

# ChatGPT szerepe a tudásmenedzsmentben

Pitlik László

egyetemi docens  
Kodolányi University  
pitlik.laszlo@kodolanyi.hu  
ORCID: 0000-0001-5819-0319

## Absztrakt

A chatGPT megjelenésével az emberek tudásról alkotott szubjektív érzete jelentősen megváltozott: ma már minden érintett számára világos, az őszinte magához, mit is jelent valójában (objektíven) pl. a „büfészak” kifejezés? Minden büfészak, ahol a chatGPT sikeres vizsgát tud(na) tenni, sikeresen diplomát tud(na) szerezni. A KJE Informatika Tanszéke a chatGPT megjelenése óta folyamatosan kísérlet-sorozatokot végez az ember-gép-szimbiózis aktuális határainak feltárása érdekében. A chatGPT bevonása nem tilos, sőt kötelező a Hallgatóknak a jegyszerezés/szakdolgoztatás (az e mögötti programozás) kapcsán (ahogy a munka világában is az racionális körülmények között). Aki nem tud triviálisan jobb lenni, mint a chatGPT, az nem is érdemel pl. jegyet/diplomát. A kísérletek (esettanulmányok) tudásmenedzsmentet érintő tapasztalatai (hasonlóan az ELTE oktatásmódszertani konferenciájára bejelentett előadásokhoz) quasi triviálisak, de éppen ezért súlykalandók: Knuth értelmében tudás/tudomány az, ami forráskódba átirtható, minden más emberi aktivitás művészet. A chatGPT egy kód, vagyis a benchmark a chatGPT megjelenése óta immár nem feltétlenül egy másik ember, hanem a chatGPT által szimulált „átlag”-ember (embergép / gépember). A chatGPT egyszerre magántanár, levelezőtárs, nyelvtanár, iskolatárs, stb. Vagyis a chatGPT a vele kompatibilis személyiségek számára az önmegvalósítás eszköze. A chatGPT-vel való együttműködés a kérdezni tudás kultúrájának finomhangolását jelenti. A chatGPT meghaladásának alapja a kritikai érzék (a konzisztens gondolkodás, a bizonyításkultúra magas szintje).

**Kulcsszavak:** *ember-gép-szimbiózis, konzisztencia, hasonlóságelemzés, Knuth, automatizálás, munkaerő piac, információs többletérték, minőségbiztosítás, kockázatelemzés, mesterséges intelligencia*

## 1. Bevezetés

Az előadás és a cikk a chatGPT megjelenése óta több mint 100 egyetemi Hallgatóval folytatott oktatói kooperáció tapasztalatait mutatja be esettanulmány jelleggel. A tapasztalatok kiterjednek a chatGPT használatának sikeres és sikertelen scenárióira, ahol sikerességről akkor beszélünk, ha az emberi döntéshozó (munkaadó) integrálni tudja/tudta a

chatGPT teljesítményeket a napi céges folyamatokba. Sikereség esetén is vizsgálendő volt, az emberi munkaerő milyen részletek tekintetében mennyivel több bérköltség fejében tud többet nyújtani, azaz értelmezhető-e az információs többletérték fogalma. A sikertelenség esetén vizsgálendő volt, vélemezhetően milyen chatGPT-jellegű fejlesztések mellett hidalhatók át a hiányosságok a jövőben, illetve, ha vannak legalább részleges áthidalási lehetőségek, akkor ezek piaci megjelenése mikorra várható? A chatGPT munkaerő piaci értelmezései mellett a chatGPT más mesterséges intelligenciákkal való kapcsolata is vizsgálatra került.

## 2. Előzmények

A chatGPT által generált publicitás mellett szakirodalmi hiány formálisan nem merülhet fel (vö. pl. csak az eltérő nyelveken elérhető Youtube-objektumok):

- [https://www.youtube.com/results?search\\_query=chatgpt+risks+chances](https://www.youtube.com/results?search_query=chatgpt+risks+chances)
- [https://www.youtube.com/results?search\\_query=chatgpt+risks+chances](https://www.youtube.com/results?search_query=chatgpt+risks+chances)
- [https://www.youtube.com/results?search\\_query=chatgpt+kock%C3%A1zatok+%C3%A9s+es%C3%A9lyek](https://www.youtube.com/results?search_query=chatgpt+kock%C3%A1zatok+%C3%A9s+es%C3%A9lyek)

Átfogó szakirodalmi művekre hivatkozni a tételes kísérleti „jegyzőkönyvek” helyett itt és most irreleváns, lévén a már ismertté vált esetek leíró jellegű értelmezése a quasi végtelen mennyiségűvé és egyre differenciáltabbá váló (jövőbeli) sokaságra vonatkozóan ennek adott részhalmaza alapján csak önkényes lehet.

