

Rédey Gábor – Neumann Attila – Sütő Zoltán

## Információkeresés

**A cikk óvatos kezdeményezés az információkeresés nyelvének, és ezen keresztül az egész információkeresési folyamatnak az „átértelmezésére”. Bár a gondolkodás jórészt nyelvfüggetlen, ezzel szemben az információcsere, és ennek következtében az információkeresés folyamata is erősen nyelvhez kötött, nyelvfüggő. A természetes nyelvet ma még nem szokványos önmagában információkereső nyelvként felhasználni. Annak ellenére, hogy az ötlet ígéretes, meglehetősen sok és nehéz problémát vet fel. A szakirodalomban kb. az 1990-es évek elejétől olvashatunk ilyen célú kutatásokról és eredményekről. A hagyományos információkereső nyelveknek a természetes nyelvekhez képest szegényes a szintaktikai szerkezetük. Az ezeken a nyelveken feltett, olykor „homályos” kérdésre sokszor „zajos” (irreleváns) vagy nagy információvesztésű válasz érkezik. Ebben a helyzetben előrelépést csak egy rugalmasabb szintaxisú információkereső nyelvtől várhatunk, amely nemcsak az egyedi fogalmakat (vagy azok valamilyen együttesét), hanem azok természetes nyelvi relációit is képes modellezni. A cikk célja, hogy rövid áttekintést nyújtson az új típusú, a természetes nyelvek logikai finomstruktúráját hűen leképező ismeretreprezentációs nyelvekről, elemzi helyük és alkalmazásuk lehetőségeit az információkeresés területén.**

Az információkeresés színhelye hagyományosan a könyvtár, művelője a könyvtáros, tudománya a könyvtártudomány. A könyvtári információkeresés módszerei még ma is erősen kötődnek az információtárolás hagyományos technológiájához. A számítástechnika megjelenésével azonban ez a kizárólagosság fellazult. A számítástechnika tömegessé válása előtt a fejlődés még szerves volt (Kunszt: LOGEL rendszere [1]), később az új technológia rohamos elterjedése nyomán a háttérben egyebek mellett újraértelmeződött az információtárolás és -keresés fogalma is. Ez a helyzet mára némileg megváltozott (szemantikus web), a hagyományos és az újabb szemléletmód között bizonyos közeledés tapasztalható.

### Mi az információ?

Az információ olyan alapfogalom, amely több nézőpontból is vizsgálható. Az egyik nézőpontot az információelmélet képviseli, amely a statisztikai valószínűség elvein alapul, és a kibernetika egyik ágának számít. Az információelmélet tárgya ebben az értelemben az információ mennyiségi vonatkozása, amely jelenlegi szempontunkból nézve kevésbé érdekes. Ehelyett érdeklődésünk középpontjában az információ minőségi oldala, vagyis a

szemantikai információ áll, amit a következőképpen definiálnak az irodalomban:

- „... az információ valamely szövegnek olyan struktúrája, amely alkalmas arra, hogy változást idézzon elő a befogadó képstruktúrájában” [2].
- „Információ fogalmán a viselkedést befolyásoló, új ismeretet nyújtó adatok tartalmi jelentését értjük. Az adatok és hírek csupán információhordozók. Az információ határozatlanságmennyiség-megváltozást okoz, jelentése az ismeretszint különbség” [3].

Az információ elvi meghatározásán túl nem vonatkoztathatunk el gyakorlati megközelítéstől sem, vagyis attól, ahogyan az információra a mindennapokban sokszor nem tudatosan gondolunk:

- „Az információ megjelenési formája általában természetes nyelvű szöveg, amelyből csak megfelelő szövegértelmező, illetve -feldolgozó képességgel tud a tanuló ismereteket szerezni” [4].
- „... információnak nevezünk mindent, amit a rendelkezésünkre álló adatokból nyerünk. Az információ olyan tény, amelynek megismerésekor olyan tudásra teszünk szert, ami addig nem volt a birtokunkban. Az információ legkisebb egysége a bit. A számítástechnikában a programok is 1 bites információkból épülnek fel” [5].

Könnyű észrevenni, hogy az információ elvi meghatározásában nincs nagy különbség a különböző szakterületek között. A különbség inkább az információfogalom gyakorlati megközelítésekor bukkan felszínre, ami leginkább az információ reprezentálásának mikéntjében ölt testet.

