

A mondat, amely többet ígér

Slug: a-mondat-amely-tobbet-iger | Publikálva: 2026. 07. 09. | Tags: Állításbiztonság, AI-audit, AI-governance, Információs integritás

A bemutatóteremben minden mondat rendben van. Az új AI-rendszer negyven százalékkal gyorsítja az ügyintézés, csökkenti a hibákat, igazságosabb döntéseket támogat, és megfelel az európai szabályozásnak. A diák elegánsak, a grafikonok felfelé mennek, a lábjegyzetek pedig olyan apr...

Állításbiztonság, AI-szöveg és a magabiztos forma kockázata

A bemutatóteremben minden mondat rendben van. Az új AI-rendszer negyven százalékkal gyorsítja az ügyintézés, csökkenti a hibákat, igazságosabb döntéseket támogat, és megfelel az európai szabályozásnak. A diák elegánsak, a grafikonok felfelé mennek, a lábjegyzetek pedig olyan aprók, hogy már-már dekorációként működnek.



Illusztráció: A mondat, amely többet ígér

A négy állítás közül akár mindegyik lehet részben igaz. Ettől még egyik sem használható automatikusan úgy, ahogy a bemutató sugallja. A negyven százalékhöz kell kiinduló folyamat, mérési időszak és összehasonlítható feladat. A hibacsökkenéshez tudni kell, milyen hibát mértek, és mit hagytak ki. Az igazságosság nem egyetlen technikai mutató. A jogszabályi megfelelés pedig nem olyan tulajdonság, amelyet egy termék egyszer megszerez, aztán örökre viselhet a dobozán.

A mondat itt nem feltétlenül hazudik. Inkább előre elkölt olyan bizonyosságot, amelyet a bizonyíték még nem termelt meg.

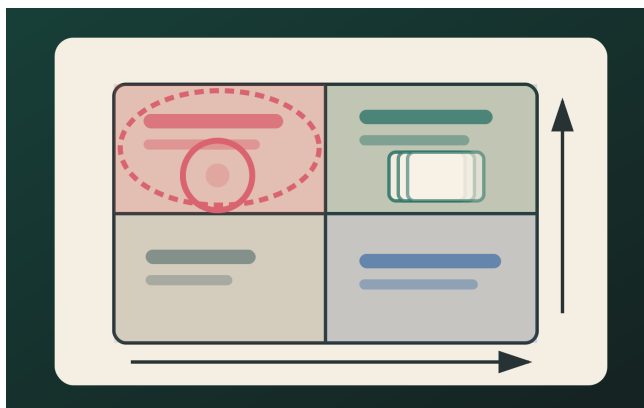
Ez a probléma régebbi a generatív AI-nál. Reklám, politikai kommunikáció, szakértői nyilatkozat és ügyfélszolgálati levél régóta tudja, hogyan lehet egy részben igaz állítást teljesebbnek, általánosabbnak vagy véglegesebbnek mutatni. A nyelvi modellek annyiban változtattak a helyzeten, hogy ugyanezt a formai magabiztosságot ipari sebességgel képesek előállítani. A mondat gördül. A bizonyíték néha csak kullog utána.

A forma nem viseli a bizonyítás terhét

Olvasás közben ritkán állunk meg minden mondatnál, hogy külön leltárt készítsünk a premisszáiról. A világos szerkezet, a nyugodt ritmus és a szakmai hang megkönnyíti a befogadást. Ez hasznos. Egy rosszul szerkesztett, zavaros szöveg akkor is akadály, ha történetesen pontos.

A gond ott kezdődik, amikor a könnyű olvashatóságot a megbízhatóság jelének tekintjük. A természetesnyelv-generálás kutatása több feladatban dokumentálta, hogy a rendszer képes a forrásban

nem szereplő vagy ténytyszerűen hibás részleteket is folyékonyan előállítani [1][2]. Ebből nem következik, hogy minden AI-szöveg hamis, és az sem, hogy emberi szerzőknél ne létezne ugyanilyen probléma. A szűk, de fontos következtetés ennyi: a nyelvi minőség és a forráshűség két külön vizsga.



