan lenne, mintha kijelentenénk, hogy lehetetlenség az olyan jármű, amely gyorsabban halad, mint ahogy az ember megy. Tíz - húsz év múlva már nem lesz vita tárgya."

Kérdés: Szükséges-e és lehet-e az ilyen magasan fejlett automatáknak valamiféle erkölcsrendszert közvetíteni, mely visszatartaná oket attól, hogy a társadalomnak kárt okozzanak ?

Steinbuch: "Ez a legizgalmasabb kérdés. Mondottam, hogy van egy dualisztikus felfogás, amely az emberi és a gépi intelligencia között alapvető különbséget feltételez. Ebből kiindulva azt kell mondanunk, nincs mód arra, hogy ellenőrzésünk alatt tartsuk a gépeket. Hiszen akkor a gépek morálisan nem tudnak cselekedni. Veszélyes azonban itt megrekedni! Igenis meg kell kísérelni programnyelvre fordítani azokat a fogalmakat, amelyeket mi morál, etika, társadalmi együttélés, stb. szavakkal jelölünk, s ezeket aztán bevinni a számítógépbe. Csak így van esélyünk arra, hogy a jövőben ne a számítógépek uralkodjanak felettünk. A morált nem szabad valami gyakorlatilag érthetetlen dologként feltüntetnünk, hanem inkább részletesen elemezni kell, majd a számítógépnek pontosan meg kell mondani, melyek azok a dolgok, amiket nem szabad tenni."



Vélemények az MI mellett és ellen ..



*"Az eredeti kérdés : Tudnak-e a gépek gondolkodni? véleményem szerint túl értelmetlen ahhoz, hogy megérdemeljen egy vitát. Mindamellett úgy gondolom, hogy századunk végére a szavak és az általánosan oktatott álláspontok olyan erosen meg fognak változni, hogy bárki ellentmondás kiváltása nélkül beszélhet majd a gépi gondolkodásról." (Alan **Turing**,1950)*

Turing jóslata: 2000 körül egy 10^9 elemu tárolóval rendelkezo számítógépet már kell tudnunk úgy programozni, hogy egy kérdezovel 5 percig tudjon társalogni és a dötnöknek nem lesz 70%-tól nagyobb esélye a helyes döntésre.

„I believe that in about fifty years' time it will be possible, to programme computers, with a storage capacity of about 10^9 , to make them play the imitation game so well that an average interrogator will not have more than 70 per cent chance of making the right identification after five minutes of questioning.”

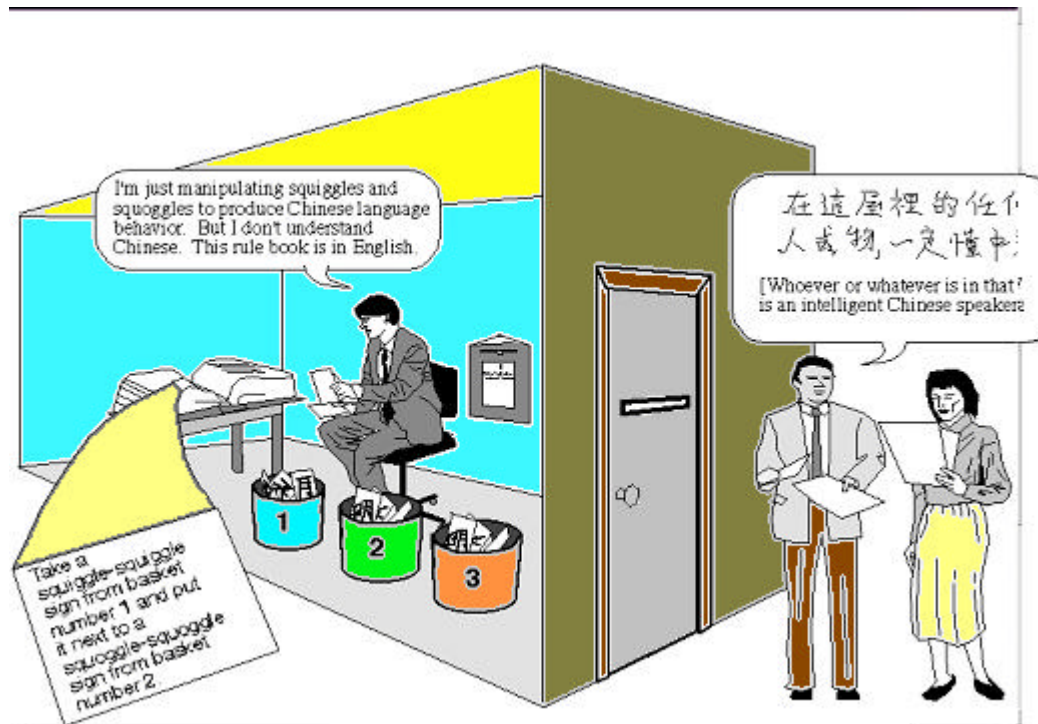
A jóslat helyességét ítéljük meg magunk.

† Vélemények az MI mellett és ellen ..