## Információreprezentáció

A szemantikus információ elvi megközelítése az általános emberi információreprezentációhoz kötődik, a szöveg (adat)struktúrájában látja annak megjelenését. A szöveg a maga tömegével azonban ebből a szempontból tekintve hatalmas, struktúrálatlan halmaz. Ehhez az adattömeghez hagyományosan két különböző módon lehet viszonyulni. A két eddig említett megközelítési mód – amit az egyszerűség kedvéért „könyvtári” és „számítástechnikai” megközelítésnek nevezünk – az információfogalom gyakorlati értelmezésében, a reprezentáció módjában tér el alapvetően egymástól.

## Hagyományos információkereső nyelvek

A szöveg közvetlen információtartalmának vizsgálatától távolodik el a könyvtári információfogalom azzal, hogy az információ reprezentálására szabványosított, strukturált információkereső nyelvet alkalmaz. A könyvtári információkereső nyelv olyan mesterséges vagy természetes nyelven alapuló nyelv, melynek szavai vagy nem a természetes nyelv szavai, vagy természetes nyelven kifejezett szavak ugyan, de a szavakat szabályozott formában (pl. főnév, alanyeset, egyes szám, hátravetett értelmezős homonimák, kiiktatott szinonimák) használják, és az e szavak által megnevezett fogalmak bizonyos meghatározott relációk által részben rendezettek. Világos, hogy az információkereső nyelvek kifejező ereje a természetes nyelvekéhez viszonyítva jelentősen romlik, nem beszélve az egyéb mellékhatásokról, azonban célja nem is a szövegek finom információtartalmának, finomszerkezetének megjelenítése, hanem éppen a durva tartalomnak, struktúrájának a feltárása a globális tájékozódás segítése céljából.

Ugyanilyen eltávolodás figyelhető meg a számítástechnikai információfogalom esetében, azonban egészen más okból és más eredménnyel. A számítástechnikai információ reprezentációja formális, „bitközpontú”, hiszen célja is csak bizonyos jelso-rozatok előfordulásainak megtalálása. A könyvtári információkereső nyelvekkel ellentétben a számí-

tástechnika számára egy szöveg pusztán egy nyelv (összefüggés nélküli) szavainak összessége. A reprezentációs veszteség mibenléte itt is azonnal szembetűnik.

Összegezve: mindkét eddig tárgyalt megközelítés a lehetőségek talajáról kiindulva, jelentős veszteséggel reprezentálja a szemantikus információt, ami eleve meghatározza a keresés minőségét és eredményességét. Ezt a képet némileg árnyalják elsősorban a könyvtári információkereső nyelvek továbbfejlesztési törekvései. Kunszt már említett tanulmánya [1] a jellemzően kétargumentumú generikus, partitív stb. ontológiai relációkkal strukturált keresőnyelvet megkísérli kiegészíteni a többargumentumú grammatikai relációkkal is, amely így elvileg képes lenne nyelvtanilag összetett keresőkifejezések képzésére is, ezáltal jobban megközelítve a természetes nyelvek kifejezőképességét. Talán érdemes itt kiemelni, hogy Kunszt reprezentációs módszere nagy hasonlóságot mutat a közel ugyanebben az időben Sowa által publikált fogalmi gráfok (conceptual graphs) [6] módszerével; kezdeményezése azonban egyelőre visszhangtalan maradt.

Kíváncsún lenné tehát, hogy az információkereső nyelvek is képesek legyenek a szöveg belső, szintaktikai összefüggéseinek a kifejezésére. A nem gépi információkeresés céljaira előállított eszközökben (pl. különféle speciális mutatókban) voltak és vannak erre szolgáló eszközök, de a gyakorlatban alkalmazott információkereső rendszerekben – legyenek azok akár hagyományos katalógusok vagy mutatók, akár online számítógépes információkereső rendszerek – ilyenek használata csak igen ritkán, kivételesen fordul elő [7].

Az idők folyamán azonban a számítástechnikai megközelítés információfogalma sem maradt változatlan. Az utóbbi években nagymértékű közeledést tapasztalhatunk a könyvtári információfogalomhoz. Itt különösen arra a változásra gondolunk, amely a mesterségesintelligencia-kutatások nyomán, az ontológiák megjelenésével a *teljes szöveges kereséstől a szemantikus web* fogalmáig vezetett.

Ugyanakkor e két módszer a gyakorlatban meglehetősen el is különül egymástól, kialakult alkalmazási területeik inkább kiegészítik, mint átfedik egymást. Ez természetes módon veti fel azt a problémát, hogy az információkeresés mégiscsak egységes szemléletű, nem függhet attól, hogy éppen mit, miben, milyen céllal keresünk. A következőkben ezt az eredeti célt tarjuk szem előtt.

