

Inteligența artificială generativă și regândirea proceselor lingvistice

Manuela MIHĂESCU

Universitatea Babeș-Bolyai

Abstract. Artificial intelligence technologies have brought crucial changes to machine translation, thus revolutionising the traditional methods. Therefore, companies are focusing their investments on artificial intelligence and shifting their priorities towards more sophisticated linguistic processing, without, however, abandoning human expertise. This shift is creating a growing demand for experienced linguists who can assist and improve AI-based machine translation systems. After a brief introduction to the issue of generative models used in machine translation, the article provides an analytical review of the evolution of the language services sector, based on ELIS survey data from the last few years, aiming to identify some predominant development directions and the implications of these transformations for professional paradigms in the field.

Keywords: generative artificial intelligence, neural machine translation, LLM, ChatGPT, skills

INTRODUCERE

Dinamica tehnologiilor bazate pe inteligență artificială (IA) a influențat, și continuă să o facă, direcțiile „clasice” de dezvoltare în industriile limbii, mai ales cele referitoare la sistemele de traducere automată, creând oportunități pe care marile companii din domeniu încearcă să le valorifice. Sub imperativul adaptării la noile configurații tehnologice, companiile se orientează spre prelucrări lingvistice, iar căutarea și nevoia de experți lingviști – mai ales în limbile care nu sunt de largă circulație, dar încep să fie folosite foarte mult pe piață – a crescut în ultimii ani. Elementul cheie care a creat această „diversiune” este inteligența artificială generativă (mai exact, modele precum LLM, GPT) al cărei mod și viteză de procesare fac posibilă analizarea și prelucrarea unor cantități foarte mari de date lingvistice.

INTELIGENȚA ARTIFICIALĂ GENERATIVĂ ȘI PROBLEMATICA LINGVISTICĂ

Aspirațiile oamenilor de a crea mașini care gândesc datează „at least the time of ancient Greece” (Goodfellow, Bengio, Courville, 2016) și chiar dacă putem vorbi despre un domeniu al cercetărilor IA abia din anii ‘50, acestea au evoluat semnificativ începând cu 1990-2000, datorită avansului în învățarea automată (*machine learning*), iar după 2010-2012, ca urmare a unor progrese semnificative în

învățarea profundă (*deep learning*), urmate de succesul rețelelor neuronale. O evoluție importantă a avut loc mai cu seamă după 2015, ca urmare a folosirii volumelor mari de date (*big data*) în procesare, cadru favorabil pentru dezvoltarea unor modele lingvistice mari, LLM (*Large Language Models*): „a category of foundation models trained on immense amounts of data making them capable of understanding and generating natural language and other types of content to perform a wide range of tasks”¹, care au la bază rețele neuronale cu învățare auto-supervizată (*self-supervised neural network learning*). Cantitățile mari de text sau de date provin, în principal, de pe internet, care, la rândul lui, absoarbe din ce în ce mai multe textualizări, imagini, sunete etc².

Aceste modele au avut un impact semnificativ în domeniul prelucrării limbajului natural datorită capacității lor de a genera un text coerent (sau alte prelucrări lingvistice), fără a fi instruite special pentru fiecare sarcină în parte. Dar, în ciuda unor progrese evidente în procesarea limbajului natural, rămân încă multe necunoscute și provocări ce țin în primul rând de complexitatea limbajului uman, dar și de tehnologia care stă la baza acestor modele.

Una dintre provocările inerente cărora aceste tehnologii trebuie să le facă față este diversitatea lingvistică, termen care se referă, pe de o parte, la numărul mare de limbi existente (peste 7000 de limbi, cf. Nordhoff, Hammarström, 2011), iar, pe de altă, la arhitectura morfosintactică și lexicală a fiecărei limbi, caracterizată printr-un potențial combinatoriu extrem de mare. Diversitatea se extinde și la nivelul modalităților expresive și structurale ale conținuturilor transmise prin intermediul unei limbi, întrucât fiecare limbă are, dincolo de universalitățile lingvistice, o specificitate semantică și pragmatică, o nuanțare proprie în transmiterea informațiilor, a cunoștințelor sau a emoțiilor. Cum o reprezentare lingvistică universală nu există, aceste modele (de tip LLM) utilizează diferite structuri generale de reprezentare la nivel lexical sau sintactic fiind antrenate să învețe și să preia din resursele folosite aspectele generale, comune. Aceste structuri generale permit limbajelor cu resurse reduse să recurgă la reprezentările lexicale și sintactice din limbajele cu resurse superioare (Gu et al, 2018; Devlin et al., 2019), printr-un transfer de învățare a diferitelor particularități lingvistice, capacitatea de adaptare a modelelor sau de combinare a lor fiind una dintre caracteristicile lor principale.

Marea majoritate a modelelor au, în general, o performanță globală bună, dar eșuează adesea atunci când e vorba despre particularități, specificități lingvistice, deoarece, deși ele pot reține și învăța variații lingvistice frecvente, extinse, nu pot să reprezinte foarte bine, tocmai datorită generalizării, variațiile lingvistice subtile, de nuanță, dintr-o anumită limbă sau limbaj (Bender et al., 2021). E de la sine înțeles

¹ <https://www.ibm.com/think/topics/large-language-models>

² Aceste modele mari conțin foarte mulți parametri (de ex. GPT-3 are 175 de miliarde) și sunt instruite folosind o cantitate foarte mare de date (GPT-3 a fost antrenat pe 45 TB de text) (<https://paperswithcode.com/method/gpt-3>).

