

Συνέδρια της Ελληνικής Επιστημονικής Ένωσης Τεχνολογιών Πληροφορίας & Επικοινωνιών στην Εκπαίδευση

(2024)

8ο Πανελλήνιο Επιστημονικό Συνέδριο «Ένταξη και Χρήση των ΤΠΕ στην Εκπαιδευτική Διαδικασία»



ΠΑΝΕΠΙΣΤΗΜΙΟ ΘΕΣΣΑΛΙΑΣ

ΕΛΛΗΝΙΚΗ ΕΠΙΣΤΗΜΟΝΙΚΗ ΕΝΩΣΗ ΤΕΧΝΟΛΟΓΙΩΝ ΠΛΗΡΟΦΟΡΙΑΣ & ΕΠΙΚΟΙΝΩΝΙΩΝ ΣΤΗΝ ΕΚΠΑΙΔΕΥΣΗ

8ο Πανελλήνιο Επιστημονικό Συνέδριο

Ένταξη και Χρήση των ΤΠΕ στην Εκπαιδευτική Διαδικασία

Βόλος, 27-29 Σεπτεμβρίου 2024

Διοργάνωση

Πανεπιστήμιο Θεσσαλίας

Παιδαγωγικό Τμήμα Ειδικής Αγωγής
Παιδαγωγικό Τμήμα Προσχολικής Εκπαίδευσης
Παιδαγωγικό Τμήμα Δημοτικής Εκπαίδευσης
Τμήμα Επιστήμης Φυσικής Αγωγής & Αθλητισμού

Ελληνική Επιστημονική Ένωση Τεχνολογιών Πληροφορίας & Επικοινωνιών στην Εκπαίδευση

Επιμέλεια

Χαράλαμπος Καραγιαννίδης
Ηλίας Καρασαββίδης
Βασίλης Κάλλας
Μαρίνα Παπαστεργίου

etpe2024.uth.gr

ISBN: 978-618-5866-00-6

Μεγάλα Γλωσσικά Μοντέλα: προβληματισμοί, προκλήσεις και στο βάθος η Ελληνική Γλώσσα

Αριστείδης Βαγγελάτος

Βιβλιογραφική αναφορά:

Βαγγελάτος Α. (2025). Μεγάλα Γλωσσικά Μοντέλα: προβληματισμοί, προκλήσεις και στο βάθος η Ελληνική Γλώσσα. *Συνέδρια της Ελληνικής Επιστημονικής Ένωσης Τεχνολογιών Πληροφορίας & Επικοινωνιών στην Εκπαίδευση*, 839-845. ανακτήθηκε από <https://eproceedings.epublishing.ekt.gr/index.php/cetpe/article/view/8500>

Μεγάλα Γλωσσικά Μοντέλα: προβληματισμοί, προκλήσεις και στο βάθος η Ελληνική Γλώσσα

Αριστείδης Βαγγελάτος
vagelat@cti.gr
Ερευνητής, ΙΤΥΕ «Διόφαντος»

Περίληψη

Η Τεχνητή Νοημοσύνη, και η «παραγωγική» εκδοχή της, που περιλαμβάνει κατά κύριο λόγο τα Μεγάλα Γλωσσικά Μοντέλα (ΜΓΜ) και τις εφαρμογές τους, έχουν αλλάξει μέσα σε μικρό χρονικό διάστημα, τον τρόπο που σκεφτόμαστε, ενεργούμε και κυρίως «γράφουμε» (παράγουμε γραπτό κείμενο). Και όταν αναφερόμαστε στην γραφή, η έμφαση, στο πλαίσιο της εργασίας, δίνεται σε ό,τι έχει σχέση με την εκπαιδευτική διαδικασία με επιπλέον στόχο την αξιολόγηση των νέων δεδομένων, σε σχέση με την Ελληνική γλώσσα. Στην παρούσα έρευνα, μέσω κυρίως βιβλιογραφικής ανασκόπησης, προσεγγίζουμε τα ΜΓΜ και τις ανησυχίες που έχουν δημιουργήσει κατά κύριο λόγο στην εκπαιδευτική κοινότητα, ενώ παράλληλα σχολιάζουμε ακροθιγώς, την «επάρκεια» των ΜΓΜ σε γλώσσες λιγότερο διαδεδομένες, όπως είναι τα Ελληνικά.