## Tudásreprezentációs nyelvek

Az eddigiek alapján felvetődik a kérdés: vajon létezik-e olyan gyakorlati információfogalom, amely az előzőeknél jobban megközelíti az információ elvi értelmezését? Abból indulhatunk ki, hogy az információ reprezentálására a természetes nyelvénél alkalmasabb eszköz nem létezik. Ez indokolja, hogy a természetes nyelveket modellező mesterségesintelligencia-rendszereket tekintsük az információt leghívebben reprezentáló nyelveknek, amelyek képesek az információ legmélyebb szemantikai összefüggéseinek tükrözésére.

A természetes nyelvek szemantikai információtartalmának reprezentációja régi keletű törekvés, egyben a logika tárgya. A modern szimbolikus logika kezdetét a XIX. század végétől számítják. Ez nem jelenti azt, hogy az ókori vagy a középkori logika eredményei mellőzhetők lennének, éppen ellenkezőleg, valójában messzemenően azokra az eredményekre is támaszkodhatunk. Mindenesetre azzal az igénnyel, hogy a logikai következtetések az aritmetika módjára kiszámíthatók legyenek, először *Leibniz* lépett föl, célját azonban – legalábbis részben, a matematika nyelvére korlátozva – csak *Frege* érte el két századdal később. Mindezekkel arra utalunk, hogy a logikai ismeretreprezentáció célkitűzései és eredményei felelnek meg leginkább a szemantikai információ olyan igényű reprezentálásának, ami lehetővé teszi, hogy adott esetben egy szöveges információbázis számára feltett információkereső kérdés egyáltalán kiértékelhető legyen.

A logikai ismeretreprezentáció a logika nyelvén valósul meg. Ez a nyelv ma sokak számára a szimbolikus logikának a XIX. és XX. század fordulóján kialakult nyelvét jelenti, amelyet *Boole*, *Frege*, *Russell*, *Peirce*, *Peano* és mások az aritmetika nyelvénél mintájára alkottak meg. A természetes nyelvek és az aritmetika nyelve azonban bonyolultságukban nagyon is eltérnek egymástól. A hagyományos logika nyelve – bár voltak erre kísérletek – nem alkalmas a természetes nyelvek logikai szerkezetének modellezésére. Nem azért, mert a feladat nem volna így megoldható, hanem mert az eredmény gyakorlatilag nem használható. Lásunk ennek szemléltetésére egy példát *Ruzsa Imre* könyvéből [8]:

*Egyetlen fiú sem csak Marit szerette.*

$P \rightarrow \{ \sim \exists x (fiú\ x) \ \& \ [(\lambda y. szeret\ x\ y) \equiv \lambda y(y = Mar)] \}$

Az illusztráció azt mutatja, hogy a formula előállítása és visszaolvasása egyaránt nehézséget okoz, aminek az az oka, hogy a leírt formula a magyar nyelvű mondat szemantikai információtartalmát ugyan pontosan tükrözi, szintaktikai szerkezetét azonban nem. A logikai szintaxis kissé szegényes a természetes nyelvek szintaxisához képest. Így az algoritmus, amelyet a magyar nyelv egy töredékének formalizálására *Ruzsa* javasol, amelynek segítségével tehát egy természetes nyelvű mondatból a hozzá tartozó logikai formula előállítható, kilátástalanul bonyolult. Ez indokolja egy olyan logikai nyelv szükségességét, amely nemcsak a természetes nyelvű mondatok szemantikai tartalmának hű leképezésére képes (mint ahogyan ezt a hagyományos logika nyelve teszi), hanem a nyelv szintaktikai viszonyainak hű leképezésére is. Ekkor ugyanis elvárható, hogy – alkalmas természetes nyelvi elemző közbeiktatásával – a természetes nyelvű mondat szintaktikai egységei könnyen átfordíthatók legyenek a logikai nyelv szintaktikai egységeire. Vagyis a *természetes nyelvű szöveg*  $\rightarrow$  *reprezentált szöveg* közötti fordítás – a számítógépes nyelvészet meglévő eredményeit felhasználva – gépesíthető.