A chatGPT szövegeinek értelmezése lektori, dramaturgi feladat, így minden korábbi minőségi szövegalkotási tudás/kompetencia a chatGPT-vel való emberi együttműködés releváns rétege: vö.

- Academic writing skills (AWS): pl. [https://wilson.fas.harvard.edu/files/jeffreywilson/files/jeffrey\\_r.\\_wilson\\_academic\\_writing.pdf](https://wilson.fas.harvard.edu/files/jeffreywilson/files/jeffrey_r._wilson_academic_writing.pdf)
- Roter Faden: [Beate Leßmann - Material - Roter Faden: Aufgaben, Vorlagen, Tests \(beate-leßmann.de\)](https://www.beate-leßmann.de/material/roter-faden-aufgaben-vorlagen-tests/beate-leßmann.de)

A saját kutatások kényszerűen eddig a szövegalkotás minőségéből kiindulva a szövegek (megértés-szimulálás nélküli) értékességének összehasonlító elemzése irányába mozdultak el: vö.

- Szabvány-jellegű elvárások és értelmezések: <https://miau.my-x.hu/myx-free/index.php3?x=test1>, ill. <https://miau.my-x.hu/myx-free/index.php3?x=test11>
- Szöveges pályamunkák robotizált értékelése: <http://miau.my-x.hu/miau/140/la140.doc>
- Robotíró, írásminőség-elemzés:

- [https://miau.my-x.hu/miau/300/otdk\\_biralok\\_biralata.docx](https://miau.my-x.hu/miau/300/otdk_biralok_biralata.docx)
- [https://miau.my-x.hu/miau/273/otdk\\_2021\\_abstract\\_standards.docx](https://miau.my-x.hu/miau/273/otdk_2021_abstract_standards.docx)
- [https://miau.my-x.hu/digeco/2020/2020osz/digeco\\_tdk\\_publication.docx](https://miau.my-x.hu/digeco/2020/2020osz/digeco_tdk_publication.docx)
- [http://miau.my-x.hu/miau/224/jo\\_fogalma\\_otdk\\_biralat\\_anonimizalt\\_2.docx](http://miau.my-x.hu/miau/224/jo_fogalma_otdk_biralat_anonimizalt_2.docx)
- [http://miau.my-x.hu/miau/200/otdk\\_v2.doc](http://miau.my-x.hu/miau/200/otdk_v2.doc)
- [http://miau.my-x.hu/miau/200/otdk\\_v1.doc](http://miau.my-x.hu/miau/200/otdk_v1.doc)
- [http://miau.my-x.hu/miau/182/etdk2013\\_baczai\\_gyorgyne.pdf](http://miau.my-x.hu/miau/182/etdk2013_baczai_gyorgyne.pdf)
- ...

### 3. Módszertani alapok

Ebben a cikkben/előadásban tehát nem egy önkényes rendszertan-teremtési kísérlet kerül bemutatásra, hanem sokkal inkább a kísérletben spontán résztvevő egyetemi Hallgatók, mint leendő döntéshozók/munkavállalók gondolkodásmódjának ösztönös chatGPT-interpretációs mintázatai kerülnek intuitív módon értelmezésre annak érdekében, hogy az ember-gép-szimbiózisról remélhessünk többet megtudni.

A téma (bár látszólag) informatikai jellegű, klasszikus humán-etológiai kísérletként egyelőre nem fogható fel. A jelen fázis sokkal inkább jövőbeli humán-etológiai kísérletek előkészítését teszi lehetővé, ahogy interjúkkal kérdőíves felmérések alapozhatók meg.