Λέξεις κλειδιά: Μεγάλα Γλωσσικά Μοντέλα, Επεξεργασία Φυσικής Γλώσσας, Τεχνητή Νοημοσύνη

Εισαγωγή

Τα τελευταία δυο χρόνια, παρατηρούμε μια τεράστια εξέλιξη σε αυτό που όλοι, ίσως για ευκολία, ονομάζουμε τεχνητή νοημοσύνη (ΤΝ), και απεικονίζεται στην καθημερινότητά μας κυρίως με εφαρμογές Μεγάλων Γλωσσικών Μοντέλων (ΜΓΜ) – Large Language Models (LLMs).

Κάθε χρήστης έχει πλέον δοκιμάσει τουλάχιστον μια σχετική εφαρμογή (π.χ. ChatGPT, Gemini, Copilot, LLaMA), με τους πιο θαλερούς να το αξιοποιούν στην καθημερινότητά τους: από τα πιο καλά παραδείγματα, οι «πρωτοπόροι» φοιτητές που γράφουν μια διατριβή, ή κάποιο επιστημονικό άρθρο, σε πολλές περιπτώσεις αναζητούν βοήθεια, στα συστήματα αυτά (Yan et al., 2024). Οι σχετικές εφαρμογές κάνουν κυριολεκτικά θραύση, σε κάτι που υπό άλλες συνθήκες, θα χαρακτηρίζαμε πολύ εύκολα από λογοκλοπή έως και αντιγραφή, ενώ σήμερα, δυσκολευόμαστε ακόμα και να το αναγνωρίσουμε: έχουν βλέπετε κάνει την εμφάνισή τους εφαρμογές «παραφράσης» που αναλαμβάνουν να κάνουν ένα κείμενο που έχει παραχθεί από εφαρμογή ΤΝ, μη αναγνωρίσιμο ως τέτοιο (από ... άλλες σχετικές εφαρμογές αναγνώρισης της λογοκλοπής – πλαγιαρισμού).

Μέσα σε αυτό το δυστοπικό (αλλά και καινοφανές) τοπίο, πολλές είναι οι παράμετροι που πρέπει να διερευνηθούν καλύτερα: να τις κατανοήσουμε με στόχο να τις αξιοποιήσουμε στην εκπαίδευση και όχι απλά να χρησιμοποιηθούν με σκοπό την εξαπάτηση. Και σίγουρα υπάρχει μεγάλο πεδίο έρευνας προς την κατεύθυνση αυτή.

Στο παρόν άρθρο, πέρα από τη συζήτηση για τις προκλήσεις αυτές, εξετάζουμε και ένα θέμα που έχει να κάνει με την δυνατότητα (ή την αδυναμία) που έχουν τα ΜΓΜ να υποστηρίξουν στον ίδιο βαθμό, το Αγγλόφωνο κοινό (η πλειοψηφία των ΜΓΜ έχουν εκπαιδευτεί σε μεγάλο βαθμό σε κείμενα – σώματα κειμένων – που είναι γραμμένα στην Αγγλική γλώσσα) και ταυτόχρονα το μη Αγγλόφωνο κοινό και ειδικότερα άλλες, λιγότερο χρησιμοποιούμενες γλώσσες, όπως είναι και η Ελληνική.

Στη συνέχεια της εργασίας, αρχικά κάνουμε μια εισαγωγή στα ΜΓΜ, κατόπιν παρουσιάζουμε τις ιδιαιτερότητες της Ελληνικής γλώσσας σε σχέση με την Επεξεργασία Φυσικής Γλώσσας, μετά σχολιάζουμε τις ανησυχίες και τις προκλήσεις που φέρνουν οι εφαρμογές των ΜΓΜ, και πριν τα συμπεράσματα, εξετάζουμε αν τα ΜΓΜ που έχουμε σήμερα, είναι το ίδιο «ικανά» για όλες τις φυσικές γλώσσες.