A vázolt problémára az irodalomban több megoldás is létezik. Anélkül, hogy részletekbe bocsátkoznánk, csak egy-egy példát villantunk fel az egyes módszerek legszembetűnőbb sajátosságainak illusztrálására. A részletek iránt érdeklődők számára a meglehetősen gazdag irodalomra utalunk. Sowa már említett fogalmi gráfok (*conceptual graphs* = CG) néven ismert reprezentációs nyelvet [6] *Peirce* egzisztenciális gráfnyelvéből vezeti le:

*All trailer trucks are eighteen wheelers.*

[trailerTruck :  $\forall$ ]  $\rightarrow$  (part)  $\rightarrow$  [wheel : { $\ast$ }@18]

*Iwańska UNO-nyelve* [9] (a betűszó az *Unification* és a *NegO* szavakból származik) már kifejezetten a nyelvtani szerkezetre épít, bizonyos alaprelációkkal kiegészítve:

*Every student works hard.*

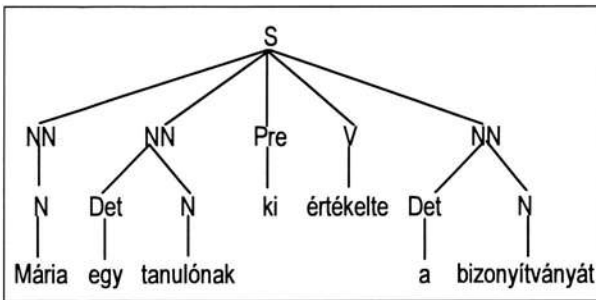
$np(det \Rightarrow every, n \Rightarrow student) == [work(adv \Rightarrow hard)]$

*Rédey* intenzionális szövegrepresentációs nyelve (*iCTRL* = *Intensional Conformal Text Representation*) [10] szintén a mondat nyelvtani relációit tükrözi, abból az alapfeltételezésből kiindulva, hogy a nyelvtani szerkezet a logikai szerkezetet teljes egészében magában foglalja:

Mária egy tanulóknak kiértékelte a bizonyítványát.

(((((értékelte x y)z w,  
a y, bizonyítványát y),  
<ki>w,)  
egy z, tanulóknak z),  
Mária x.

Azt, hogy ez utóbbi esetben a szintaxis a természetes nyelvek alapvető nyelvtani relációira (állítvány, alany, tárgy, jelző, határozók) épül, példamondatunk nyelvtani elemzése szemlélteti (1. ábra).



1. ábra A „Mária egy tanulóknak kiértékelte a bizonyítványát” mondat nyelvtani elemzése a MorphoLogic Kft. Moose számítógépes nyelvtani elemző rendszerével

A nyelvtani és logikai szerkezet ilyen szoros kapcsolata garantálja, hogy a természetes nyelvű szöveg ↔ reprezentált szöveg közötti fordítás gépi úton valóban könnyen végrehajtható. A gépi reprezentációra fordítás lehetősége olyan mozzanat, amelynek hiánya értelmetlenné tenné a szóban forgó reprezentációs nyelv minden más esetleges előnyét. Emellett ugyanebből – tehát hogy a reprezentáció mind szintaktikailag, mind szemantikailag, minden részletében követi a természetes nyelv szerkezetét – következik, hogy a reprezentált szöveg minden részlete a keresés számára elvileg hozzáférhető.

A következőkben vázoljuk az információkeresés elvét, továbbá a szöveghű ismeretreprezentációs nyelvekre alapozható információkereső rendszerek architektúráját.

## Mi az információkeresés?

Információkeresésen általában azt értik, amikor valamilyen formalizált információt hasonlítanak egy már rendelkezésre álló, formalizált információhalmaz elemeihez. Ennek háttérben az áll, hogy a keresést mindig valamilyen tudáshiány váltja ki,

ami vagy valamilyen feltételezés (hipotézis) formáját ölti, amelynek ismeretlen igazságértékét verifikálni kell, vagy valamilyen, bizonyos konkrét tulajdonságokkal rendelkező ismeretlen létezésének a feltételezését jelenti, amit a rendelkezésre álló adatok alapján szintén igazolni kell. A keresés mindig valamilyen előzetes, többnyire nyilvánvalónak gondolt (ezért általában nem kifejezett) ismeretre épül. A keresés után a talált információ – optimális esetben – növeli a kereső már meglévő ismeretszintjét.