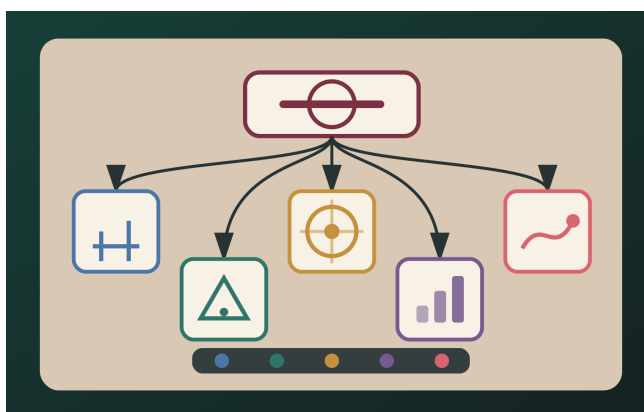
Illusztráció: A mondat, amely többet ígér

Bender és munkatársai arra figyelmeztettek, hogy a nyelvi teljesítmény könnyen összekeverhető a megértés, a szándék és a társadalmi felelősség emberi fogalmaival [3]. A vita messzebbre nyúlik annál, amit egyetlen cikkben rendezni lehet. A gyakorlati következmény viszont kézzelfogható. Egy emberinek ható válasz mögött nem feltétlenül van emberi értelemben vett mérlegelés, emlékezet vagy felelősségvállalás. Az olvasónak ezért nem a hang alapján kell eldöntenie, mit bír el a szöveg.

A rendszer ugyanazzal a kellemes nyugalommal képes leírni ellenőrzött tényt, valószínű becslést, általános tanácsot és kitalált részletet. A hangnem nem jelzi, melyik mondat melyik kategóriába tartozik. Ha a szöveg ezt sem jelzi, a különbség az olvasó nyakába kerül.

Egy mondatban több állítás is elbújhat

"A megoldás biztonságosabbá teszi a működést." Elsőre egyetlen állításnak látszik. Valójában legalább öt kérdést sűrít össze.



Illusztráció: A mondat, amely többet ígér

Mihez képest biztonságosabb? Milyen káreseményt tekintünk relevánsnak? Milyen környezetben mérték? Mekkora a változás? Ki viseli a fennmaradó kockázatot?

Ha ezekre nincs válasz, a mondat jelentése lebeg. Az egyik olvasó kevesebb adatvesztést ért alatta, a másik kisebb jogi kockázatot, a harmadik ritkább emberi hibát. A közlés közben úgy tesz, mintha ugyanarról beszélnének.

Ugyanez történik a számokkal. "A rendszer harminc százalékkal növelte a pontosságot." Ez lehet

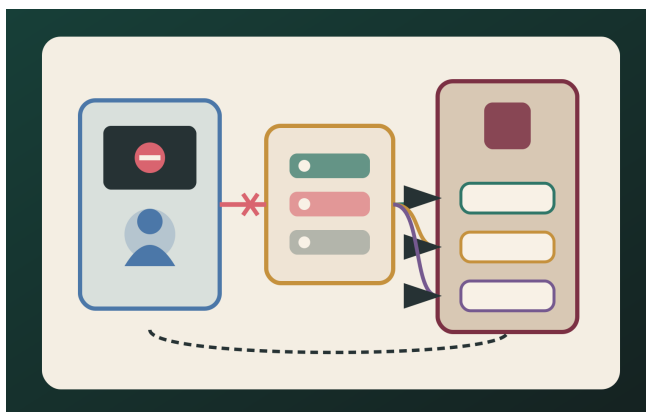
látványos javulás, és lehet ügyes optikai trükk. Nem mindegy, hogy 50-ről 65 százalékra, vagy 97-ről 97,9-re változott az eredmény; ahogy az sem, hogy száz tesztesetből vagy egymillió valós ügyből számolták. A relatív változás önmagában ritkán mond eleget a gyakorlati hatásról.