Τι είναι τα ΜΓΜ (LLM-Large Language Models)

Ένα μεγάλο γλωσσικό μοντέλο (LLM) είναι ένα πακέτο λογισμικού που διακρίνεται για την ικανότητά του να μπορεί να κατανοήσει και να παράγει γλώσσα γενικού σκοπού (Wikipedia, 2024). Τα ΜΓΜ αποκτούν αυτές τις «ικανότητες» αποθηκεύοντας («μαθαίνοντας» κατά άλλους) στατιστικές σχέσεις από μεγάλης έκτασης ψηφιακά κειμενικά έγγραφα μέσω μιας υπολογιστικά εντατικής διαδικασίας εκπαίδευσης είτε με αυτοεπιβλεψη (self-supervised) είτε με ημιεπιβλεψη (semi-supervised). Τα ΜΓΜ μπορούν να χρησιμοποιηθούν για την παραγωγή κειμένου, μια μορφή παραγωγικής τεχνητής νοημοσύνης, λαμβάνοντας ένα κείμενο εισόδου και παράγοντας κείμενο με νόημα, προβλέποντας επανειλημμένα το επόμενο σύμβολο, λέξη ή φράση.

Τα ΜΓΜ είναι χρησιμοποιούν στην πράξη τεχνητά νευρωνικά δίκτυα. Τα μεγαλύτερα και πιο ικανά, από τον Μάρτιο του 2023, κατασκευάζονται με αρχιτεκτονική που βασίζεται σε αποκωδικοποίηση βασισμένη σε μετασχηματισμούς (decoder-only transformer-based architecture).

Μερικά αξιοσημείωτα και ευρέως χρησιμοποιούμενα ΜΓΜ είναι η σειρά μοντέλων GPT της OpenAI (π.χ. GPT-3.5 και GPT-4, που χρησιμοποιούνται στο ChatGPT και στο Microsoft Copilot), το Gemini της Google, η οικογένεια μοντέλων LLaMA της Meta, τα μοντέλα Claude της Anthropic και τα μοντέλα της Mistral AI.

Πώς λειτουργούν τα μεγάλα γλωσσικά μοντέλα;

Τα ΜΓΜ ακολουθούν μια σύνθετη προσέγγιση που περιλαμβάνει πολλαπλά επίπεδα.

Στο θεμελιώδες επίπεδο, ένα ΜΓΜ πρέπει να εκπαιδευτεί σε έναν μεγάλο όγκο δεδομένων - μερικές φορές αναφέρεται ως σώμα δεδομένων - που συνήθως έχει μέγεθος petabytes. Η εκπαίδευση μπορεί να περιλαμβάνει πολλαπλά βήματα, ξεκινώντας συνήθως με μια προσέγγιση μάθησης χωρίς επιβλεψη (unsupervised learning). Στην προσέγγιση αυτή, το μοντέλο εκπαιδεύεται σε μη δομημένα δεδομένα και σε δεδομένα χωρίς ετικέτες. Το πλεονέκτημα της εκπαίδευσης σε μη επισημειωμένα δεδομένα (χωρίς ετικέτες) είναι ότι συχνά υπάρχουν πολύ περισσότερα δεδομένα διαθέσιμα σε αυτή τη μορφή. Σε αυτό το στάδιο, το μοντέλο αρχίζει να αντλεί σχέσεις μεταξύ διαφορετικών λέξεων και εννοιών (πάντα σε στατιστική μορφή).

Το επόμενο βήμα για ορισμένα ΜΓΜ είναι η εκπαίδευση και η τελειοποίηση με μια μορφή αυτοεπιβλεπόμενης (self-supervised) μάθησης. Στο βήμα αυτό, έχει πραγματοποιηθεί κάποια επισημείωση δεδομένων, βοηθώντας το μοντέλο να αναγνωρίζει με μεγαλύτερη ακρίβεια τις διαφορετικές έννοιες.

Στη συνέχεια, το ΜΓΜ εφαρμόζει βαθιά μάθηση (deep learning) καθώς περνάει από τη διαδικασία του μετασχηματισμού στα νευρωνικά δίκτυα. Η αρχιτεκτονική του μοντέλου μετασχηματιστή (transformer) επιτρέπει στο ΜΓΜ να κατανοεί και να αναγνωρίζει τις σχέσεις και τις συνδέσεις μεταξύ λέξεων και εννοιών χρησιμοποιώντας έναν μηχανισμό αυτο-ενημέρωσης ή αυτο-προσοχής (self-attention mechanism). Αυτός ο μηχανισμός είναι σε θέση να αποδίδει μια βαθμολογία, που συνήθως αναφέρεται ως βάρος, σε ένα δεδομένο στοιχείο - που ονομάζεται token - προκειμένου να προσδιορίσει τη σχέση.