Az előzők alapján nem meglepő tehát, hogy a kereső nyelv sajátosságai meghatározzák a keresés eredményének várható minőségét is. Ha csak karaktersorozat tudunk keresni egy másik karaktersorozatban, akkor annál többet nem várhatunk, mint hogy meg is találjuk. Ha a kereső nyelvünk szavai részben rendezettek bizonyos relációkra nézve, akkor jogosan feltételezhetjük, hogy ez a keresés eredményében is tükröződik.

A keresés minőségét alapvetően befolyásolja az a háttérismeret, amire támaszkodni lehet. A pusztán karaktersorozat-keresés esetében semmilyen háttérismeret nem tudunk felhasználni, ellenben a könyvtári információkeresés vagy a szemantikus web masszív háttérismeretre támaszkodik. Ez a háttérismeret azonban általános, statikusan rögzített, és csak lassan bővül. Mindezeket túl az az információtömeg, ami a keresés bázisát jelenti, nem, vagy csak viszonylag szűk hányadban vesz részt a keresésben. Mivel a keresés mindig csak a reprezentációs (információkereső) nyelven hajtható végre, ez más megvilágításban azt jelenti, hogy a hagyományos információkeresés számára a szöveg jelentős része „elérhetetlen” marad, pontosan annyi információ érhető el a kereséskor, amennyit a reprezentációs nyelv „felbontóképesége” megenged. A reprezentációs nyelv tehát eleve meghatározza a keresés minőségét, ami magától értetődően támasztja alá az információkereső/-reprezentációs nyelv célszerű megválasztásának alapvető jelentőségét, hiszen „... az információ annyit ér, amennyi megtalálható belőle” [11].

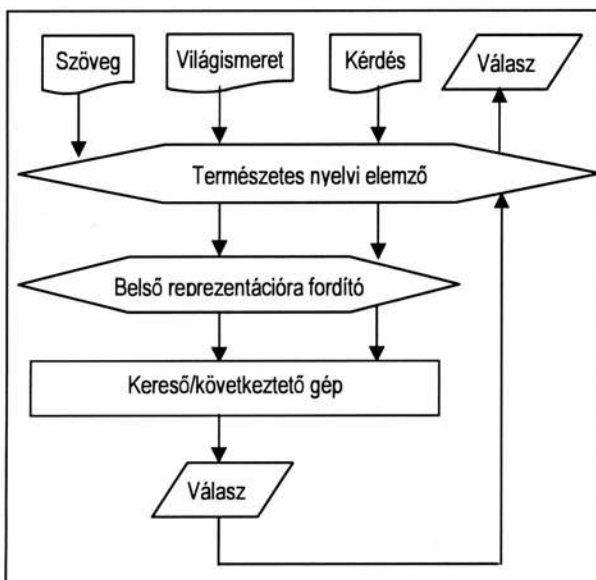
## Egy lehetséges kérdés-válasz rendszer

A következőkben egy olyan rendszert vázolunk, amely minőségi előrelépést jelent az információkeresésben. Az alkalmazott ismeretreprezentáció módszere tekintetében nincs korlátozás, elvileg bármely, szövegek szemantikai információtartalmának reprezentálására alkalmas módszer alkalmazható. Ilyenek pl. a már említett CG, illetve az

UNO reprezentációs módszerek, továbbá az iCTRL, amelyet munkánkban alkalmazunk.

Egy ilyen rendszerrel szemben a következő követelményeket állítjuk:

- A rendszer képes megtalálni bármely természetes nyelven megfogalmazott terminust, és megmutatja, hogy azt mely szövegösszefüggésben találta meg.
- A terminus keresése közben lehetőség van asszociációkra (az eredetivel valamilyen relációban lévő terminusok keresésére).
- A rendszer számára természetes nyelvű kérdéseket lehet megfogalmazni, és azokra ugyanazon a nyelven válasz érkezik.
- A rendszer megmagyarázza, hogy az adott kérdésre adott válaszhoz milyen közbenső lépéseken jutott el.
- A rendszer a kérdések megfogalmazásához segítséget nyújt: az ember számára érthető formában mutatja meg az általa használt fogalmakat és a közöttük lévő összefüggéseket.



2. ábra A kérdés–válasz rendszerek egy működési sémája

Mi egy ilyen rendszer lényege? Működtet egy „értelmező motort”, amely képes értelmezni egy természetes nyelven megfogalmazott állítást vagy kérdést, sőt képes értelmezni természetes nyelven tárolt szöveget is, továbbá képes létrehozni a kérdés, valamint a szöveg között a kívánt relációk szerinti megfeleltetést, vagyis szövegben szemantikus információt keresni. Lényegében ez azt jelenti, hogy egy gép képes nagyobb információhalmazt is átolvasni az ember helyett, és képes abból ki-

emelni pl. a kívánt relációknak megfelelő szövegrészeket. A teljes folyamat vázlatát a 2. ábra szemlélteti.