A jogi és szabályozási állítások különösen könnyen viselnek kölcsönzött tekintélyt. Az "EU AI Act-kompatibilis" felirat jól mutat egy prezentációban, de csak akkor informatív, ha tisztázza, milyen rendszerre, szereplőre, felhasználási módra és kötelezettségre vonatkozik. Az uniós AI-rendelet kockázati és szerepkörökhöz kapcsolódó követelményeket állapít meg [8]. Ebből nem következik egyetlen általános pecsét, amely minden kontextusban ugyanazt jelenti.

A mondat tehát gyakran nem azért túlterhelt, mert hamis elemeket tartalmaz. Azért, mert több eltérő bizonyítási feladatot egyetlen sima felszín alá rejt.

A "rendszer nem engedi" sajátos állatfaj

Az ügyfélszolgálati kommunikáció egyik kedvelt mondata, hogy a rendszer nem tesz lehetővé valamit. Ez lehet tényszerű leírás. Egy operátor valóban nem tud új futárt küldeni, jóváírni egy összeget vagy módosítani egy szerződést. Csakhogy a technikai jogosultság hiánya nem azonos a szolgáltató teljes lehetetlenségével, és még kevésbé bizonyítja, hogy a felkínált rendezés elégséges.



Illusztráció: A mondat, amely többet ígér

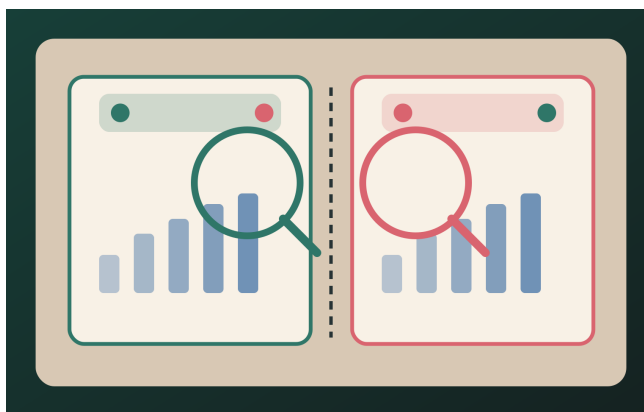
A mondat két szintet csúsztat egymásra. Az egyik a jelenlegi folyamat: az adott munkatárs és felület mit tud végrehajtani. A másik a felelősség: a szervezetnek milyen megoldást kellene létrehoznia. A kettő között lehet költség, belső szabály, szervezési nehézség vagy egyszerű üzleti döntés. Ezek valós korlátok, csak nem ugyanazok, mint a fizikai lehetetlenség.

Az AI könnyen tovább erősíti ezt a csúsztatást. Egy panaszlevél-generátor jogi következtetést tehet a történet végére, mert a szövegkörnyezet alapján oda illik. Egy ügyfélszolgálati bot együttérző bekezdéseket gyárt, miközben a tényleges problémára nincs műveleti útvonala. A kommunikáció ilyenkor puhább lesz, a helyzet viszont nem feltétlenül rendeződik.

A sablonos empátia önmagában nem bűn. Sok ember számára kifejezetten fontos, hogy ne hideg gépnnyelven kapjon választ. A probléma ott jelenik meg, ahol az érzelmi keret a megoldás helyére lép, vagy elfedi, hogy a rendszer ugyanazt a lezáratlan állapotot adja vissza szebb mondatokkal.

A technikailag igaz mondat is félrevezethet

Tversky és Kahneman klasszikus kísérletei megmutatták, hogy azonos következmények eltérő keretezése más döntéseket válthat ki [4]. Ebből nem érdemes azt a látványos következtetést levonni, hogy minden megfogalmazás manipuláció. Keret nélkül nincs közlés. Minden mondat kiválaszt, sorrendbe rendez és hangsúlyoz.