Τα ΜΓΜ έχουν γίνει όλο και πιο δημοφιλή επειδή έχουν ευρεία εφαρμογή για ένα πλατύ φάσμα εργασιών ΕΦΓ, συμπεριλαμβανομένων των ακόλουθων: Παραγωγή κειμένου, Μετάφραση, Περίληψη περιεχομένου, Παραλλαγή περιεχομένου, Ταξινόμηση και κατηγοριοποίηση, Ανάλυση συναισθήματος, Διαλογική τεχνητή νοημοσύνη και προσωπικοί βοηθοί (chatbots).

Η Ελληνική γλώσσα και η Επεξεργασία Φυσικής Γλώσσας

Τα ελληνικά είναι η επίσημη γλώσσα της Ελλάδας, μία από τις δύο επίσημες γλώσσες της Κύπρου, καθώς και μία από τις 24 επίσημες γλώσσες της Ευρωπαϊκής Ένωσης (ΕΕ). Έχοντας τουλάχιστον 13 εκατομμύρια φυσικούς ομιλητές, όπως αναφέρεται στην έκδοση του 2020 (23η) του Ethnologue, μιας βάσης δεδομένων αναφοράς γλωσσών που δημοσιεύεται από την SIL International (Eberhard, Simons, & Fennig, 2019), η ελληνική κατατάσσεται στην 89η θέση μεταξύ των 200 πιο διαδεδομένων γλωσσών παγκοσμίως. Το ελληνικό αλφάβητο αποτελείται από 24 γράμματα και δύο διακριτικά, ένα για την αναπαράσταση του σημείου τονισμού (π.χ. «ί») και ένα για τα διαλυτικά, που αποτελείται από δύο τελείες πάνω από ένα φωνήεν και υποδηλώνει μια ξεχωριστή συλλαβή (π.χ. «ϊ»).

Η ΕΦΓ ήδη από τα τέλη της δεκαετίας του 1980 αναγνωρίστηκε ως μια αποφασιστική ευκαιρία για την ελληνική γλώσσα να συμμετάσχει στο γλωσσικό τοπίο της Ευρωπαϊκής Ένωσης επί ίσοις όροις με άλλες πιο διαδεδομένες γλώσσες. Υπό το πρίσμα αυτής της πολιτικής και με τη βοήθεια της Ευρωπαϊκής Ένωσης και εθνικών χρηματοδοτήσεων, έχουν αναπτυχθεί πολλαπλά εργαλεία και πόροι για την «ελληνική» ΕΦΓ. Ωστόσο, οι σχετικές εργασίες που διερευνούν την περιοχή είναι περιορισμένες (Gavrilidou et al., 2012; Papantoniou & Tzitzikas, 2020). Σύμφωνα με την πιο πρόσφατη (Papantoniou & Tzitzikas, 2020), αναφέρονται 79 δημοσιευμένες εργασίες που αντλήθηκαν τόσο από την ερευνητική όσο και από την “γκρίζα” βιβλιογραφία και αφορούν πόρους και εργαλεία ΕΦΓ για τα νέα ελληνικά. Οι εργασίες αυτές περιγράφουν εξελίξεις για διάφορα επίπεδα επεξεργασίας, όπως φωνητική, μορφολογία, σύνταξη, ενσωμάτωση λέξεων, σημασιολογία, ανάλυση συναισθήματος και απάντηση ερωτήσεων. Όπως γίνεται φανερό από αυτή τη συνοπτική επισκόπηση, ένας μεγάλος αριθμός από τις αναφερόμενες εργασίες, αν και καινοτόμες όσον αφορά τις ερευνητικές προσεγγίσεις, δεν φαίνεται να έχουν ενσωματωθεί ή αξιοποιηθεί σε ένα ευρύτερο πλαίσιο εφαρμογής. Ακόμη πιο σημαντικό είναι ότι οι πόροι και τα εργαλεία που σχετίζονται με πιο εξειδικευμένες περιοχές (π.χ. Υγεία, κ.α.) μπορούν να μετρηθούν στα δάχτυλα του ενός χεριού (Vagelatos et al., 2011; Vagelatos et al., 2022; Zervopoulos et al., 2019).