Az információkeresésnek ez az elképzelhető legkényelmesebb módja. Az ember röviden elbeszélget egy géppel, és eredményként megkapja egy nagy szöveghalmazból a számára fontos szövegrészeket, a feltett kérdéseire adott válaszokat.

Mindez, bár meglehetősen futurisztikusan hangzik, megoldható az előbbieken vázolt nyelvi elemző és reprezentációs módszerek alkalmazásával, amelyek jól követik azt az absztrakciós folyamatot, amelyet az ember a természetes nyelvekben használ.

### Az információkereső rendszerek hatékonysága

A mesterségesintelligencia-módszerek már vázolt gyakorlati alkalmazása további két figyelemre méltó szempontot vet fel: a nyelvi analízist és a gépi reprezentációra fordítást végző algoritmus elképzelhető sebességének kérdését, és az ebből következő gazdaságossági kérdéseket.

Tekintettel arra, hogy bármilyen hatékonynak is képzelünk egy nyelvi analízist és gépi reprezentációra fordítást végző algoritmust (aminek hatékonysága egyébként nyilván fokozható előfeldolgozási, szűrési eljárásokkal), a nyelvi analízis és a fordítás csak komplex logikai műveletként gondolható el, elemzési, illesztési és összehasonlítási műveletek halmazával, amelyek időszükséglete az algoritmus fokozatos csiszolásával ugyan nyilván folyamatosan egy minimum felé szorítható, ez a minimum azonban a jelenlegi és a jövőbeni hardverképeségek mellett is mindenképpen jelentős érték marad. Eddigi tapasztalataink szerint kielégítő teljesítmény lenne, ha egy átlagos összetett mondat kiértékelését az algoritmus 0,001 szekundum körül el tudná végezni. Ez azt jelenti, hogy egy átlagos könyv (300 oldal) „átolvasása” az algoritmusnak 2,5–3 másodpercet vesz igénybe. Ez az eredmény már mindenképpen használhatónak mondható, mert pl. mentesítheti az embert attól, hogy fölöslegesen elolvassa a számára irreleváns irodalmat. Ugyanakkor ilyen módszerrel nekilátni egy könyvtáryi anyag feldolgozásához egyetlen kérdés miatt, egyszerűen kilátástalan. (Az Országos Széchényi Könyvtárban kb. 4,5 millió információs egységet tárolnak, amelynek a túlnyomó többsége könyv, így ez a munka mintegy 140 na-

pig tartana.) Az ilyen nemzeti könyvtárban található információmennyiséget nagyságrendekkel meghaladó információtáraknak (mint amilyen az internet) hasonló módszerekkel nekiesni még akkor is értelmetlen, ha a jövő ígéretébe, a kvantumszámítógépek világába képzeljük magunkat, akár több nagyságrenddel megnövelt számítási kapacitással.

Hasonló gondolatmenettel feltételezhetjük, hogy egy már feldogozott szöveges állomány esetében tetszőleges kérdés kiértékelése átlagosan legalább ugyanennyi, vagy akár nagyságrendekkel több időbe kerül. Ennyiből talán nyilvánvaló, mennyire főleg az ábrándot kerget az, aki nagy ismeretbázisok online faggatását tűzi ki célul. Ez a felismerés valamiképpen a formalizált, szisztematikus kérdések rendszerében rejlő lehetőségek felértékelődéséhez vezet.

Némi megfontolás után kiderül, hogy az említett formalizált kérdések halmaza lényegében azonosnak tekinthető a szóban forgó háttértudás egy részével: az ontológiák, teauruszok, osztályozási rendszerek által tárolt és rendszerezett fogalmakkal. Vagyis, hatékonyabbá tehető a keresés, ha első lépésben veszünk egy jól rendszerezett fogalomtárat, és az algoritmusunkkal ismeretbázisunkat e fogalomtár szerint rendezzük. Ez a módszer első körben elveszti a tetszőleges kérdés feltételének közvetlenségét, de a keresést, a felhasznált fogalomtár hierarchikus rendezettségét kihasználva, hatékonyabban és gyorsabban hajtja végre.