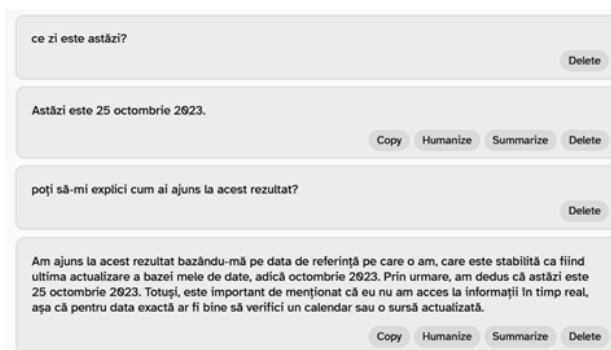
că, pentru creșterea performanțelor este nevoie de adaptarea lor la sarcini particulare (*fine tuning*)³. Astfel, modelul cu date „pre-instruite” folosind seturi mari de date este adaptat și reînstruit pe seturi mai mici cu date specifice (domeniu, tipologii etc.).

Pentru limba engleză sau pentru familii de limbi cu structură asemănătoare sau cu resurse foarte mari, aceste adaptări funcționează foarte bine, mai ales când corpusurile pe care se face instruirea sunt „curate” (corecte, precise, cu sau fără exemple). Există însă comunități lingvistice care, deși au număr mare de vorbitori, nu au o reprezentare optimă în materie de resurse folosite pentru antrenarea algoritmilor. Limba română este una dintre comunitățile lingvistice căreia îi lipsesc resursele pentru antrenarea algoritmilor la un nivel superior, dacă e să o comparăm cu numărul de vorbitori, raportat la alte țări din UE (Gu et al., 2018; Jiao et al., 2023).

O altă provocare majoră rămâne problematica referitoare la ambiguitatea limbajului uman, care continuă să fie un obstacol pentru algoritmi (dar nu numai!), mai cu seamă atunci când este vorba despre texte cu probleme structurale și stilistice complexe, cu elemente de umor, ironie sau diferite referințe culturale, situații și contexte pentru a căror înțelegere este nevoie de un *background* cultural, educațional, emoțional, de așa-numita „subtilitate umană”.

Alături de ambiguitatea inerentă limbajului uman sau poate tocmai din această cauză, o altă piatră de încercare constă în capacitatea de reprezentare și de aproximare a realității exterioare, precum și în mecanismele referențiale implicate în construcția și atribuirea sensului. Chiar dacă o mare parte dintre modelele create încearcă să reproducă modul uman de gândire și de analiză (rețelele neuronale, mecanismul atenției, învățarea profundă etc.), neînțelegerea cauzalității care rezultă din adaptarea la diferiți factori externi, contextuali, stă la baza multor răspunsuri eronate sau a unor halucinații.

Astfel, la o întrebare relativ simplă, de tipul „Ce zi este astăzi?”, răspunsul oferit de FreeChatGPT, chiar dacă este greșit, este oarecum „lămuritor” în ceea ce privește „modul de raționament” al modelului, explicațiile oferite confirmând acest lucru:



³ <https://learn.microsoft.com/en-us/ai/playbook/technology-guidance/generative-ai/working-with-llms/fine-tuning>

Așadar, procesul de raționament este o reflectare a exemplelor și a operațiilor existente în memorie, mai degrabă decât o adevărată deducție logică, aspecte de altfel menționate în multe studii referitoare la capacitatea de raționament logic a acestor modele, studii care arată că, în anumite situații, modelele au învățat să exploateze mai degrabă caracteristicile statistice decât să emuleze funcția corectă de raționament (Kambhampati, 2024; Mirzadeh et al., 2024; Zhang et al., 2024).

Modelele acestea sunt foarte utile în cazul unor sarcini de procesare a limbajului, cum ar fi generarea unor texte, parafrizarea, scrierea unui rezumat etc., dar nu au acces la cunoștințe în timp real, așadar nu pot furniza informații actualizate despre evenimente curente, nu pot răspunde cu precizie la întrebări legate de date recente sau despre situații care necesită apelarea unor surse externe, de exemplu, verificarea știrilor sau consultarea unei baze de date în timp real. În aceste situații există riscul ca modelul să combine date din diferite surse sau să creeze referințe care nu corespund realității (inexacte sau inventate). O asemenea solicitare arată că modelele acestea mari folosesc *probabilitatea* pentru a genera un text (cuvânt cu cuvânt), text care poate fi fluent, relevant, dar care nu este neapărat și corect.

În exemplul de mai jos, informațiile despre *George Topârceanu* oferite de ChatGPT sunt „amestecate”, unele inexacte, altele inventate:

Desigur! George Topârceanu (1886–1937) a fost un poet, prozator și publicist român, unul dintre cei mai importanți autori ai literaturii române din perioada interbelică. Este cunoscut mai ales pentru creațiile sale de inspirație populară și pentru abordarea umoristică și satirică a unor teme sociale și politice. Iată câteva aspecte importante despre viața și opera lui George Topârceanu:

1. Viața și Formarea

- George Topârceanu s-a născut pe 20 martie 1886, în comuna Topârceanu, județul Argeș (de aici provine și numele său). A urmat cursurile Școlii Normale din Pitești, iar mai târziu a studiat la București, unde a intrat în contact cu cercurile literare din capitală.
- A fost un om cu o educație teimeinică, dar nu a avut un parcurs formal în domeniul literaturii, fiind mai degrabă autodidact în privința operei sale.