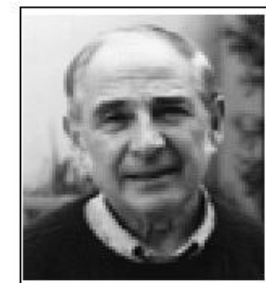


"Bizonyára egy számítógép a program irányítása alatt áll, de ez nem jelent megszorítást, egyszerűen amiatt, hogy programozhatjuk a számítógépet arra, hogy kreatív legyen." (**Computers and Thought**)

"Számítógépprogram lenne az emberi elme? Nem. Egy program csupán manipulációkat végez a szimbólumokkal, ám az agy jelentést is kapcsol hozzájuk." (John R. **Searle**)



Searle a mesterséges intelligenciával kapcsolatos fenntartásait a híres, "**Kínai szoba**" néven ismertté vált gondolkísérletével fejezte ki.



John Searle

Forrás:

<http://www.macrovu.com/CCTGeneralInfo.html>



Vélemények az MI mellett és ellen ..



"Vegyünk egy nyelvet, amelyet nem értünk. Én például nem tudok kínaiul. Számomra a **kínai** írás megannyi értelmetlen kriksz-kraksz. Most tegyük fel, hogy magamra hagynak egy szobában, ahol kosarakban kínai írásjelek vannak, ezenkívül pedig kapok egy angol nyelvű szabálygyűjteményt, amely megmondja, hogyan kell bizonyos kínai írásjeleket összekapcsolni más kínai írásjelekkel. A szabályok a jeleket teljes egészében az alakjukkal azonosítják, és nem teszik szükségessé, hogy bármelyiket is értsem közülük. A szabályok például ilyen eloirásokat tartalmazhatnak: 'végy egy így és így tekergo jelet az egyes számú kosárból és tegyél mellé egy ilyen meg ilyen kacskaringót a kettes számú kosárból.'

Képzeld el, hogy kínaiul tudó emberek kis szimbólumhalmazokat adnak be kívülről, én pedig válaszképpen átalakítom ezeket a szabálygyűjtemény utasításai szerint, majd az új szimbólumhalmazt visszajuttatom. Ez esetben a szabálygyűjtemény a 'számítógépprogram', akik írták, azok a 'programozók', én pedig a számítógép vagyok. A szimbólumokkal teli kosarak alkotják az 'adatbázist', a beadott szimbólumhalmazok a 'kérdések', az általam visszajuttatottak pedig a 'válaszok'.

Most tegyük fel: a szabálygyűjteményt úgy állították össze, hogy a 'kérdésekre' adott 'válaszaim' megkülönböztethetetlenek egy született kínai feleleteitől. ... Ezzel **kiálltam azt a Turing próbát**, amely azt vizsgálja, tudok-e kínaiul. Ennek ellenére valójában egyáltalán nem tudok kínaiul. ... A digitális számítógépek (is) csak manipulálják a formális szimbólumokat a programban foglalt szabályok szerint."



Vélemények az MI mellett és ellen ..



Searle érvei számtalan bírálatot váltottak ki.

- **A rendszerelvu válasz** (Berkeley)

„Miközben igaz az, hogy a szobába zárt egyetlen személy nem érti meg a történetet, valójában a helyzet az, hogy o csak egy része a teljes rendszernek, és a **rendszer az, amely a történetet megérti...**”

- **A robot válasz** (Yale)

„...Tételezzük fel, hogy egy **számítógépet teszünk a robot belsejébe**, és ez a számítógép...ténylegesen oly módon működtetné a robotot, hogy amit a robot csinálna, az nagyon hasonló lenne az érzékeléshez, járáshoz, mozgáshoz, szögbeveréshez, evéshez, iváshoz – bármihez, amit csak akarunk....”

- **Az agyszimulátor válasz** (Berkeley és MIT)

„Tételezzük fel, hogy olyan programot tervezünk, amely...a **neuronkiszülések tényleges sorozatát szimulálja** a szinapszisoknál, ahogyan az az anyanyelvu kínai agyában zajlik,...Nos, ebben az esetben bizonyára azt kellene mondanunk, hogy a gép megértette a történeteket, és amennyiben ezt tagadjuk, vajon nem kellene-e azt is tagadnunk, hogy az anyanyelvu kínaiak megértették a történeteket?”



Vélemények az MI mellett és ellen ..



- **A kombinációs válasz** (Berkeley és Stanford)
„**Egyesítsük a három elozo választ!** Képzeljünk el egy robotot, amelynek a koponyaüregébe egy agy formájú számítógépet helyeztek, képzeljük el, hogy a számítógép az emberi agy minden szinapszisát programozza, képzeljük el, hogy a robot egész viselkedése megkülönböztethetetlen az emberi viselkedéstől, és az egészet egy egységes egészként gondoljuk el, ... Ebben az esetben a rendszernek minden bizonnyal intencionalitást kellene tulajdonítanunk.”
- **Más elmék válasz** (Yale)
„Honnan tudjuk, hogy más emberek megértik a kínait, vagy bármi mást?
Csak a viselkedésükből. A számítógép pedig ugyanúgy átmegy a viselkedési teszteken, mint ok (elvileg), így ha kogníciót tulajdonítunk más embereknek, elvileg a számítógépnek is azt kell tulajdonítanunk.”
- **Sok ház válasz** (Berkeley)
„**Legyenek bármik** is azok az oki folyamatok, amelyeket alapvetőnek tekint az intencionalitáshoz, ...végülis képesek leszünk olyan eszközöket készíteni, amelyek rendelkeznek ezekkel az oki folyamatokkal és mesterséges intelligenciával....”