Az ismerethalmaz  $n$  elemű fogalomtár szerinti rendezése elvben azt jelenti, hogy az  $n$  kérdést a teljes ismerethalmazon végigfuttatva előáll az az  $m$  ( $< n$ ) elemű szignifikáns fogalomtár, amely a szóban forgó ismerethalmazt pontosan jellemzi. Ez annak ismeretében válik fontossá, hogy pl. az ETO középkiadása mintegy 80 000 nyelvi egységet tartalmaz, ami az implicit információk figyelmen kívül hagyása esetén a teljes ismeretbázison 80 000 kérdés végigfuttatását feltételeznél, a fenti alapadatokat figyelembe véve mintegy 30 000 év időszükséglettel.

Úgy hisszük, hogy ez a gondolatmenet kellően alátámasztja a jelenlegi könyvtári keresőrendszerek alapvető szerepét. Eszerint a meglévő vagy azokhoz hasonló rendszerek nélkülözhetetlenek az információkeresés első, az adekvát forrás meghatározásának fázisában, amit a jövőben egy kötetlenebb, párbeszéd jellegű finomkeresési fázis követhet. E második fázis feladata lesz integrálni a meglévő navigációs lehetőségeket, és az újonnan

rögzítendő indexelési adatokat, mint keresésgyorsító és -pontosító eszközöket a belső szerkezet finomstruktúrájával, amely a szűkített adathalmazon való értelemszerinti kereséssel ténylegesen megvalósítja az intelligens kérdés–felelet funkció követelményeit.

## **Összegzés, jövő, feladatok**

A fentiekben vázoltuk a könyvtári információkeresés elvi és gyakorlati háttérét, különös tekintettel a számítástechnikai eszközök felhasználására. A létező keresőmódszerek két, egymástól lényegesen különböző szempontot megvalósító családba sorolhatók, úgymint a hagyományos könyvtártudomány vonalát követők, illetve a gépi számítástechnikai szemléletűek. Előbbi reprezentálja a hosszú idő alatt felgyűlt specifikus könyvtártudományi ismereteket, és tartalmi keresésnek is nevezhetjük, utóbbi alkalmazza a számítógépek nyers technikai és algoritmikus képességeit, és mint ilyet, formai keresésnek is tekinthetjük. Sajnálatos, hogy e két megközelítés sokáig távol állt egymástól, konfliktusossá téve a kapcsolatot a két szakterület között. A probléma az 1990-es évektől kezdve tudatosult e határterület művelői között, és számos megoldási kísérlet született a tartalmi szempontokat adekvát módon kiszolgáló számítástechnikai eszközök megalkotására. E módszerek a mai napig nem érték el azt a szintet, amely egy gördülékeny keresőeljárás megalkotásához nélkülözhetetlennek látszik. A jelen megoldások fő hiányossága a szükséges jelentős manuális előkészítő munka, a természetes nyelvű szövegek és a gépi ábrázolásuk közti lényeges különbség miatt.

Vázoltunk egy alternatív lehetőséget, amely a természetes nyelvű szövegek automatikus tartalmi/logikai leképezését képes megvalósítani, illetve ilyeneken keresést, következtetést végezni. Lehetőség nyílik a létező tudásanyagok integrálására, a rajtuk történő navigálással együtt. Az élő gyakorlatra tekintettel elemeztük a tartalmi keresőrendszerek hatékonyságproblémáját is, ami a különböző rendszerek egymást kiegészítő, párhuzamos alkalmazásának fontosságára mutat rá.

A hivatkozott szövegrepresentációs módszer jelenleg fejlesztési fázisban van. Implementációja a feladatok széles köréhez adhat alkalmas eszközt. Közvetlen célként a létező könyvtári keresőrendszerekbe való automatikus szövegbesorolás képességét céloztuk meg. Később specifikus ismeretanyagra vonatkozó szakértői rendszer kiépítését

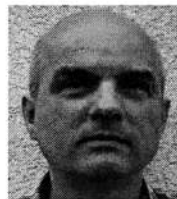
tervezzük. Távlati célként többnyelvű tudásháttérre alapozott, rugalmas ember-gép, kérdezz-felelek kommunikáció megvalósítását tervezzük. Alkalmazási területként elsősorban olyan tudományágak kerülhetnek szóba, amelyek világos, egyértelmű, rögzített fogalomrendszerrel fejtik ki tárgyukat (pl. a jogi, orvosi, mérnöki tudományok).