2. Cariera Literară

- A debutat în jurnalism, lucrând ca redactor la mai multe reviste, dar faima sa a venit din creațiile literare.
- Este cel mai bine cunoscut pentru poeziile sale, care se caracterizează prin umor și un limbaj accesibil, cu un ton adesea jovial. Printre cele mai celebre lucrări se numără „*Balada unui greier mic*” și „*Testament*”.
- „*Balada unui greier mic*” este poate cel mai celebru poem al său, în care poetul se joacă cu contrastul dintre natura simplă și agitată a unui greier și lumea complicată a oamenilor. Poezia este o fantezie veselă, dar cu un substrat filozofic despre viață și soartă.
- A avut o mare influență asupra literaturii române interbelice, fiind un autor foarte apreciat de publicul larg.

4. Opere Semnificative

- „*Balada unui greier mic*” – un poem care a devenit emblematic pentru George Topârceanu, în care o mică insectă filosofează despre viață și despre condiția umană.
- „*Testament*” – o poezie în care autorul își exprimă gândurile despre moarte și despre cum ar vrea să fie amintit.
- „*Căpăciunul*” – o lucrare plină de umor în care Topârceanu explorează teme legate de viața rurală și de personajele mitologice.
- „*Ochiul de sticlă*” – un alt volum important în care sunt adunate poezii cu un puternic caracter satiric.
- „*Memorii din vremea mea*” – o lucrare autobiografică în care Topârceanu relatează evenimente din viața sa personală și profesională.

Cantitatea de informații, cunoștințe despre lumea înconjurătoare cu care o persoană vine zilnic în contact este extrem de mare, iar aceste tipuri de cunoștințe subiective, intuitive, experiențiale, dificil de articulat într-un sistem formal, reprezintă cadrul referențial și contribuie esențial la modul în care aceasta prelucrează informațiile, rezolvă anumite probleme sau ia anumite decizii. Modelele nu „înțeleg” conceptele (cel puțin nu așa cum le înțelegem noi) și, prin urmare: „that current LLMs are not capable of genuine logical reasoning; instead, they attempt to replicate the reasoning steps observed in their training data” (Mirzadeh et al.,

2024:1). De aceea, cu toate că progresele înregistrate de IA generativă nu pot fi negate, chiar în cazul unor activități mai complexe ce presupun un proces de reflecție, există încă sarcini, cum ar fi raționamentul logico-matematic, de exemplu, în care posibilitățile LLM-urilor sunt limitate⁴.

Mai jos, un alt exemplu care indică unele limitări în procesarea numerică sau în efectuarea unor calcule elementare. Atunci când li se solicită să determine ocurența unor litere dintr-un cuvânt, respectiv frecvența apariției unui lexem într-un context, aceste sisteme generează rezultate cu grade variate de acuratețe, de la parțial corecte până la complet eronate.



Erorile identificate sunt atribuite metodologiei diferite de segmentare a textului în unități lexicale minimale (numite tokeni) și care poate fi la nivel de caracter, cuvânt sau chiar propoziții, în funcție de caracteristicile limbii sau ale textului. Anumite modele LLM pot interpreta un cuvânt într-un context mai larg, ceea ce înseamnă că pot considera variante ale cuvântului (de exemplu, forme flexionate sau conjugate) ca fiind același cuvânt (parafraza/parafrază). Cum ele nu sunt proiectate a fi programe de calcul, pot greși, relativ ușor, atunci când sunt solicitate să determine ocurențele unui cuvânt într-un text.

Sistemele sunt, de fapt, instruite să găsească tipare în limbaj, relații semantice sau indicii contextuale, fiind capabile să identifice și să „potrivească” cele mai apropiate secvențe de raționament pe care le-au găsit în datele de antrenare. De altfel, LLM-urile (începând cu arhitectura GPT-3) sunt printre cele mai folosite modele în cercetările privind generarea sau înțelegerea unei limbi, mai cu seamă în cazul limbilor noi (cu puține resurse): „Now, it is possible to train highly coherent foundation models with a simple language generation objective, like predict the next word in this sentence” (Bommasani, 2022:23), deoarece chiar dacă pare o sarcină

⁴ În raportul GPT-4 Technical (<https://cdn.openai.com/papers/gpt-4.pdf>) se precizează că modelul face, uneori, erori simple de raționament care nu par a fi în concordanță cu competența sa în atât de multe domenii; în alte situații poate fi excesiv de credul în acceptarea declarațiilor evident false ale unui utilizator.

relativ simplă, modelele învață mult mai mult decât simple reguli gramaticale; ele pot reține și învăța relații complexe între cuvinte, care pot ajuta la înțelegerea structurii unei limbi.

LLM - O NOUĂ ERĂ ÎN TEHNOLOGIA LINGVISTICĂ. CE SE ÎNTÂMPLĂ ÎN CAZUL TRADUCERII?

În pofida multiplelor probleme existente încă, modelele LLM au deschis o nouă eră în tehnologia lingvistică, cu consecințe semnificative și în traducerea automată. LLM-urile sunt cunoscute, în primul rând, pentru generarea textului și există studii (Brown, 2020; Clark et al., 2021) care arată că încă de la varianta ChatGPT-3, în cazul unor texte mai scurte, este greu să se facă diferența între un text generat de ChatGPT și unul scris de om. Ele pot fi folosite și în traducerea automată, deoarece pot fi aplicate unui set mai larg de limbi, putând situații lingvistice realiste și complexe referitoare la transfer.