Searle megválaszolta ezeket a felvetéseket.



Vélemények az MI mellett és ellen ..



Más oldalról közelítette meg a problémát és építette fel bírálatát nagyszerű könyvében **Roger Penrose**. Kvantumfizikai meggondolásokkal arra jutott, hogy az emberi gondolkodásban **nemalgoritmikus** (Turing-géppel el nem végezhető, azaz számítógépre nem átültethető) összetevőknek is kell lenniük. Hosszasan kutatja a tudat jelenségét, a biológiai és mesterséges tudat fizikai helyét. Az *ihlet, meglátás, eredetiség* mentális fogalmak esetében felveti a nemtudatos értelemről való eredésüket. (V.ö. Freud gondolataival.) Több példával támasztja alá azon meggyozódását, hogy a **gondolkodás** működhet **szavak nélkül**, nem verbális tudati képek formájában is. Igaz, kérdojeles formában, de felveti az **állati tudatosság** lehetőségét is.

Roger Penrose: A császár új elméje Akadémiai Kiadó, Budapest, 1993. p489.

Az MI legádázabb bírálóinak egyike **Hubert Dreyfus**.

Kiindulásában a számítógépet egy **számdaráló** egységnek fogja fel, mely csak elkülöníthető, elemi, alternatív választásokat tud kezelni. Az MI kutatók kezdeti eredményeit játékoknak nevezte, mindaddig, míg a 60-as évek végén ki nem kapott a MacHack nevű sakkprogramtól. Legutóbbi, 1970-es állítását, miszerint az MI alkalmazások sohasem lesznek képesek **kitörni abból a konzisztens ismerethalmazból** álló mikrovilágból, amelyet beprogramoztak nekik, nehezebb megdönteni. (A megoldást az öntanuló rendszerek adják.)

† Vélemények az MI mellett és ellen ..