A verselemzés-jellegű módszertan felvállalásának további okai:

- A chatGPT maga is változik (pl. a fejlesztői által ad hoc jelleggel, önkényesen kinyitott vagy bezárt funkcionálisitási kapuk tetten érhetőségei révén, ill. a chatGPT működésének önreflexív volta révén).
- Az érintett emberek é a chatGPT viszonya képes jelentősen változni bármilyen pozícióból bármilyen pozícióba az egyre újabb és újabb esettanulmányok megismerése nyomán.

Mit is várhat akkor az Olvasó ettől a cikktől:

- A saját kísérletek publikus URL-jeinek listáját, vagyis sok-sok érdekes nyers esetleírást (vö. emberi inputok + chatGPT-outputok, mint egy fajta nyers log-adatok).
- A chatGPT-vel való ember együttműködés lehetséges vetületeire vonatkozó (pl. korrekúra-alapú, lábjegyzet-alapú, kommentár-alapú) nyers (log-)adat-feldolgozás típushelyzeteinek bemutatását.
- A chatGPT munkaerőpiaci értékét becslő kockázat-elemzéseket, minőségbiztosítási lépéseket, információs többletérték-bebecsléseket.

- A tudás(inkl. nem-tudás) és a tudásmenedzsment fogalmának újragondolását: vagyis az ember-gép-szimbiózisok által megjelenő új (cyborg)-tudásformák karakterisztikáira való rámutatást.

A cikk szerversen kapcsolódik más (quasi párhuzamos) konferenciák (vö. IKSAD, ELTE) dokumentumaihoz: pl.

- <https://miau.my-x.hu/miau2009/index.php3?x=e181>
- <https://miau.my-x.hu/miau2009/index.php3?x=e0&string=r%C3%A1cz.l>

#### 4. Eredmények

A chatGPT kapcsán a saját kísérletek dokumentálása megfelel az adatbázis-építés, a log-képzés, a példatár-fejlesztés elveinek. Ezek a gyűjtemények elsődlegesen kontextus-orientáltak.

A nyers chatGPT-kísérletek értelmező korrektúrái olyan emberi észrevételek, kritikák, teljesítmény-fokozó emberi beavatkozások, kommunikációs trükkök (vö. az emberi kérdés stílus-árnyalatainak hatásait felmérő alternatívák chatGPT-válaszok, adott ember-gép-konverziós állapotból következő gráf-jellegű konverzió-vezetési alternatívák tételes mibenlétének dokumentálása, stb.), melyek a nyers chatGPT-állapotok szignifikáns feljavításához vezetnek azáltal, hogy az ezekkel szembesülő emberi szakértő számára megnyílnak saját helyzetére alkalmas adaptációkat megalapozó lehetőségek:

#### 5. URL-katalógus, avagy nyers LOG-ok és ezek értelmezései

Saját kísérletek katalógusa a cikk lezárásának pillanatáig (2023.11.07.) – inkl. korrektúrák/lábjegyzetek/stb.:

1. A MY-X knuth-ató csoport szervere külső szemmel:  
<https://www.google.com/search?q=chatGPT+site%3Amiau.my-x.hu>
2. A MY-X team saját hírfolyama:  
[https://miau.my-x.hu/miau2009/index\\_tki.php3?\\_filterText0=\\*chatGPT](https://miau.my-x.hu/miau2009/index_tki.php3?_filterText0=*chatGPT)
3. A MY-X team saját wikipedia-jellegű szolgáltatásai:  
<https://miau.my-x.hu/mediawiki/index.php?title=Speci%C3%A1lis%3AKeres%C3%A9s&search=chatGPT>
4. A MIAÚ már katalogizált chatGPT-orientált bejegyzései/cikkei:  
<https://miau.my-x.hu/miau2009/index.php3?x=e0&string=chatgpt>
5. Katalogizálásra váró munka-anyagok:  
<https://miau.my-x.hu/miau/304/>
6. Humán-etológiai kísérletek téma:  
[https://miau.my-x.hu/mediawiki/index.php/ChatGPT\\_cyborg-etology\\_human-machine-interactions](https://miau.my-x.hu/mediawiki/index.php/ChatGPT_cyborg-etology_human-machine-interactions)

7. Oktatást támogató IT-megoldások részleges katalógusa:  
[https://miau.my-x.hu/mediawiki/index.php/Oktatast\\_tamogato\\_IT\\_megoldasok](https://miau.my-x.hu/mediawiki/index.php/Oktatast_tamogato_IT_megoldasok)