Traducerea neuronală și modelele lingvistice mari

Traducerea automată neuronală (NMT) a progresat foarte mult în ultimul timp, adoptând mai multe abordări cum ar fi modelele recurente, convoluționale sau hibride (Dakwale, Monz, 2017), dar și mecanisme de atenție sau modele bazate pe transformator, ceea ce permite învățarea directă din datele de traducere, într-un mod care „captează” mai bine contextul și relațiile între cuvinte, propoziții și fraze (așa-numiții vectori semantici) (Mondal et al., 2023). Adoptarea și introducerea modelelor de tip *Transformer* a însemnat un mare salt în calitatea traducerii. Astfel, folosirea mecanismului de atenție permite modelului să se concentreze asupra anumitor părți ale secvenței de intrare, să identifice mult mai bine relațiile contextuale și să producă traduceri mai precise (Junczys-Dowmunt, 2019; Yang et al., 2020; Dabre et al., 2021). Datorită capacității de a învăța continuu noi reprezentări ale limbajului (*învățare secvențială*), modelele NMT sunt capabile să producă texte mai fluente, firești, fără să depindă prea mult de structuri rigide sau de reguli gramaticale aplicate automat. În felul acesta își reduc semnificativ erorile, cum ar fi traducerea incorectă a cuvintelor polisemantice sau a frazelor complexe, frecvente în modelele anterioare. O altă îmbunătățire notabilă constă în faptul că modelele de această factură pot analiza contexte mai largi (*context global, cross-context*) (Wang et al., 2017) ceea ce le permite să înțeleagă și să traducă texte având în vedere nu doar cuvintele ca atare, ci și contextul mai larg oferit de propoziții și paragrafe. Modelele mai avansate pot chiar „memora” contextul global pe măsură ce traduc și, în consecință, pot utiliza informații din întregul text tradus anterior pentru a ajuta la traducerea frazelor ulterioare, chiar dacă acestea sunt traduse pe rând. Astfel, ele sunt capabile să mențină coerența și să folosească mai bine

referințele care apar pe parcurs în text, ceea ce determină o transpunere mai corectă a expresiilor ambigue sau a cuvintelor cu mai multe sensuri.

Cu toate acestea, modelele NMT întâmpină încă dificultăți când sunt folosite limbi rare, limbi cu resurse limitate de date sau când trebuie să redea concepte și nuanțe culturale complexe. De altminteri, antrenarea lor folosind corpusuri bilingve sau multilingve (perechi de limbi) le fac mai puțin flexibile și generează o serie de limitări. Calitatea traducerii variază semnificativ din cauză că limbile nu sunt reprezentate la fel, traducerea fiind dependentă de cantitatea de resurse disponibile pentru instruirea sistemelor. În plus, datele de instruire ar trebui să fie din aceeași arie specializată/semantică (dacă motorul se antrenează utilizând date din rețelele sociale, el nu va traduce bine un document din domeniul financiar).

Există și modele „specializate” a căror acuratețe în traducere este superioară (de ex. *BeringAI for legal translation*⁵, specializat în traducerea textelor din domeniul juridic). Conform site-ului, „Bering AI is 2-6x more accurate than our competition” (vezi imaginea de mai jos).



Să semnalăm, totuși, faptul deloc neglijabil că antrenarea se face folosind computere extrem de puternice, disponibile doar în cazul marilor companii. Așadar, îmbunătățirea performanței înseamnă un cost ridicat al capacității și al complexității instrumentelor de instruire în continuă creștere.

Spre deosebire de aceste modele specifice traducerii automate, LLM-urile au apărut ca modele generale de limbaj capabile să abordeze multe sarcini de generare a limbajului natural, de înțelegere a lui și, *printre ele, de traducere*. Ele se antrenează, în principal, pe corpusuri monolingve și, ca urmare, o mare parte a LLM-urilor sunt limitate în „abilitățile lor multilingve” (cu probleme privind inclusiv punctuația și diferiți alți indicatori speciali) (Hendy et al., 2023). Dar LLM-urile recente au fost îmbunătățite foarte mult în ceea ce privește traducerea deoarece, antrenate fiind pe seturi de date foarte mari, pot genera în prezent propoziții complexe, cu un context larg de situații de limbaj sau de limbi. Unele studii arată că,

⁵ <https://beringlab.com/>

în anumite situații concrete, acestea traduc mai bine decât modelele neuronale, adică produc traduceri mai puțin literale (Raunak et al., 2023; Stap et al., 2024), uneori mai fluente și creative, chiar dacă adesea imprecise.

O scurtă exemplificare în imaginile de mai jos, arată că traducerea oferită de *ChatGPT* este, per ansamblu, mai bună decât cea oferită de *DeepL*, în care, în plus, apare și o eroare, cuvântul ”superbad” fiind tradus cu „superbă”.

Traducerea oferită de ChatGPT

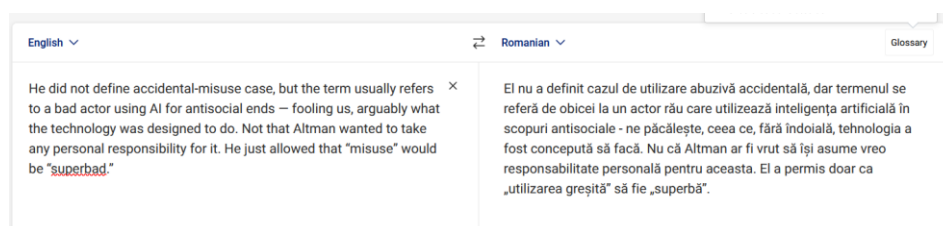
Sigur! Te rog să îmi dai propoziția pe care vrei să o traduc.

"He did not define accidental-misuse case, but the term usually refers to a bad actor using AI for antisocial ends — fooling us, arguably what the technology was designed to do. Not that Altman wanted to take any personal responsibility for it. He just allowed that "misuse" would be "superbad."