## Irodalom

- [1] KUNSZT György: A tudományos kutatás logikai modellezése és tematikai irányítása. Budapest, Akadémiai Kiadó, 1975.
- [2] FÜLÖP Géza: Az információ. Bukarest, Kriterion, 1990.
- [3] BÁLYA Dávid: Az informatika kihívása a teszt-technológiában. [Budapest], BME TIO, 1997.
- [4] DÁN Krisztina-HARALYI Ervinné: Könyvtárhasználati ismeretek a kerettantervben. <http://www.om.hu/letolt/kozokt/konyvtar.doc>
- [5] VINCZE Tamás: Hálózati kislexikon. [http://gisfigyelo.geocentrum.hu/informatika/kisokos\\_informacio.html](http://gisfigyelo.geocentrum.hu/informatika/kisokos_informacio.html)
- [6] SOWA, John. F.: Knowledge Representation: Logical, Philosophical and Computational Foundations, Pacific Grove, CA, PWS Publ. Co., 1999.
- [7] UNGVÁRY Rudolf-VAJDA Erik: Könyvtári információkeresés. Budapest, Typotex, 2002.
- [8] RUZSA Imre: Logikai szintaxis és szemantika. 2. köt. Budapest, Akadémiai Kiadó, 1988.
- [9] IWAŃSKA, Lucja M.-SHAPIRO, Stuart C. eds.: Natural language processing and knowledge representation. Cambridge, MIT Press, 2000.

- [10] RÉDEY Gábor: iCTRL: Intensional conformal text representation language. = Artificial Intelligence, 109. köt. 1-2. sz. 1999. p. 33-70.
- [11] PROKKNÉ PALIK Mária: A tartalmi feltárás problémái online könyvtári katalógusokban. = Tudományos és Műszaki Tájékoztatás, 52. köt. 11-12. sz. 2005. p. 525-527.

Beérkezett: 2006. XI. 8-án.



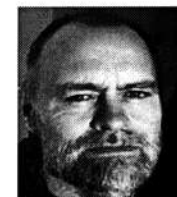
### Rédey Gábor

az Országos Atomenergia Hivatal vezető főtanácsosa.  
E-mail: [redegy@iif.hu](mailto:redegy@iif.hu)



### Neumann Attila

a Neumann Fivérek Kft. ügyvezetője.  
E-mail: [neumann.attila@chello.hu](mailto:neumann.attila@chello.hu)



### Sütő Zoltán

a TotalZoom technológia kifejlesztője.  
E-mail: [suto.zoltan@gmail.com](mailto:suto.zoltan@gmail.com)

## Nem lesz EU-szintű a szerzői jogdíj szabályozása

A korábbi elképzelésekkel ellentétben meglepő gyorsasággal visszavonta az egységes szerzői jog kialakítására vonatkozó javaslatát az Európai Bizottság. Az előterjesztés ellen korábban főleg Franciaország és különböző alkotói szervezetek tiltakoztak. Az Európai Bizottság tervei szerint egy minden mai igényt kielégítő, korszerű, egységes európai szerzői jogi szabályozást vezettek volna be a jelenleg hatályos, számos országban eltérő szerzői jogi törvények helyett. Jelenleg az Európai Unió egyes tagországaiban teljesen vegyes a kép, hogy mely készülékekre van szerzői jogdíj. Az unióban csupán három olyan ország van, ahol egyáltalán nincs szerzői jogdíj: Nagy-Britanniában, Írországon és Luxemburgban.

Hazánkban az *Artisjus Magyar Szerzői Jogvédő Iroda Egyesület* oldalán közölt felsorolás szerint kötelező jogdíjat fizetni az audio- és a videokazet-

ták, a hang- és a video-képhordozó nyersanyag import, az írható CD- és DVD-lemezek, illetve DVD-RAM-ok, az integrált tárolóegységgel rendelkező zenelejátszók a kép-, illetve hanghordozóként használható memóriakártyák és a minidiszkek után.

„A téma összetettsége miatt döntött úgy az Európai Bizottság, hogy egyelőre elveti a szerzői jogdíj egységes európai szabályozását” – nyilatkozta a döntéssel kapcsolatban Pia Ahrenkilde Hansen szóvivő.

„Számunkra ezzel a döntéssel bizonyossá vált, hogy az Európai Bizottság minden olyan elképzelést meghiúsított, amelyek a szerzői jogdíjak igazságosabb kiszabását, beszedését és elosztását lehetővé tette volna” – reagált a hírre Mark McGann, a CLRA szóvivője.

<http://www.sg.hu/cikkek/49225/>