## 6. Információs többletérték-bebecslés: minőség és kockázat

Ahogy a konvencionális mezőgazdaság mellett megjelent a környezetvédelmi-/bio-/ökológiai-aspektusok tömege, úgy a chatGPT is egy fajta újszerű kiegészítésként hat a tudásmenedzsment eddigi szakterületére. Ahogy a mezőgazdaság esetén sem változott meg a jövedelem, haszon, változó költség, állandó költség, fedezeti hozzájárulás, eredménykimutatást, stb. logikája, hiszen az új költség és/vagy bevételi pozíciók megjelenése magán a kalkulációs sémán érdemben nem kellett, hogy változtasson, úgy az információs többletérték-bebecslés módszertana, gyakorlata sem változik meg a chatGPT kedvéért: vö. <https://miau.my-x.hu/miau2009/index.php3?x=e0&string=added>

Az ember chatGPT-vel való együttműködése érdemben nem más, mint a gyaloglás helyett a lóra/biciklire/rekuperáló elektromos autóra való átszállás esete: a chatGPT új teljesítménykomponenseket (pl. egy ember, mint karmester által quasi párhuzamosan instruált gép-hadsereg: vö. looper-alapú egyszemélyes zene-alkotás) és új kockázatkomponenseket (pl. hallucinálás, felerősödő tudományos IZÉ-kultúra) teremt meg, ahol a felerősödő tudományos IZÉ-kultúra nem más, mint az emberek által mindig is saját gondolati instabilitásuk (vö. kellő hozzá nem értésük, pillanatnyi motiválatlanságuk, pedagógiai járatlanságuk, a befogadóval való „együttrezgés” hiánya, stb.) leleplezésére alkalmas szövegbányászati log-adatként értelmezhető magas absztrakciós szintű nyelvi elemek súlyos halmozódásának kockázata (vö. buzzword-jellegű, divatosan semmitmondó szavak, határozatlan számnevek, stb.): pl. chatGPT saját magáról adott magyar nyelvű önértelmezésében szereplő kifejezések kritikája (vö. nyers és ízé-színkódokkal ellátott konverziós értelmezések:

[https://miau.my-x.hu/miau/304/chatgpt\\_onvallomas\\_tudomanyos\\_ize\\_retegei.docx](https://miau.my-x.hu/miau/304/chatgpt_onvallomas_tudomanyos_ize_retegei.docx)).

A chatGPT-alapú (start-up-jellegű) üzleti modellekre példa a SOCIAL AI projekt: <https://www.youtube.com/watch?v=zrq5ssPQK5M> - A SOCIAL AI projekt kockázatelemzése már laikus (potenciális jövőbeli alkalmazók) által elvégzett tesztelés keretében is triviálisan tetten érhető, ahol a kockázatelemzés és a minőségbiztosítás ugyanazon érme két oldalát jelenti: a minőségbiztosítás feltárja a gyenge pontokat, ezek hálózátát, s a kockázatelemzés igyekszik megoldási alternatívákat kínálni a gyengeségekre az erősségek el-nem-rontása mellett: pl. tudományos-IZÉ-kultúrát képviselő szavak/szerkezetek tételes feltárása, ill. ezek tiltása a chatGPT számára.

## 7. Cyborg-tudás és ennek menedzsmentje

A cyborg-tudás sokkal inkább nevezendő adaptációs kompetenciának, mint klasszikus értelemben magolással, gyakorlással ténylegesen, jól tesztelhetően elsajátítható, ill. tanítható tudásnak: a cyborg-létforma olyan, mint az abszolút zenei hallás, mely, ha nincs, akkor a szolfész-gyakorlatok egy fajta álcahálóként úm. segítenek a teljes laikusságból való kiemelkedésben, de ettől még nem lesz valaki mestere a zenének, ill. a gép-emberszimbiózisnak. Szerencsére nem csak azok élnek meg a zenélésből, akiknek abszolút hallása van, hanem még nagyon sokan mások is, akiknek egyértelműen korlátozottak a zenei adottságai, de pl. csapatban (együttesben, kórusban, stb.) szerves részévé tudnak válni produkcióknak. Így a cyborg-létforma (a zene-ipar analógiájára) hasonlóképpen ki fogja termelni a szinteket, kasztokat, stb., ill. az ezek közötti tetszőleges átjárásokat, ahol az átjárás sikerességének gyakoriság a megtett úttal fordítottan lesz itt is arányos. Következésképpen minden zsumnaliszta eufória és/vagy katasztrófa-hangulatkeltés felesleges, álságos és így kerülendő a chatGPT és minden hasonló erőtér kapcsán is. Az emberi faj, mint eszközhasználó, feltaláló faj eredete kezdetei óta cyborg-létet folytat, s teszi ezt abban az értelemben sikeresen, hogy a gép-emberszimbiózis gyakorisága, mértéke, komplexitása az évezredek alatt csak egyre nőtt és környezeti/egyéb összeomlás nélkül továbbra is nőni fog.