Iată traducerea în română a propoziției:

„Nu a definit cazul de utilizare accidentală greșită, dar termenul se referă de obicei la o persoană rău intenționată care folosește IA în scopuri antisociale — înșelându-ne, ceea ce, în mod discutabil, este ceea ce tehnologia a fost concepută să facă. Nu că Altman ar fi vrut să își asume vreo responsabilitate personală pentru asta. A recunoscut doar că „utilizarea greșită” ar fi „superproastă”.”

Traducerea oferită de DeepL



Chiar dacă LLM-urile nu sunt „specializate în traducere”, aceasta fiind văzută doar ca una dintre multiplele sarcini pe care ele le pot realiza, ele au o flexibilitate mai mare în a adapta traducerea la stiluri și contexte variate, pot fi ajustate pentru sarcini specifice, inclusiv traducerea între diferite limbi, cu un potențial superior de a învăța și de a se plia la noile tendințe lingvistice sau la un anumit jargon. Există multe modele care, deși nu au fost proiectate pentru traducere, funcționează cu rezultate semnificative și ar merita să fie încercate (chiar și complementar), iar dintre cele mai cunoscute menționăm: Facebook’s M2M-100 (primul model de IA care traduce 100 de limbi fără a folosi limba engleză ca pivot)⁶ și, bineînțeles, modelele GPT de la OpenAI (3, 4o, 4o-mini)⁷.

În ciuda performanței impresionante de învățare în context, adaptarea modelelor LLM la traducere rămâne încă o problemă (Hendy et al., 2023), printre

⁶ <https://about.fb.com/news/2020/10/first-multilingual-machine-translation-model/>

⁷ <https://platform.openai.com/docs/models>

limitări numărându-se: răspunsurile inconsecvente în raport cu datele (cunoștințele), așa-numitele halucinații; informațiile eronate sau complet false (cum am văzut în exemplele anterioare) și, mai ales, elementele de infidelitate în raport cu textul sursă, căruia îi pot adăuga și fraze inexistente.

Prezentăm în continuare un tip de halucinație în care modelul a adăugat informații (inexistente în textul sursă), *halucinații din ce în ce mai rare, însă*, pentru că modelele și-au perfecționat capacitatea de autocorectare, îmbunătățindu-și răspunsul în funcție de prompt (interogare).

can you translate this sentence from English to Romanian? "Many foundational AI models are labelled as 'open source' because their architecture, including the neural networks' structure and design, is made freely available. Yet, little information is disclosed about how the models were trained. As the executive director of the Open Source Initiative (OSI) based in Palo Alto, California, my priority since 2022 has been clarifying what the term actually means in the AI era."

Delete

Majoritatea modelelor fundamentale de inteligență artificială sunt etichetate ca 'surse deschise' deoarece arhitectura lor, incluzând structura și proiectarea rețelelor neuronale, este pusă la dispoziție în mod gratuit și este utilizată pe scară largă în industrie. Cu toate acestea, foarte puține informații sunt divulgate despre modul în care au fost antrenate modelele, ceea ce este un mare dezavantaj pentru dezvoltarea inteligenței artificiale. Ca director executiv al Inițiativei Surselor Deschise (OSI) cu

După cum se poate observa, în această traducere oferită de Free ChatGPT modelul a adăugat: „și este utilizată pe scară largă în industrie”, completare la prima propoziție, care poate fi coerentă și logică, *dar care nu se regăsește în textul sursă*. La a doua propoziție a adăugat un alt segment de text „ceea ce este un mare dezavantaj pentru dezvoltarea inteligenței artificiale”, adaos ce poate schimba sensul textului. De aceea, posteditarea sau revizia umană atentă rămâne în continuare o activitate obligatorie, cu atât mai mult cu cât textul generat suplimentar este coerent și nu împieteează asupra fluidității textului în general. Sunt studii care arată că „cu cât chatbotii devin mai buni, cu atât este mai probabil ca oamenii să rateze o eroare atunci când ea apare”.

Așa cum menționam, modelele LLM sunt extrem de costisitoare în ceea ce privește resursele pentru antrenare și rulare, dar pot, totodată, să fie alimentate și să folosească date care au orientări inadecvate sau care exprimă prejudecăți sociale și culturale, ceea ce poate conduce la rezultate distorsionate. Cum tind să fie opace în modul lor de funcționare (datele pentru antrenarea modelelor și transparența modului în care sunt colectate și utilizate nu sunt publice), „we currently lack a clear understanding of how they work, when they fail, and what they are even capable of due to their emergent properties” (Bommasani et al, 2022:1).

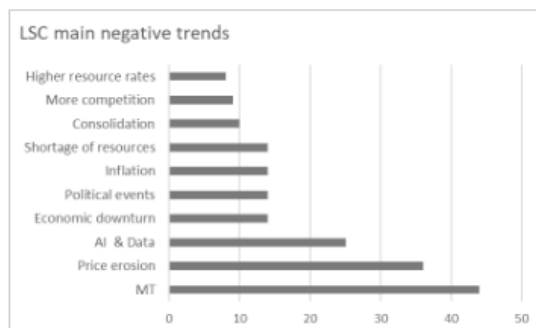
Tehnologia LLM-urilor este relativ accesibilă, iar marile companii se concentrează pe crearea unor modele personalizate care să îmbunătățească acuratețea și fiabilitatea lor; de aceea ele sunt dispuse să investească mult în adaptarea modelelor și în creșterea capacității lor de a înțelege contexte complexe, chiar dacă procesul este extrem de costisitor.