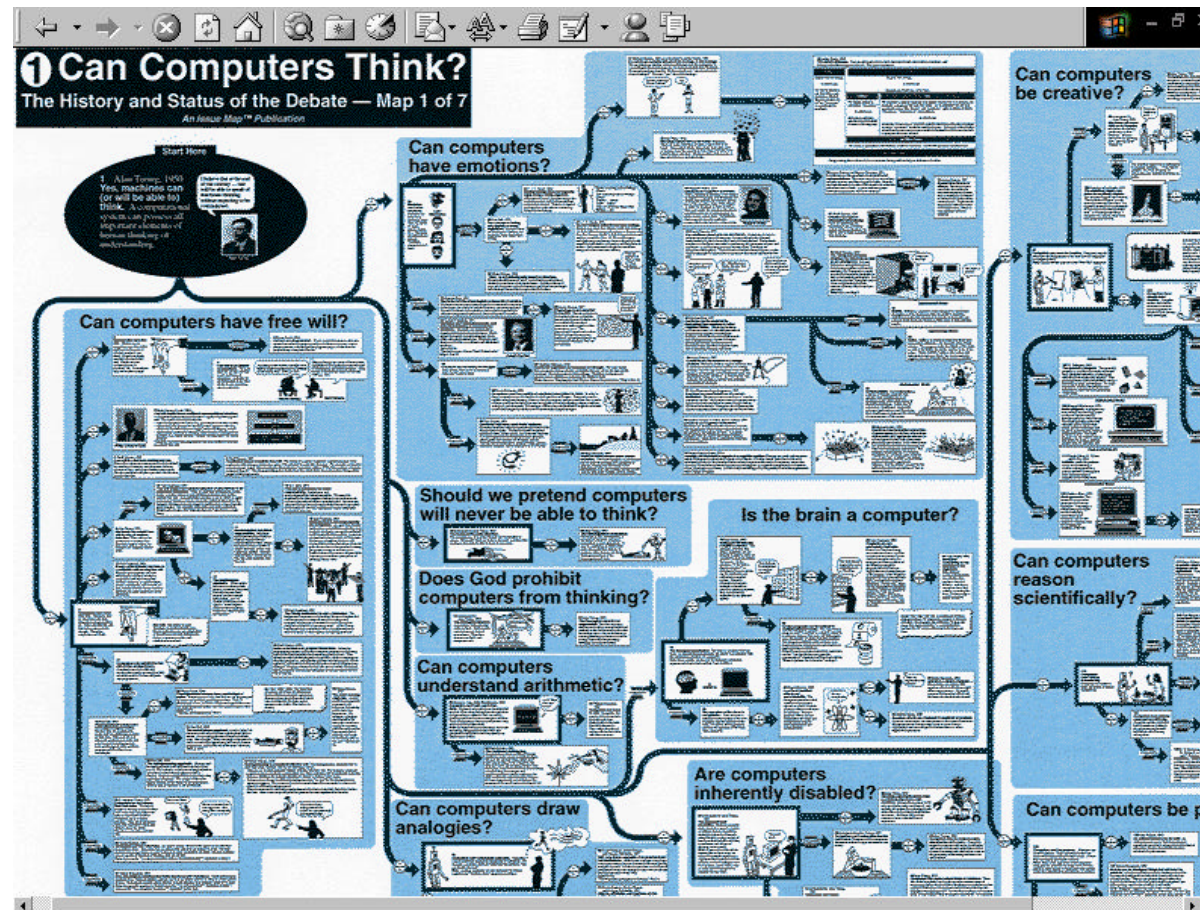


Joseph Weizenbaum és **Terry Winograd**, maguk is számítógépes kutatók, szintén az MI bírálói közé tartoznak. Weizenbaum az ELIZA program megalkotásával tulajdonképpen a Turing teszt könnyű teljesíthetőségét kívánta bizonyítani. (Ami azért teljesen nem sikerült.) Mivel programja több embert is megtévesztett, Weizenbaumban az a meggyőződés alakult ki, hogy **a számítógépet vissza kell tartani az emberi társadalom bizonyos területeitől**. Ezeket a problémákat olyan előjelekként fogta fel, melyek a számítógépek jövőbeni erőteljes invázióját vetítik előre. Szemléletére az a katonai irányzat is rányomta bélyegét, melyet az amerikai Védelmi Hivatal MI irányú kutatásai tükröztek. Kritikai beállítódása a morális kérdések felé vett irányt. Véleménye szerint, bár egy robot rendelkezhet azokkal a tapasztalatokkal, melyekkel egy ember, soha nem fog rendelkezni az emberek genetikai örökségével, a szülő - gyermek viszony, vagy a nemiség érzéseivel. **Szavaiból az emberiség számítógéptől való féltése érződik ki.**

† Vélemények az MI mellett és ellen ..



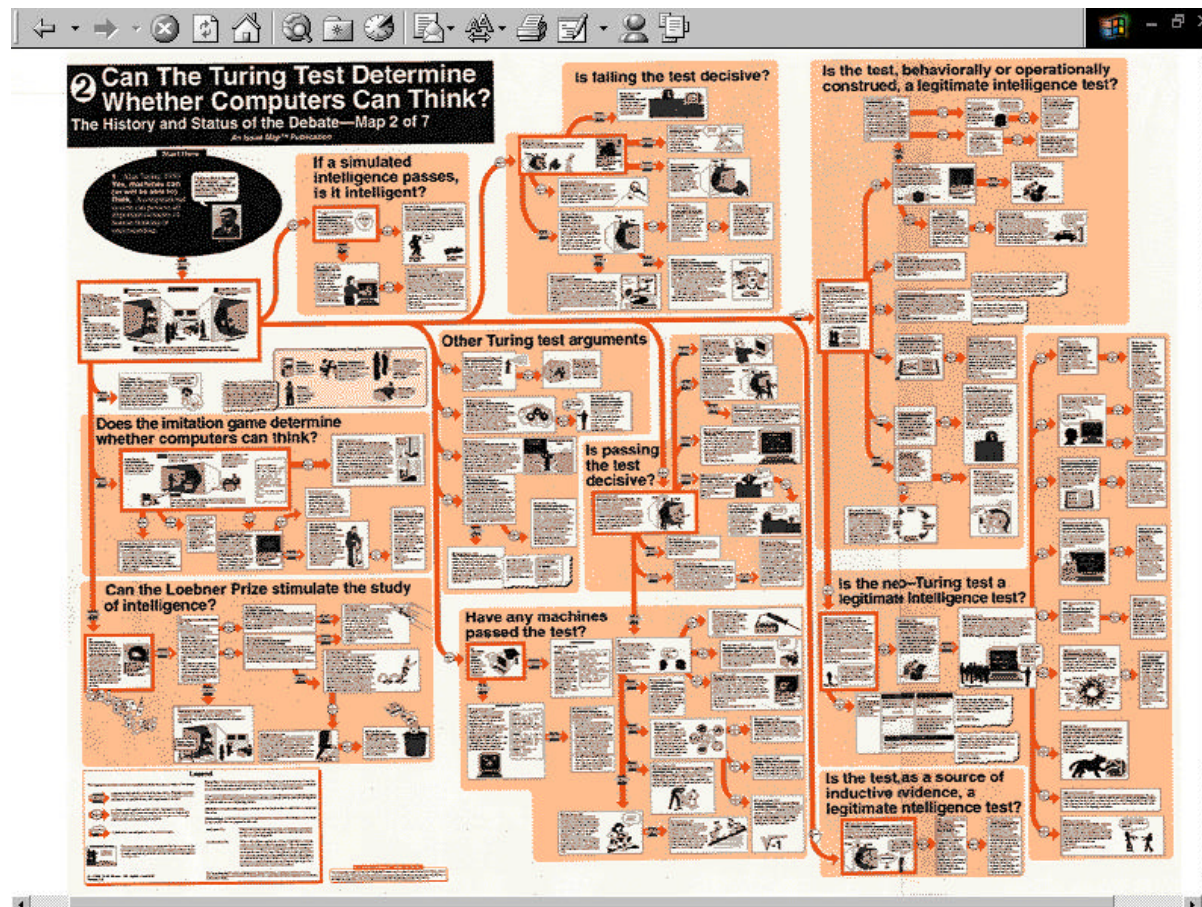
Fantasztikus kidolgozottságú tablótak alkottak a gépi intelligencia melletti pro- és kontra érvekről a **MacroVU Inc.** kutatói.