## 8. Humán magatartásminták a cyborg-lét határán

A chatGPT kísérletek alapján félreérthetetlen emberi magatartásformák ismerhetők fel:

- A chatGPT-korlátainak felismerése: a potenciális alkalmazók már egyetlen egy kérdés-felelet pár alapján spontán gyanúkat képesek és akarnak megfogalmazni (vö. Turing-teszt - <https://hu.wikipedia.org/wiki/Turing-teszt>).
- Kíváncsiság, minőségbiztosítás: A spontán/tudatosan értelmezett emberi gyanúk kíváncsiságból (minőségbiztosítási szempontokból) több tesztet is elvezethetnek adott ember esetén.
- A tudatosan felismert / spontán megérzett chatGPT-korlátok feloldása: Akinek van ötlete arra, hogyan támogassa meg más inputokkal (kérdésekkel, triggerekkel) a chatGPT-t, az erre is kísérletet tesz – zömmel és azonnal egynél többet is.
- A feloldhatatlannak érzett kompetencia-hiányok tudatosítása: Amennyiben a saját kíváncsiság szülte kísérletek és/vagy a mások kísérleti tapasztalatai alapján az emberi felhasználó hitét veszti arra vonatkozóan (vö. pl. sikertelen Turing-teszt), hogy adott jellegű kihívás esetén a chatGPT bevonása sikerrel kecsegtethet, akkor kialakul egy fajta aknamező-térkép az emberi felhasználó fejében.
- Az eddig feloldhatatlannak hitt problémátípus újszerű megvilágításba helyezése: Amennyiben új, triviálisan az eddigi zavarokat feloldani képes impulzus kerül elő, akkor az emberi hit helyre áll, az aknamező térképe újra rajzolódik az ember fejében.

- A közgazdász-beállítottságú mérlegelés: Effektív chatGPT siker esetén is, vagyis amikor az elvitathatatlan az ember számára, hogy a chatGPT-képesség megállja a helyét pl. a Turing-teszt értelmében, egyes emberek alap-beállítottságukból fakadóan párhuzamosan azt is vizsgálják: megéri-e az ember helyett csak a chatGPT (vö. social AI-projekt potenciális automatizmusának lehetőségei/kockázatai) és/vagy a chatGPT és az ember együttesen mennyivel növeli meg az ember hatékonyságát (egy időegység alatti sikereinek számát)?

## 9. Vita

Az a kérdés természetesen racionális kérdés, vajon a cyborg-lét komplexitásának növekedése egy fenntartható pályának indikátora-e? A vélelem a nagy összeomlásig: igen. S ha lenne is nagy társadalmi összeomlás bolygó szinten, akkor is vizsgálendő: mi az ok, s mi az okozat. Természetesen ezen ok-okozatiság nem utólagos bölcselkedések formájában értékes az emberiség számára, hanem előrejelzések formájában (vö. klímaváltozás), mely előrejelzéseket azonban úgy kell megtenni, hogy az idő múlásával az előrejelzett értékek és a tények összevetésének eredményeként az előrejelzések adott pillanatában vélt világmegértési rendszer (a modell) értéke tetten érhető legyen mindenkor (vö. ismét csak klímamodellek, időjárás-előrejelző-modellek, tőzsdei modellek, stb.). Az emberiség, a tudomány maga egyelőre még mindig szenved a JÓ fogalmának KNUTH-i jellegű<sup>1</sup>, azaz robotokba való átadni tudása, akarta kapcsán: vagyis mind a mai napig nem tudható tetszőleges biztonsággal előre, hogy több modell közül a jövőben majd melyiket fogják utólagosan a legjobbnak tekinteni (vö. céltalanság tétele). S mivel ez a tétel a mesterséges intelligenciának nevezett (zömmel soft) matematika fejlődését is permanensen áthatja, így az MI bármilyen feletti eluralkodása irracionális (zsurnaliszta) félelem, hiszen a céltalanság tételének feloldása inkább tűnik lehetetlennek, mint lehetségesnek – még a konzisztencia-alapú modellezés komplexitás-optimalizáló elvárása és automatizációs törekvései mellett is...