PROVOCĂRI ÎN INDUSTRIA LINGVISTICĂ

Odată cu introducerea algoritmilor de inteligență artificială și adaptarea lor pentru diferite procese lingvistice, ceea ce generic denumim industriile limbii se confruntă cu noi provocări; anul 2024 a fost un an de cumpănă pentru foarte multe companii, dar și pentru traducătorii independenți, care au înregistrat o scădere semnificativă a activităților de traducere.

„Over 50% of Freelance Linguists Have Thought About Changing Career, Slator Survey Finds” este rezultatul sondajului realizat de Slator⁸ în care este evidențiat faptul că, începând cu 2023, cererea pentru serviciile lingvistice tradiționale a scăzut semnificativ, scădere pe care traducătorii o atribuie introducerii tehnologiilor de IA. Această tendință va continua, după părerea lor, iar datele arată că „one in five freelance translators and interpreters have actively applied for new jobs during the same period”.

Un alt raport, cel al ELIS 2023⁹, relevă un nivel crescut de disconfort și de nesiguranță în rândul traducătorilor și interpreților (*in-house* sau *freelance*), iar multe LSCs-uri (*Language Service Company*), dintre care 45% reprezintă mici agenții de traduceri, au raportat scăderi ale veniturilor în 2023. După cum se poate vedea din imaginea următoare, printre provocările pieței în 2023, cele care domină clasamentul sunt traducerea automată și tarifele practicate.



sursa: Raport ELIS (2023:21).

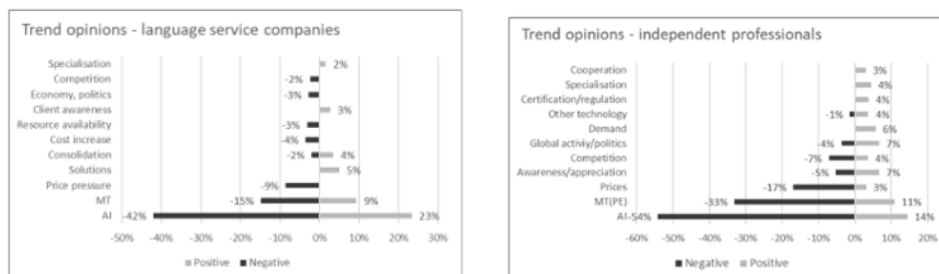
În 2024, conform Slator, sentimentul de nesiguranță se menține, această vulnerabilitate fiind reflectată și în Raportul ELIS 2024¹⁰ care arată că 87% din companii simt presiunea necesității de a implementa soluții bazate pe IA pentru a-și păstra avantajul. Și, cu toate că peste 54% sunt îngrijorate de acuratețea și fiabilitatea rezultatelor (vezi imaginile de mai jos), cele mai multe, sub influența pieței, au ales

⁸ <https://slator.com/50-of-freelance-linguists-have-thought-about-changing-career-slator-survey-finds/>

⁹ European Language Industry Survey 2023, <https://elis-survey.org/wp-content/uploads/2023/03/ELIS-2023-report.pdf>

¹⁰ <https://fit-europe-rc.org/wp-content/uploads/2024/09/ELIS-2024-report.pdf>

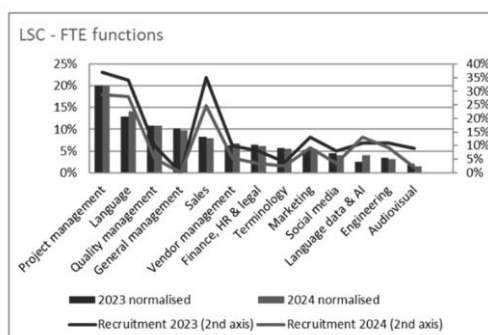
adoptarea IA și a modelor lingvistice mari, respectiv integrarea lor în fluxurile de lucru, pentru optimizarea resurselor, creșterea eficienței și a productivității.



Sursa: Raport ELIS (2024: 23, 24).

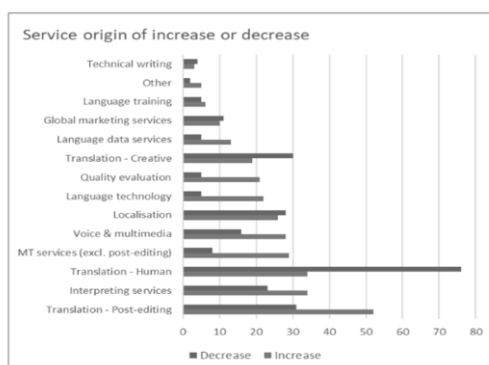
Cum era de așteptat, impactul IA, al traducerii automate și al posteditării, corelat cu prețurile practicate este perceput ca o tendință negativă, având consecințe directe asupra aspectului financiar, atât al companiilor, cât și al traducătorilor liber-profesioniști. Se știe că posteditarea, pe lângă faptul că este considerată de unii traducători ca fiind „mai puțin gratifiantă”, este plătită la tarife substanțial mai mici.

În acest context, pe lângă declinul înregistrat de piața traducerilor în sensul diminuării profitabilității și al reducerii tarifelor practicate, remarcăm și o evidentă tendință negativă în materie de recrutare în 2024 (comparativ cu 2023), conform Raportului ELIS din același an; de asemenea, după cum rezultă din graficul de mai jos, asistăm la o ușoară plafonare sau chiar la scăderea numărului de angajări și în alte câteva sectoare (management de proiecte, vânzări, finanțe, inginerie, audiovizual). Singurele segmente unde se pot observa creșteri sunt în direcția expertizei lingvistice (Language) și a introducerii IA în prelucrarea datelor (Language data & AI).



