

Bódog András

Webinárium a mesterséges intelligencia könyvtári alkalmazásáról

2020. szeptember 24-én az Egyesült Királyságbeli Időszaki kiadványok Csoportja (*United Kingdom Serials Group*) webináriumot rendezett *AI: Empowering Libraries & Making It Real (MI: a könyvtárak megerősítése & tegyük valósággá!)* címmel, a téma a mesterséges intelligencia (MI) könyvtári felhasználása volt.

Az első előadó, Ken Chad könyvtári tanácsadó volt, aki általános jelleggel ismertette a mesterséges intelligencia kérdéskörét. A science-fiction alkotásokban ábrázoltaktól eltérően a jelenlegi MI korántsem gondolkodó gépeket jelent, hanem a háttérben dolgozó, olykor láthatatlan algoritmusokat, amelyek felgyorsítják, és az emberi, manuális munkánál hatékonyabban végrehajtják az ismétlődő és olykor lassú, esetleg nehézkes folyamatokat. Napjainkban sokan nem bíznak a mesterséges intelligenciát alkalmazó rendszerekben, többnyire a tudományos-fantasztikus művekben és a populáris kultúrában ábrázolt hamis kép miatt. Chad szerint fontos kialakítani az MI etikai hátterét, amely magában foglalja mind fejlesztői, mind használói oldalról az etikus alkalmazást. A brit Guardian napilap publikált egy témába vágó vitacikket: egy egyetemi hallgató öngyilkossága után az édesapa szorgalmazta, hogy az egyetemek gyűjtsenek adatokat az egyetemi polgárokról annak érdekében, hogy kiszűrhessek a szorongó fiatalokat. Bár kétségtől pozitív szándék vezérli, ez adatvédelmi szempontból nem éppen etikus adatgyűjtés lenne, az esetleges rossz célú felhasználásról nem is beszélve. Mivel a könyvtárak még manapság is megbízható szervezetnek számítanak a társadalom szemében, az MI-rendszerek könyvtári alkalmazása követendő példaként szolgál. Ám ha az MI nem öntudatra ébredt szuperszámítógépeket vagy meghök-

kentően emberi androidokat jelent, akkor valójában mit is? Jelenleg a mesterséges intelligencia a gyakorlatban információmenedzsmentet jelent, amelyre vizualizáció, analitika, majd végül valamilyen gépi tanulás, esetleg neurális hálózat épül, eljutva a kezdeti, csak leíró adatoktól a következtető, prediktív adatokig. Gyakorlati példát jelent az MI alkalmazására a *big data*, az analitika és a természetes nyelvfeldolgozás.

Alapvetően háromféle MI-elméleti irányzat létezik: a fenti példának megfelelő, valamilyen meghatározott munkára készített, úgynevezett szűk mesterséges intelligencia (*artificial narrow intelligence* – ANI); az emberi intelligenciával vetekedő általános mesterséges intelligencia (*artificial general intelligence* – AGI); végül az embereknél is okosabb mesterséges szuperintelligencia (*artificial super intelligence* – ASI). Utóbbi kettő még a jövő zenéje, azonban az ANI-megoldások már a mindennapjaink részét képezik. Könyvtárak esetében a feladatok elemzésében és a munkafolyamatok javításában játszhatnak fontos szerepet a gépi tanulási algoritmusok. Számos felsőoktatási könyvtár gyakorlatában látunk példát szűk mesterséges intelligencia vezérelte chatbotra, az MI-t alkalmazó analitikai rendszerek pedig az oktatási adatok mellett a könyvtári használat mérőszámait is hatékonyabban kimutathatják.