† Vélemények az MI mellett és ellen ..



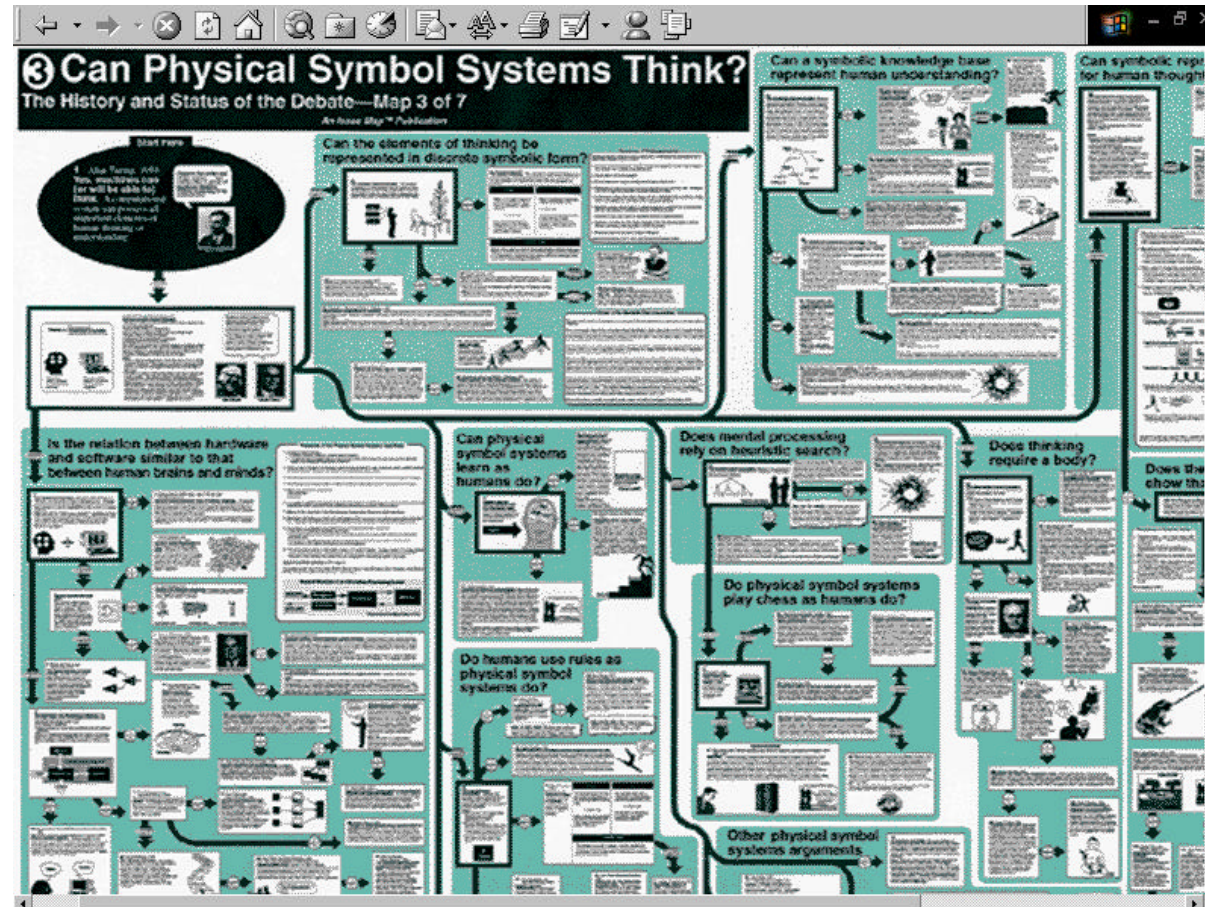
A tablók interaktívak és Internetrol elérhetok.
Elégséges-e a Turing teszt a számítógépek gondolkodási képességének meghatározására?



† Vélemények az MI mellett és ellen ..