Illusztráció: A mondat, amely többet ígér

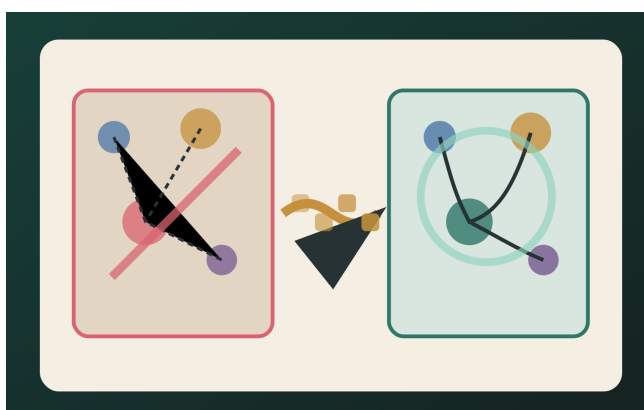
A felelősség akkor válik élessé, amikor a keret rendszeresen ugyanabba az irányba tolja a döntést, miközben a költséget, az alternatívát vagy a bizonytalanságot hátrébb rendezi.

"Akár 11 százalékos hozam." A mondat lehet szó szerint igaz, ha valamilyen feltétel mellett valóban előfordulhat ekkora eredmény. Az olvasó döntését azonban nem csak a maximum érdekli. Tudnia kellene a veszteség lehetőségéről, az időtávról, a díjakról, a kockázati eloszlásról és arról, milyen adatból származik a szám.

Grice társalgási maximái jól mutatják, miért lehet problémás egy technikailag igaz, de lényeges körülményt elhallgató közlés [5]. A hétköznapi kommunikációban nem pusztán szavakat értelmezünk. Feltételezzük, hogy a beszélő annyi releváns információt ad, amennyi a helyzethez szükséges. Ha egy reklám vagy AI-válasz tudatosan erre a feltételezésre épít, a félrevezetéshez nem mindig kell hamis mondat.

Ecker és munkatársai szerint a téves információ korrekciója azért is nehéz, mert az első magyarázat beépülhet a helyzet mentális modelljébe, és a későbbi cáfolat nem feltétlenül tölti ki az ott maradó rést [6]. Ez az AI-szövegeknél különösen fontos. A gyors, magabiztos első válasz akkor is irányt adhat a további gondolkodásnak, ha később pontosítjuk.

A javítás tehát nem merül ki abban, hogy a hibás mondat mellé odatesszük: "elnézést, ez pontatlan volt". A helyére működő magyarázat kell.



Illusztráció: A mondat, amely többet ígér

Az audit nem szövegrendőrség

Az audit szó könnyen egy hosszú ellenőrzőlistát idéz fel, amelyből a végén piros és zöld mezők maradnak. Erre néha szükség van, különösen szabályozott környezetben. A szöveg vizsgálata viszont

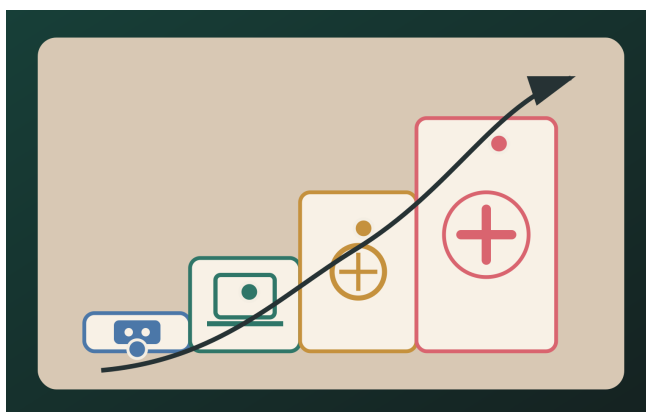
elveszíti az értelmét, ha kizárólag tiltott szavakat keres.