Ανησυχίες που εκφράζονται για τις εφαρμογές των ΜΓΜ

Η έλευση της ΤΝ και πιο συγκεκριμένα της παραγωγικής ΤΝ, με τη μορφή που αναφέρθηκε παραπάνω, φέρνει μαζί της και αρκετές ανησυχίες που εκφράζονται από τους ερευνητές. Ας δούμε τις πιο σημαντικές από αυτές.

Τα ΜΓΜ μπορούν να χρησιμεύσουν ως προσωπικοί δάσκαλοι αγγλικών, υποστηρίζοντας τους ερευνητές που δεν έχουν ως μητρική γλώσσα την αγγλική να ξεπεράσουν τις δυσκολίες που συχνά αντιμετωπίζουν κατά τη συγγραφή επιστημονικών εργασιών στη γλώσσα αυτή (Koga, 2023).

Ωστόσο, υπάρχουν προβλήματα στην αξιοποίηση ΜΓΜ για τη δημιουργία κειμένων, όπως η παραγωγή ψευδούς πληροφορίας (hallucinations¹), η λογοκλοπή αλλά και οι ανησυχίες για την προστασία της ιδιωτικής ζωής. Για τον μετριασμό αυτών των τρωτών σημείων, οι συγγραφείς θα πρέπει να ελέγχουν πολλαπλώς το παραγόμενο περιεχόμενο συγκρίνοντάς το με αξιόπιστες πηγές για να διασφαλίσουν την ακρίβεια, να χρησιμοποιούν ανιχνευτές γλωσσικής ομοιότητας (εργαλεία κατά της λογοκλοπής) και να αποφεύγουν την εισαγωγή σημαντικών και προσωπικών πληροφοριών στις εντολές/ερωτήματα (prompt) των ΜΓΜ.

Επιπλέον υπάρχει η προτροπή τα ΜΓΜ να χρησιμοποιούνται περισσότερο για την επεξεργασία και την βελτίωση κειμένου παρά για τη δημιουργία μεγάλης έκτασης «πρωτότυπου» κειμένου (Elsevier, 2023).

Σύμφωνα με τα πρότυπα του Elsevier (Elsevier, 2023), οι συγγραφείς θα πρέπει να χρησιμοποιούν τεχνολογίες παραγωγικής τεχνητής νοημοσύνης (TN) και τεχνολογίες με υποστήριξη TN κατά τη διαδικασία συγγραφής, μόνο για τη βελτίωση της αναγνωσιμότητας και της γλώσσας της εργασίας τους. Καθώς η τεχνητή νοημοσύνη μπορεί να παρέχει ένα έγκυρο φαινομενικά αποτέλεσμα που στην ουσία του είναι εσφαλμένο, ελλιπές ή προκατειλημμένο, θα πρέπει να χρησιμοποιείται με ανθρώπινη εποπτεία και έλεγχο και οι συγγραφείς θα πρέπει να εξετάζουν προσεκτικά και να τροποποιούν κατάλληλα το αποτέλεσμα. Το περιεχόμενο του οιοδήποτε πονήματος, είναι τελικά υπευθυνότητα του δημιουργού και μόνο αυτός είναι υπόλογος για αυτό.

Συγκεκριμένα περιοδικά και οργανισμοί ανταποκρίνονται διαφορετικά στη χρήση της TN. Για παράδειγμα, το Korean Journal of Radiology (!) ενθαρρύνει την ορθή εφαρμογή της παραγωγικής TN για να διευκολύνει τη διάδοση σημαντικών επιστημονικών γνώσεων μέσω των δημοσιεύσεων, αποτρέποντας παράλληλα την επιστημονική ανάρμοστη συμπεριφορά και τις παραβιάσεις της δεοντολογίας των δημοσιεύσεων (Park, 2023). Στον αντίποδα, οι πολιτικές των περιοδικών της Radiological Society of North America (RSNA, 2023) σημειώνουν σαφώς ότι οι συγγραφείς θα πρέπει να είναι σε θέση να δηλώνουν ότι τα άρθρα τους δεν περιέχουν λογοκλοπή, συμπεριλαμβανομένου του κειμένου και των γραφικών που παράγονται από εφαρμογές TN.