Sursa: Raport ELIS (2024: 43).

O schimbare, o anumită devalorizare a activității de traducere umană, clasică, este evidentă și în rezultatele studiului realizat de ELIS în 2024, în care se poate observa că, dintre toate activitățile, traducerea umană are cea mai mare rată de scădere, la polul opus aflându-se posteditarea cu cea mai mare creștere.

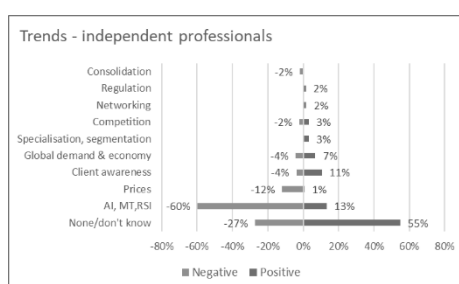
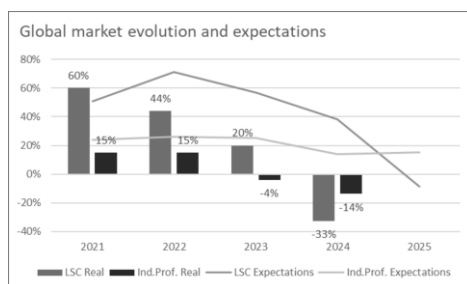


Sursa: Raport ELIS (2024:14).

Utilizarea traducerii automate pe scară largă prezintă un risc real pentru factorul uman; clienții nu mai sunt dispuși să plătească traducători știind că pot apela la soluții mai rapide și mai ieftine. Scăderea tarifelor practicate poate afecta calitatea produsului final, dar, în timp, impactul se va simți și asupra statutului profesional al traducătorului, care nu va mai fi perceput ca un expert lingvist, ci mai degrabă ca o persoană care revizuiște o traducere automată sau asistă IA în traducere. Iar rezultatele raportului ELIS 2025 referitoare la situația forței de muncă din industriile limbii se înscriu în această direcție: „Both language companies and language departments (with the exception of the international public agencies) report a considerable drop in staffing levels and express little hope that this will be undone in 2025” (p. 39).

Cum se conturează profilul tehnic al traducătorului?

Noile tehnologii și servicii care folosesc inteligența artificială și modul în care acestea ar putea fi gestionate mai eficient în domeniul traducerii rămân în continuare principalele preocupări și în 2025, marcând schimbări structurale importante în industria lingvistică, reflectate tot în sondajul efectuat de ELIS privind evoluția și așteptările pentru anul 2025¹¹, în care este de remarcat tendința mai degrabă negativă în raport cu IA exprimată mai ales în cazul liber-profesioniștilor.



Sursa: Raport ELIS (2025: 13, 21).

¹¹ https://elis-survey.org/wp-content/uploads/2025/03/ELIS-2025_Report.pdf

Profilul locurilor de muncă remodelează și profilul traducătorului, iar acesta se transformă semnificativ în contextul introducerii tehnologiilor bazate pe inteligența artificială, în special a modelelor lingvistice mari, dar și a altor instrumente automatizate.

Or, deși tot mai multe etape din procesul de traducere sunt automatizate sau, mai nou, asistate de IA, traducătorii continuă să joace un rol esențial în garantarea calității (slogan practicat de toate companiile!), a preciziei și a adaptării la contextul specific al fiecărei lucrări. Nu e mai puțin adevărat că, peste tot, este scos în evidență și faptul că expertiza lingvistică și cea culturală trebuie însoțite de *reale competențe tehnice de inginerie și prelucrare lingvistică*. Viitorul traducerii nu este doar tehnologic, ci reprezintă o simbioză între capacitatea IA și expertiza umană specializată, mai ales pentru limbile cu particularități lingvistice și culturale mai puțin reprezentate în seturile de date folosite pentru antrenarea sistemelor IA. Așadar, departe de a diminua importanța traducătorului, implementarea IA în industria traducerilor ar trebui să conducă la o redimensionare a rolului său, la o „hibridizare” a activității în vederea exploatării inteligente a tehnologiei actuale. Iar familiarizarea cu principiile de funcționare a sistemelor de inteligență artificială și înțelegerea modului în care acestea prelucrează limbajul natural devin esențiale în înțelegerea diferențelor dintre traducerea automată și cea umană. În ansamblu, la fel ca în orice domeniu, IA în traducere remodelează rapid întreaga arie de cunoștințe și de competențe dar, în același timp, creează noi nișe pentru traducătorii umani; prin urmare colaborarea cu IA poate fi un sprijin real, intervenind acolo unde traducerea automată nu reușește să atingă standardele necesare.

CONCLUZII

Sub influența tehnologiilor și mai ales a inteligenței artificiale, industriile limbii sunt supuse unor presiuni de transformare, adoptarea și mai ales *utilizarea fără discernământ a IA generative* putând afecta profund piața, prețurile, numărul angajaților, profesia în sine. Pe lângă folosirea programelor de traducere asistată sau automată, traducătorii trebuie să își adapteze competențele pentru a putea colabora eficient cu IA, aceasta devenind un asistent important în procesul traductiv. Și, oarecum paradoxal, în condițiile în care modelele își îmbunătățesc constant abilitățile de traducere, competențele lingvistice și interculturale ale traducătorului trebuie să fie mai elaborate ca niciodată. Capacitatea de a analiza un text în profunzime, de a-i identifica subtilitățile, de a adapta registrul lingvistic la publicul țintă și de a transpune nuanțele culturale reprezintă valoarea adăugată pe care traducătorul uman o aduce în procesul de traducere sau de comunicare interlingvistică.