A következő előadó, Manisha Bolina, a Yewno vállalat üzletfejlesztési alelnöke a *Yewno Discover* keresőszolgáltatást mutatta be a milánói Szent Szív Katolikus Egyetem (*Università Cattolica del Sacro Cuore*) példáján. Napjainkban olyan hihetetlen mennyiségű adat vesz körül bennünket, hogy mesterséges intelligencia támogatására van szükség az optimális információkeresés biztosítása érdekében. A *Yewno Discover* rendszere mély neurális hálózatok alkalmazásával tárja fel a hatalmas adatmennyiséget, és alkot szemantikus klasztereket, amelyek a további beérkező adatok függvényében dinamikus módon frissülnek. A Yewno által tudásgráfnak (*knowledge graph*) nevezett módszer képes a felszín alatti rejtett kapcsolatokat is felismerni. Míg a hagyományos adatbázisok nem értik a tartalom jelentését, a tudásgráf szemantikus jelentés szerint tudja kategorizálni és tématerképszerű gráfban vizuálisan megjeleníteni a fogalmakat. A hagyományos kulcsszavas keresés természetes nyelvfeldolgozással egészül ki, majd a mély neurális háló a kontextus és a jelentés értelmezésével tárja fel az egyetem állományában lévő írott tudásmennyiséget. Tulajdonképpen kulcsszavak helyett fogalmakat alkot, amelyekhez jelentés társul, növelve ezzel a keresés hatékonyságát. A vizuálisan megjelenő tudásgráf a releváns szakirodalom mellett a témához kapcsolódó egyéb témákat is feltünteti, amelyekhez a forrást az egyetem saját intézményi repozitóriuma és egyéb dokumentumai mellett – az üzleti feltételek tisztázását követően – az előfizetett adatbázisok tartalmait képezik. Röviden összefoglalva: a *Yewno* egy olyan (szűk) mesterséges intelligencia vezérelte keresőrendszer, amely egy hatalmas mennyiségű tudástárban képes keresni, összefüggéseket feltárni, és mindezt vizuálisan ábrázolni. Egy ilyen rendszert használó könyvtár olvasói fel-

használói fiókjukkal bejelentkezve testre szabhatják a kereséseket, akár egyéni tartalmakat, megjegyzéseket is hozzáadva a találatokhoz. A *Yenno Unearth* intelligens analitikus rendszer pedig a szolgáltató könyvtáraknak biztosít hatékonyságnövelő megoldásokat. (Magyarországon a Corvinus Egyetem alkalmazza a *Yenno Discover* keresőrendszerét.)

A webinarium utolsó előadója Ben McLeish volt az Altmetric tanácsadócégtől és egyben a *Dimensions* intelligens multidiszciplináris kereső fejlesztőcsapatából. Felvezetésében elárulta, hogy bármikor használjuk is a Google-t, egyidejűleg legalább négy MI-alkalmazásnak is munkát adunk a keresőkifejezések automatikus kitöltése, a releváns találatokat rangsoroló algoritmusok és a fogalmi keresés révén (utóbbi során a Google a keresőkifejezéshez hasonló vagy szinonim fogalmakat is felkínál). Nem utolsósorban, ha csak egy üres Google-oldalt töltünk be, már annak is köze van a mesterséges intelligenciához: a Google-adatközpontok szerverfarmjainak energiaellátását ugyanis a leoptimalisabb működés érdekében egy MI-vezérelt program felügyeli a nap 24 órájában.

A tudományos publikációk világában mintegy 161 millió dokumentum lelhető fel, magukban foglalva körülbelül 5 milliárd linket. Mindezek manuális kategorizálása túl lassú lenne akár egy 1500 fős szakmai stábnak is, ráadásul ekkora adatmennyiség feldolgozása egyben az új tartalmak figyelmen kívül hagyásával is járna, nem is említve a magas bérköltséget. További kihívást jelent, hogy a nagy szolgáltatók egyre inkább előnyben részesítik a tartalmak általános besorolás – például „multidiszciplináris” – szerinti rendszerezését, nehezítve a tárgy szerinti böngészést. A sajátos terminológia miatt további nehezen feldolgozható területet képviselnek a különféle tudásbázisokban a szabadalmi kategóriák, illetve a különféle pályázatok. A *Dimensions* gépi tanulás segítségével dolgozza fel a különböző interdiszciplináris forrásokat. Az aktív tanulás módszerét alkalmazó gépi tanulási modell egyaránt képes feltárni a szakemberek által feldolgozott címkézett és a feldolgozatlan, címkézetlen dokumentumokat. Az MI az emberi feldolgozó munka eredményeképpen létrehozott alapbeállítás révén alkot szemantikus kapcsolatokat a hatalmas mennyiségű tartalomban. Ezáltal lehetéssé válik a források automatikus, részletes tudományterület szerinti tartalmi besorolása és különböző szempontok szerinti szűrése, valamint az egyes intézmények azonosítása (redukálva ezzel a duplikált, felesleges munkát). A fogalomalapú keresést követően a tudományos munkához szükséges hivatkozások gyors és praktikus exportálása is lehetséges. A *Dimensions* kereső magáncélú felhasználására is mód van, az alábbi webcímen bárki számára kipróbálható: <https://app.dimensions.ai/discover/publication>.

(Ez a cikk a Könyvtártudományi Szakkönyvtár weboldalán olvasható *A mesterséges intelligencia könyvtári alkalmazásáról szóló webinarium ismertetése* című bejegyzés szerkesztett változata.)