Είναι όμως τα ΜΓΜ, το ίδιο ικανά σε όλες τις γλώσσες;

Η επιστημονική συζήτηση έχει ξεκινήσει σχεδόν με την εμφάνιση των εφαρμογών TN, σε σχέση με την «ικανότητα» των ΜΓΜ να ανταποκριθούν ανάλογα και με επάρκεια σε διαφορετικές φυσικές γλώσσες (το σημείο αναφοράς είναι φυσικά η Αγγλική που είναι η κυρίαρχη στο διαδίκτυο).

Οι Lai et al., 2023, αναφέρονται στο ChatGPT, περιγράφοντας ότι σε σύγκριση με ειδικές κειμενικές ερωτήσεις, η ανώτερη απόδοση του ChatGPT σε αγγλικές περιγραφές εργασιών για την πλειονότητα των προβλημάτων και των γλωσσών υποδηλώνει ότι το ChatGPT μπορεί να κατανοεί/αναλύει καλύτερα τις αγγλικές προτροπές και να οδηγεί σε καλύτερη παραγωγή απαντήσεων με πιο ακριβή αποτελέσματα. Επιπλέον, η συμπερίληψη αγγλικών περιγραφών εργασιών για μη αγγλικές εισόδους (προτροπές) μπορεί να θεωρηθεί ως μια προσέγγιση για τη μετατόπιση των αναπαραστάσεων των ειδικών γλωσσικών εισόδων προς τον αγγλικό χώρο που μπορεί να επεξεργαστεί καλύτερα από το ChatGPT λόγω της κυριαρχίας της αγγλικής γλώσσας στα δεδομένα εκπαίδευσής του.

¹ Βλέπε νέο ορισμό στο <https://www.cam.ac.uk/research/news/cambridge-dictionary-names-hallucinate-word-of-the-year-2023>

Ερχόμενοι τώρα σε παραδείγματα προτροπών στην Ελληνική γλώσσα, διαβάζουμε ότι στην προτροπή «Δώσε μου τρεις σημαντικές εθνικές εορτές» η απάντηση του ChatGPT (4) ήταν «1. Ημέρα της Ανεξαρτησίας, 2. Όχι Ημέρα [sic], 3. Πρωτομαγιά» (Σπινέλλης, 2024).

Λίγες εβδομάδες μετά, την ώρα που γράφεται το παρόν κείμενο, το ίδιο εργαλείο δίνει πολύ καλύτερη απάντηση (μαζί και με σχετικές αναφορά):

Σίγουρα! Εδώ είναι τρεις σημαντικές εθνικές εορτές στην Ελλάδα:

α) 25η Μαρτίου (Εορτασμός της Ελληνικής Επανάστασης του 1821). β) 1η Μαΐου (Εργατική Πρωτομαγιά). γ) 28η Οκτωβρίου (Επέτειος του ΟΧΙ).

Δεν μπορούμε να είμαστε σίγουροι για το πλαίσιο που δημιούργησε τόσο διαφορετικές απαντήσεις. Σίγουρα όμως θα πρέπει να έχει κανείς κατά νου δυο «χρυσούς κανόνες»: α) αυτό που όλοι γνωρίζαμε από την προηγούμενη «google search» περίοδο, ισχύει ακόμα: το «of course it's true, I show it on the Internet» (<https://cacm.acm.org/research/of-course-its-true-i-saw-it-on-the-internet/>) εξακολουθεί να ΜΗΝ ισχύει! Και β) σε μεμονωμένες απαντήσεις, μπορούμε να βασιστούμε μόνο αντιπαραδείγματα και ΟΧΙ συμπεράσματα! Και φυσικά, όπως αναφέρθηκε και παραπάνω, τα μοντέλα βελτιώνονται.