A "garantálja" valóban gyanús lehet, de a "hozzájárulhat" sem lesz automatikusan tisztességebb. Az utóbbi gyakran csak puhább kód. Nem támadható, mert alig mond valamit.

A használható szerkesztés megkeresi, mitől válik az állítás ellenőrizhetővé. Milyen adat áll mögötte? Milyen populációra érvényes? Mi volt az összehasonlítás? Milyen feltétel mellett bukik el? Milyen döntést próbál kiváltani, és mekkora kárt okozhat, ha téves?

A NIST AI Risk Management Framework az AI-kockázat kezelését a rendszer teljes életciklusához és a konkrét használati környezethez kapcsolja [7]. Ez a szemlélet a szövegre is átfordítható. Egy mondat kockázata nem csak a mondat belsejében van. Függ attól, ki olvassa, milyen jogosultsággal használja, milyen döntés követi, és van-e lehetőség emberi ellenőrzésre vagy korrekcióra.

Egy filmes ajánlóban elfér egy pontatlan mellékszereplőnév. Egy gyógyszeradagolási tanácsban már nem. Egy belső ötletelésben használható egy merész hipotézis, ha annak státusza világos. Egy pénzügyi ajánlatban ugyanez a mondat más terhet kap.



Illusztráció: A mondat, amely többet ígér

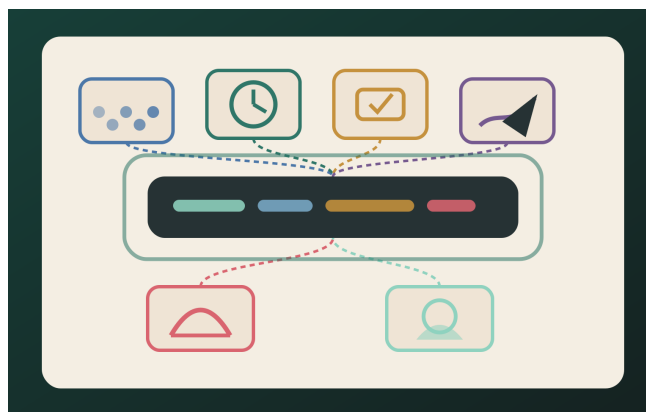
A jó audit ezért nem minden szöveget azonos szigorral mér. A tétet és a bizonyítékot kapcsolja össze.

A szerkesztőnek néha keményebb mondatot kell írnia

A bizonytalanság jelölése nem egyenlő a szöveg szétpuhításával. Sőt, a pontosabb változat gyakran keményebb, mert számon kérhető.

"Belső tesztjeink alapján a rendszer gyorsabbá teheti a feldolgozást." Ez óvatos, csak éppen információszegény.

"A 2026 márciusában futtatott, 480 számlából álló belső tesztben a rendszer emberi ellenőrzéssel átlagosan 18 százalékkal csökkentette a feldolgozási időt. A hibaarány nem változott érdemben. A teszt egyetlen dokumentumtípusra vonatkozott, ezért az eredmény más folyamatokra nem általánosítható."



Illusztráció: A mondat, amely többet ígér

A második mondat hosszabb, mégis kevesebb felesleges magabiztosságot tartalmaz. Megmutatja, mi történt, és azt is, hol ér véget az eredmény.

Természetesen ezt sem szabad díszletként használni. Ha a számok nem léteznek, nem kell feltalálni őket. Ha a minta gyenge, a helyes szerkesztés nem takarja el. A hiányzó adat maga is információ a döntéshez.

A szerzői profil szempontjából ez különösen fontos. A határozott hang nem attól hiteles, hogy mindent véglegesnek mutat. Attól, hogy pontosan megnevezi a konfliktust, végigviszi az oksági láncot, majd ott áll meg, ahol a rendelkezésre álló anyag elfogy.

Mit ad ehhez a Fraktál Dialektika?