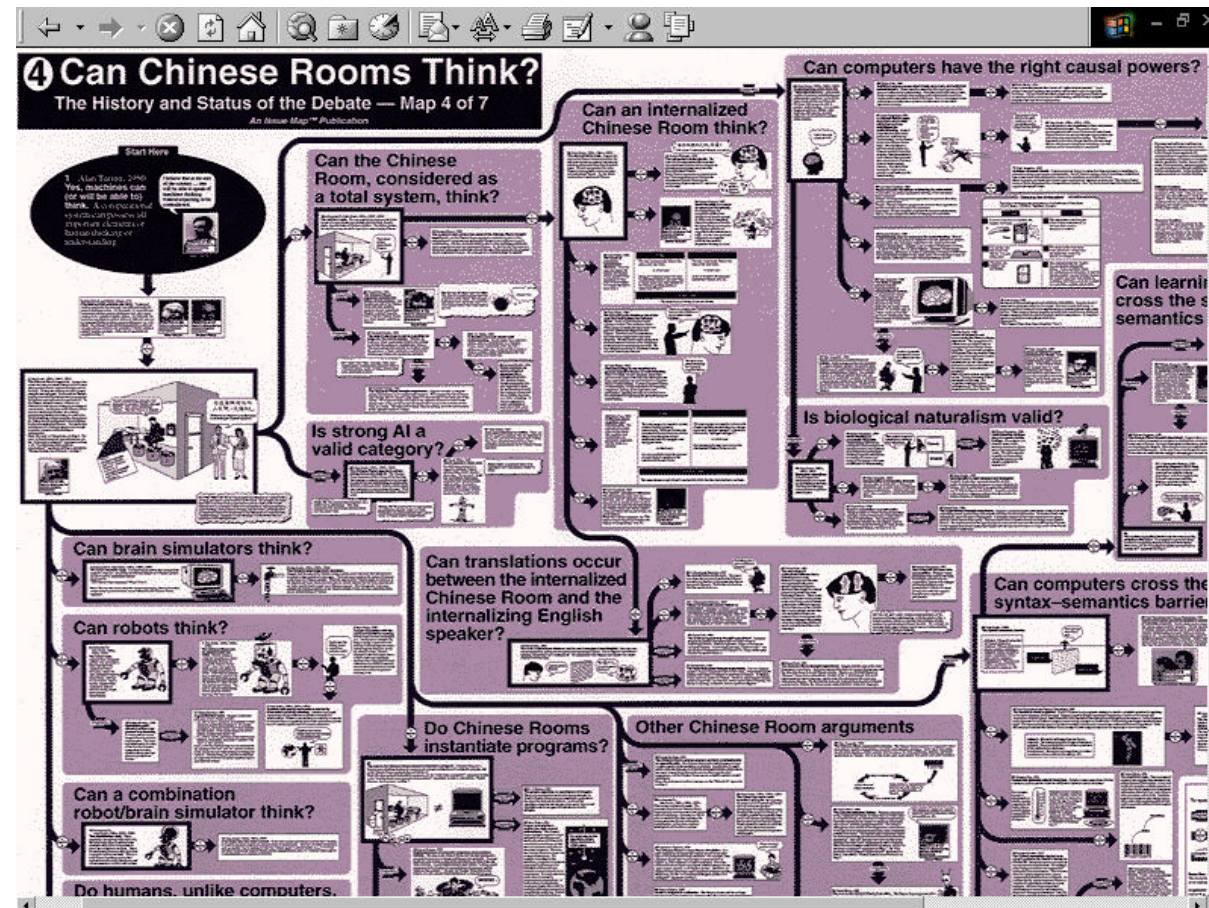
Tudnak-e a fizikai rendszerek gondolkodni?



† Vélemények az MI mellett és ellen ..



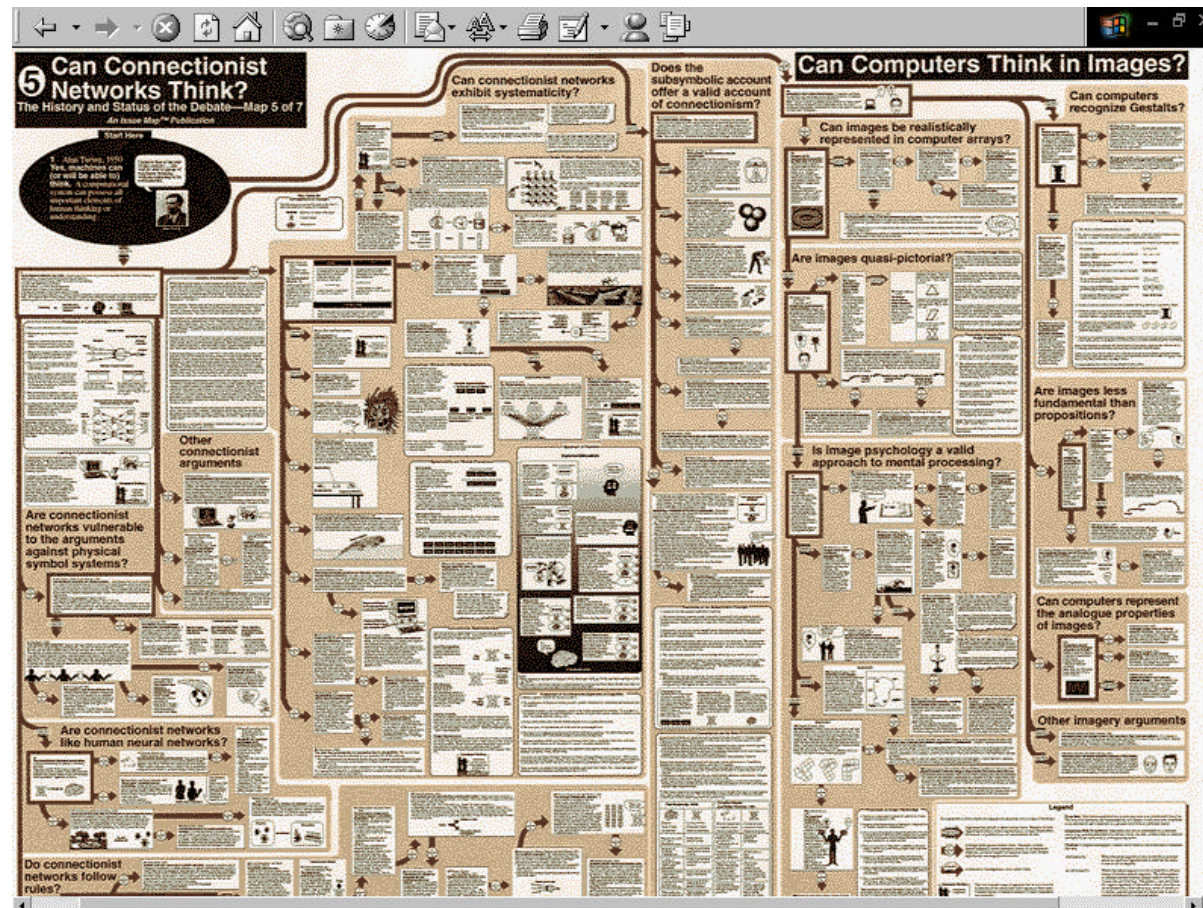
Gondolkodik-e a Kínai szoba?