ΜΓΜ προσαρμοσμένα στην Ελληνική Γλώσσα

Εδώ και αρκετό καιρό, ξεκίνησαν προσπάθειες να επεκταθούν οι δυνατότητες των ανοιχτών ΜΓΜ σε άλλες γλώσσες (π.χ. LeoLM για γερμανικά, AguilaF για ισπανικά, κ.λπ.)

Ομοίως τους τελευταίους μήνες, έχει ξεκινήσει μια προσπάθεια και από Έλληνες ερευνητές (και τα αντίστοιχα ερευνητικά ινστιτούτα) να «εξελληνίσουν» (ή καλύτερα να εκπαιδεύσουν στην Ελληνική γλώσσα) ΜΓΜ. Προς αυτή την κατεύθυνση αναπτύχθηκε και κυκλοφόρησε από το Ινστιτούτο Επεξεργασίας Γλώσσας και Λόγου του Κέντρου Έρευνας & Καινοτομίας Αθηνά, το Meltemi LLM για την ελληνική γλώσσα. Το Meltemi αναπτύσσεται ως δίγλωσσο μοντέλο, διατηρώντας τις δυνατότητές του για την αγγλική γλώσσα, ενώ επεκτείνεται για να κατανοεί και να δημιουργεί άπαιστα κείμενο στα Νέα Ελληνικά χρησιμοποιώντας τεχνικές αιχμής. Το Meltemi είναι χτισμένο πάνω στο ΜΓΜ ανοιχτού λογισμικού της Mistral-7B και έχει εκπαιδευτεί σε ένα σώμα ελληνικών κειμένων υψηλής ποιότητας. Υπάρχουν δύο παραλλαγές του Meltemi στην έκδοση 1k: το θεμελιώδες μοντέλο Meltemi-7B-v1 και το παράγωγό του, Meltemi-7B-Instruct-v1 που μπορεί να χρησιμοποιηθεί για εφαρμογές συνομιλίας (Meltemi, 2024).

Στο (Αντονοπουλος, 2024), γίνεται μια προσπάθεια συγκριτικής αξιολόγησης Ελληνικών ΜΓΜ ειδικού (αλλά και γενικού) σκοπού. Στα συμπεράσματα διαβάζουμε ότι το ειδικού σκοπού, εκπαιδευμένο μοντέλο της 11tensors, αν και σημαντικά μικρότερο, έχει τις ίδιες επιδόσεις με τα μεγαλύτερα μοντέλα. Ωστόσο, δεν αποτελεί έκπληξη... Αυτό το μοντέλο είναι ρυθμισμένο για τη συγκεκριμένη εργασία. Αυτή είναι η ομορφιά της προσαρμογής και της χρήσης μικρότερων μοντέλων για εξειδικευμένες εργασίες.

Τέλος, για να γίνει αναφορά και στο αρχικό ερώτημα: Παρόλο που τα ελληνικά υποεκπροσωπούνται στα σύνολα δεδομένων εκπαίδευσης των ΜΓΜ, σε σύγκριση με τις ευρέως ομιλούμενες (και εκπροσωπούμενες) γλώσσες, τα μοντέλα εξακολουθούν να είναι πολύ ικανά και να επιδεικνύουν μια ισχυρή απόδοση. Επιπλέον, προσφέρουν την ευελιξία για περαιτέρω τελειοποίηση ώστε να προσαρμοστούν ειδικά στις μοναδικές και ειδικότερες περιπτώσεις χρήσης που μπορεί κατά περίπτωση να απαιτούνται.

Συμπερασματικά

Στην εργασία αυτή έγινε μια πρώτη προσπάθεια αποδελτίωσης επιστημονικής βιβλιογραφίας σε σχέση με τις προκλήσεις στη χρήση ΜΓΜ γενικά και ειδικότερα για γλώσσες λιγότερο

χρησιμοποιούμενες, όπως είναι τα Ελληνικά. Από τις αναφορές που καταγράφηκαν, είναι φανερό, ότι υπάρχει ένα ζήτημα πιστότητας γενικότερα, αλλά ειδικότερα σε γλώσσες με μικρή παρουσία στο διαδίκτυο, εγείρονται αρκετά ζητήματα.