A Fraktál Dialektika ezen a területen nem igazsággépet ígér. Egy szöveg megbízhatóságát nem lehet egyetlen pontszámban maradéktalanul elintézni, mert eltérő állítások, felhasználások és következmények futnak benne együtt.

A keret inkább azt erősíti, hogy a mondatot abban a helyzetben vizsgáljuk, ahol működni kezd. Mit kapcsol össze? Milyen feltételt hallgat el? Mely bizonytalanságot tünteti el a stílus? Kinek a döntését tolja, és mi történik, ha az olvasó szó szerint veszi?

Ez magas szintű szemlélet, nem a belső FD-rendszer publikált leírása. A gyakorlati értéke abban van, hogy a szöveget nem engedi elszakadni a forrásától és a következményétől. A mondat többé felelősséget hordozó kapcsolat, és elveszíti pusztán dekoratív szerepét.

A határ itt is fontos. Egy ilyen keret nem helyettesít szakterületi vizsgálatot. Jogi megfeleléshez jogász, orvosi tanácshoz egészségügyi szakember, statisztikai következtetéshez megfelelő módszertan kell. A szerkezeti audit legfeljebb megmutatja, hol keveredtek össze a feladatok, és hol kér a szöveg olyan bizalmat, amelyhez még nincs elég alap.

A mondat akkor lesz erősebb, amikor kevesebbet játszik el

A modern szöveg nem attól modern, hogy rövid sorokban beszél, szakmai kifejezéseket villant, vagy minden második bekezdés végén odatesz egy idézhető mondatot. Ezek formai lehetőségek. Néha működnek, néha csak eltakarják az üres helyet.

A használható mondat gyorsan eljut a konkrétumig. Nem rejti el a döntéshez szükséges feltételeket. Nem kéri, hogy a hang alapján hitelegessünk meg olyan bizonyosságot, amelyet a forrás még nem igazolt. És nem tesz úgy, mintha a korlátok bevallása gyengeség volna.

A legjobb szerkesztés után nem feltétlenül marad látványosabb szöveg. Marad viszont egy állítás,

amelyről tudni lehet, mire épül, hol használható, és mi változtatná meg.

Ez kevesebb ígéretnek tűnhet. Valójában több teherbírás.

Hivatkozások

[1] Maynez, Joshua et al. (2020): "On Faithfulness and Factuality in Abstractive Summarization." Proceedings of ACL 2020, 1906-1919. DOI: 10.18653/v1/2020.acl-main.173.

[2] Ji, Ziwei et al. (2023): "Survey of Hallucination in Natural Language Generation." ACM Computing Surveys, 55(12), Article 248. DOI: 10.1145/3571730.

[3] Bender, Emily M.; Gebru, Timnit; McMillan-Major, Angelina; Shmitchell, Shmargaret (2021): "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?" FAccT '21, 610-623. DOI: 10.1145/3442188.3445922.

[4] Tversky, Amos; Kahneman, Daniel (1981): "The Framing of Decisions and the Psychology of Choice." Science, 211(4481), 453-458. DOI: 10.1126/science.7455683.

[5] Grice, H. P. (1975): "Logic and Conversation." In Cole, P.; Morgan, J. L. (eds.): Syntax and Semantics 3: Speech Acts. Academic Press, 41-58.

[6] Ecker, Ullrich K. H. et al. (2022): "The Psychological Drivers of Misinformation Belief and Its Resistance to Correction." Nature Reviews Psychology, 1, 13-29. DOI: 10.1038/s44159-021-00006-y.

[7] National Institute of Standards and Technology (2023): Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1. DOI: 10.6028/NIST.AI.100-1.

[8] Európai Parlament és Tanács (2024): Az Európai Parlament és a Tanács (EU) 2024/1689 rendelete a mesterséges intelligenciára vonatkozó harmonizált szabályok megállapításáról. Az Európai Unió Hivatalos Lapja, 2024. július 12.