## 10. Következtetések

A chatGPT kapcsán eddig katalogizált humán-etológiai kísérleti felhívások (vö. [https://miau.my-x.hu/mediawiki/index.php/ChatGPT\\_cyborg-etology\\_human-machine-interactions](https://miau.my-x.hu/mediawiki/index.php/ChatGPT_cyborg-etology_human-machine-interactions)) alapján világosan látható, hogy a chatGPT nem csodaszor: a vele való emberi együttműködés egy fajta művészet (vö. zenei analógia). A chatGPT egy zeneszerszám, mely ki fogja termelni a maga zeneiparát...

<sup>1</sup> [https://miau.my-x.hu/miau2009/index\\_tki.php3?\\_filterText0=\\*knuth](https://miau.my-x.hu/miau2009/index_tki.php3?_filterText0=*knuth)

## 11. Jövőkép

A chatGPT, mint valószínűségi alapú fecsegőgép az embernek, mint olyannak tart tükörképet: minden, ami a chatGPT-ben kockázatos annak forrása az ember: a programozó ember kevésbé, mint a log-adatokat adó szövegeket alkotó ember(tömeg). Amilyen szomorú/megnyugtató képet mutat a chatGPT logikátlansága, tudományos IZÉ-hajlama, stb., olyan szomorú/megnyugtató a biológiai ember értékességének sokszínűsége.

Így az ember maga fog változni (akarni és nem akarva is sodródni) egyre komplexebb gondolati struktúrák felé, s így egyre értékesebb log-ot fog tudni produkálni milliárdszám a chatGPT-nek, mely ennek ellenére sem fog tudni (talán soha) egymillió-dolláros matematikai tételleket bizonyítani a jelenlegi paraméterei, struktúrái mentén, mert a bizonyítások nagyszerűségéhez valami olyat (ritkát, egyedit) kell meglátni (vö. Ramanudzsán), amit egyelőre még csak a biológiai ember képes, lévén az intuíció fizikája még ismeretlen területe a tudományos világmegismerésnek (vö. vallások, ma még ezoterikus összefüggés-vélelmek, stb.).

## Összefoglalás

A chatGPT bevonása nem tilos, sőt kötelező a Hallgatóknak a jegyszerzés/szakdolgoztatrás (az e mögötti programozás) kapcsán (ahogy a munka világában is az racionális körülmények között). Aki nem tud triviálisan jobb lenni, mint a chatGPT, az nem is érdemel pl. jegyet/diplomát. A kísérletek (esettanulmányok) tudásmenedzsmentet érintő tapasztalatai (hasonlóan az ELTE oktatásmódszertani konferenciájára bejelentett előadásokhoz) quasi triviálisak, de éppen ezért súlykolandók: Knuth értelmében tudás/tudomány az, ami forráskódba átirható, minden más emberi aktivitás művészet. A chatGPT egy kód, vagyis a benchmark a chatGPT megjelenése óta immár nem feltétlenül egy másik ember, hanem a chatGPT által szimulált „átlag”-ember (embergép / gépember). A chatGPT egyszerre magántanár, levelezőtárs, nyelvtanár, iskolatárs, stb. Vagyis a chatGPT a vele kompatibilis személyiségek számára az önmegvalósítás eszköze. A chatGPT-vel való együttműködés a kérdezni tudás kultúrájának finomhangolását jelenti. A chatGPT meghaladásának alapja a kritikai érzék (a konzisztens gondolkodás, a bizonyításkultúra magas szintje).

## Irodalomjegyzék

...lásd a szövegekőben...

A speciális (pl. csoportlinkek) miatt klasszikus irodalomjegyzék generálása logikailag értelmezhetetlen.