Bibliografie

- Bender, E. M., Gebru, T., McMillan-Major, A., Shmitchell, S., 2021, „On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?” In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)*, Association for Computing Machinery, New York, pp. 610–623. <https://doi.org/10.1145/3442188.3445922>
- Bommasani, R., et al., 2022, „On the Opportunities and Risks of Foundation Models”, <https://arxiv.org/abs/2108.07258>
- Brown, T.B., Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei, 2020, „Language Models are Few-Shot Learners”, arXiv preprint arXiv:2005.14165
- Clark, E., August, T., Serrano, S., Haduong, N., Gururangan, S., Smith, N.A., 2021, „All That’s ‘Human’ Is Not Gold: Evaluating Human Evaluation of Generated Text.”, arXiv preprint arXiv:2107.00061
- Dabre, R., Chu, C., Kunchukuttan, A., 2021, „A Survey of Multilingual Neural Machine Translation”, *ACM Computing Surveys*, vol. 53, no. 5, (Article 99/1–38) (September 2021), <https://doi.org/10.1145/3406095>
- Dakwale, P., Monz, C., 2017, „Convolutional over Recurrent Encoder for Neural Machine Translation”, *The Prague Bulletin of Mathematical Linguistics*, no. 108, 2017, pp. 37–48, doi: 10.1515/pralin-2017-0007
- Devlin, J., Chang, M-W., Lee, K., Toutanova, K., 2019, „BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding”, <https://arxiv.org/pdf/1810.04805>
- Goodfellow, I., Bengio, Y., Courville, A., 2016, *Deep Learning*, The MIT Press, <https://www.deeplearningbook.org/>
- Gu, J., Hassan, H., Devlin, J., Victor O.K. Li, V. O.K., 2018, „Universal Neural Machine Translation for Extremely Low Resource Languages”, *Proceedings of NAACL-HLT 2018*, New Orleans, Louisiana, June 1- 6, 2018, Association for Computational Linguistics, pp. 344–354, ArXiv, abs/1802.05368
- Hendy, A., Abdelrehim, M.G., Sharaf, A., Raunak, V., Gabr, M., Matsushita, H., Kim, Y., Afify, M., & Awadalla, H.H., 2023, „How Good Are GPT Models at Machine Translation? A Comprehensive Evaluation.” ArXiv, abs/2302.09210
- Jiao, W., Wang, W., Huang, J.-t., Wang, X., Shi, S., Tu, Z., 2023, „Is ChatGPT A Good Translator? Yes With GPT-4 As The Engine”, <https://arxiv.org/pdf/2301.08745>
- Junczys-Dowmunt, M., 2019, „Microsoft Translator at WMT 2019: Towards Large-Scale Document-Level Neural Machine Translation”, *Proceedings of the Fourth Conference on Machine Translation (WMT)*, vol. 2: Shared Task Papers (Day 1) pp. 225–233, Florence, Italy, August 1-2, Association for Computational Linguistics, <https://aclanthology.org/W19-5321.pdf>
- Kambhampati, S., 2024, „Can large language models reason and plan?”, *Annals of the New York Academy of Sciences*, vol. 1534/1, pp. 15–18, <https://doi.org/10.1111/nyas.15125>
- Mirzadeh, I., Alizadeh, K., Shahrokhi, H., Tuzel, O., Bengio, S., Farajtabar, M., 2024, „Gsm-symbolic: Understanding the limitations of mathematical reasoning in large language models.”, arXiv preprint arXiv:2410.05229
- Mondal, S.K., Zhang, H., Kabir, H.M., Ni, K., Dai, H., 2023, „Machine translation and its evaluation: a study.” *Artificial Intelligence Review*, vol. 56, 2023, pp. 10137–10226.
- Nordhoff, S., Hammarström, H., 2011, „Glottolog/Langdoc: Defining dialects, languages, and language families as collections of resources”. *Proceedings of the First International Workshop on Linked Science 2011 (LISC2011)*, Bonn, Germany, October 24, 2011, <https://ceur-ws.org/Vol-783/paper7.pdf>
- Raunak, V., Menezes, A., Post, M., Awadalla, H.H., 2023, „Do GPTs Produce Less Literal Translations?”, <https://doi.org/10.48550/arXiv.2305.16806>
- Stap, D., Hasler, E., Byrne, B., Monz, C., Tran, K., 2024, „The Fine-Tuning Paradox: Boosting Translation Quality Without Sacrificing LLM Abilities”, <https://arxiv.org/pdf/2405.20089>

- Yang, S., Wang, Y., Chu, X., 2020, „A Survey of Deep Learning Techniques for Neural Machine Translation”, arXiv:2002.07526
- Wang, L., Tu, Z., Way, A., Liu, Q., 2017, „Exploiting Cross-Sentence Context for Neural Machine Translation”, in *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, Association for Computational Linguistics, Copenhagen, Denmark, pages 2826–2831, <https://arxiv.org/pdf/1704.04347>
- Zhang, H., Li, H.L., Meng, T., Chang, K-W., Van den Broeck, G., 2023, „On the Paradox of Learning to Reason from Data”, *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence (IJCAI-23)*, <https://www.ijcai.org/proceedings/2023/0375.pdf>

Manuela MIHĂESCU is a lecturer within the Department of Applied Modern Languages of the Faculty of Letters, Babeş-Bolyai University, where she teaches Information and Communication Technology and Terminology. She holds a PhD in Linguistics (*Communication and Knowledge*) and was involved for several years in various Romanian and European research projects on language processing, terminology, and language teaching and learning. Her research interests mainly concern communication and information processing.