† Vélemények az MI mellett és ellen ..



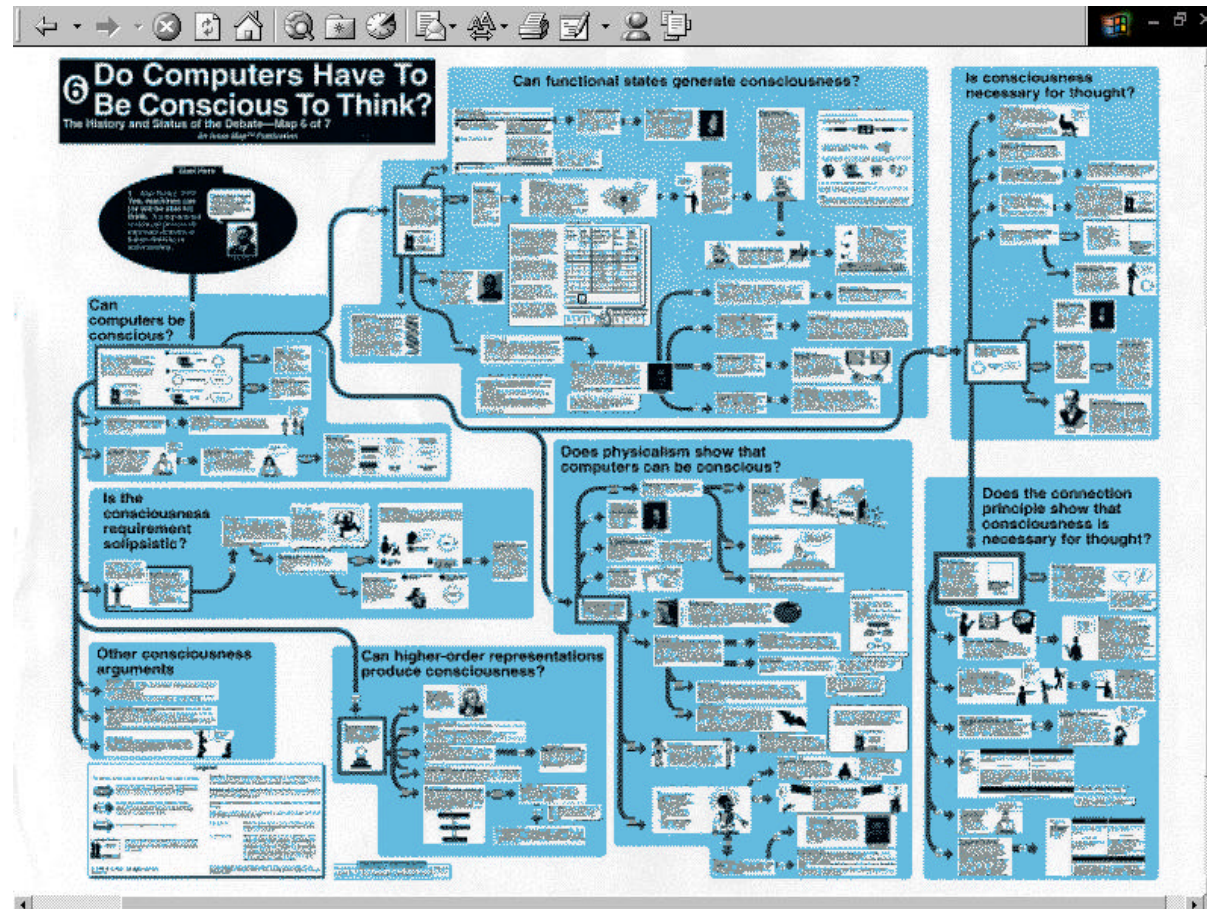
Gondolkodnak-e a neurális hálózatok?



† Vélemények az MI mellett és ellen ..



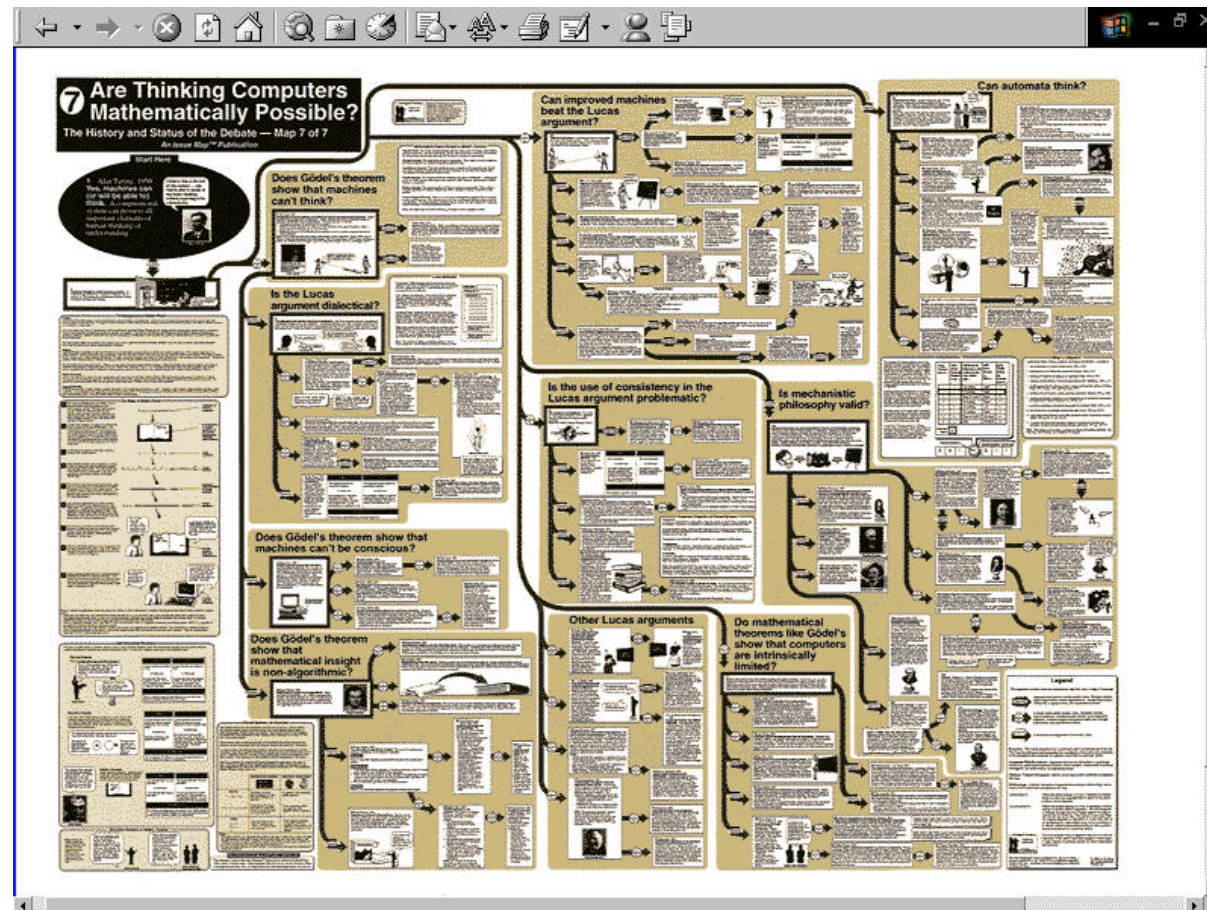
Tudatosan gondolkodnak-e a számítógépek?



† Vélemények az MI mellett és ellen ..



Egyáltalán, lehetséges-e matematikailag a gondolkodó gépek?



MacroVU®, Inc. The Power of Visual Language at Work

321 High School Road NE--Box 366, Bainbridge Island, WA 98110

Phone (206) 780-9612 | Fax (206) 842-0296

info@macrovu.com | www.macrovu.com

† Vélemények az MI mellett és ellen ..



- Eredeti gondolatok is találhatók a gépi intellegenciával kapcsolatosan **Kosara Thoughts on Artificial Intelligence, Gondolatok a mesterséges intelligenciáról** című publikációjában. Hangsúlyozza az emberi test, az evolúciós fejlődés, a folytonos tanulás, a gondolkodás társadalmiságának fontosságát az emberi intelligencia esetében. Legfontosabb felvetése azonban az **általános intelligencia kategória**.