Για να γίνει δυνατό όμως να προκύψει ένα τεκμηριωμένο συμπέρασμα, θα πρέπει να γίνει πιο συντεταγμένη έρευνα με βάση τα μοντέλα που χρησιμοποιούνται, έχοντας πάντα κατά νου, το γεγονός ότι ακόμα και αυτά, σε τακτά χρονικά διαστήματα, αλλάζουν (βελτιούμενα).

Μέχρι τότε, καλό είναι να μην εμπιστευόμαστε *a priori* τίποτε, αν δεν το διασταυρώσουμε κατάλληλα, μια που είτε μπορεί να είναι παραπληροφόρηση (είναι τρομερή η άνεση των ΜΓΜ να αδυνατούν να «απαντήσουν»: δεν γνωρίζω – δεν έχω αρκετή πληροφορία), είτε μπορεί να πρόκειται για προκατάληψη (*bias*), είτε για μια σειρά από άλλες περιπτώσεις.

Εν κατακλείδι, διαβάζουμε στο (Γιαλούση, 2023): «δύσκολα μπορεί ακόμα κανείς να αμφισβητήσει ότι πρόκειται για “μία μορφή ευφροσύνης”, τουλάχιστον για κάποιον ορισμό της τελευταίας λέξης. Προφανώς μπορείς να συνεννοηθείς με αυτή τη μηχανή σε ένα πλαίσιο συμφραζομένων».

Η δημοσίευση εκφράζει την άποψη και γνώμη του συντάκτη αυτής και δεν αντικατοπτρίζει απαραίτητα τις απόψεις της Ευρωπαϊκής Ένωσης ή της Ευρωπαϊκής Επιτροπής. Η Ευρωπαϊκή Ένωση και η Ευρωπαϊκή Επιτροπή δεν ευθύνονται για οποιαδήποτε πιθανή χρήση της πληροφορίας αυτής.

Η Εργασία εκπονήθηκε στο πλαίσιο της δράσης «Κέντρα Καινοτομίας σε 13 ΠΔΕ» που είναι ενταγμένη στο Ταμείο Ανάκαμψης και Ανθεκτικότητας και συγχρηματοδοτείται από την Ε.Ε.



Βιβλιογραφικές Αναφορές

- Antonopoulos, V. (2024). Are Llama3 and Mixtral open models ready for RAG in Greek? Retrieved from <https://medium.com/11tensors/are-llama3-and-mixtral-open-models-ready-for-rag-in-greek-7cc1ab7579ed>
- Gavrilidou, M., M. Koutsombogera, A. Patrikakos, and S. Piperidis. (2012). *The Greek language in the digital age*. New York: Springer. <http://www.meta-net.eu/whitepapers/e-book/greek.pdf>.
- Eberhard, D. M., Simons, G. F. & Fennig, C. D (Eds.) (2020). *Ethnologue: Languages of the world (23rd edition)*. Dallas, Texas: SIL International. Retrieved from <https://www.ethnologue.com>.
- Elsevier (2023). *Publishing ethics: the use of generative AI and AI-assisted technologies in scientific writing*. Retrieved October 7, 2023 from <https://www.elsevier.com/about/policies/publishing-ethics>.
- Koga, S. (2023). The Integration of large language models such as ChatGPT in scientific writing: harnessing potential and addressing pitfalls. *Korean J Radiol.*, 24, 924–925.
- Lai, V.D., Ngo, N.T., Veyseh, A.P.B., Man, H., Dernoncourt, F., Bui, T., Nguyen, T.H., (2023). Chatgpt beyond english: Towards a comprehensive evaluation of large language models in multilingual learning'. arXiv preprint arXiv:2304.05613.
- Meltemi (2024). *Meltemi, το πρώτο ανοιχτό Μεγάλο Γλωσσικό Μοντέλο για τα Ελληνικά*. Ανακτήθηκε 15 Μαΐου 2024 από <https://www.ilsp.gr/news/meltemi/>
- OpenAI. (2023). GPT-4 technical report. arXiv preprint arXiv:2303.08774.
- Papantoniou, K. & Tzitzikas, Y. (2020). NLP for the Greek Language: A Brief Survey. In *Proceedings of the 11th Hellenic Conference on Artificial Intelligence (SETN 2020)*, Athens, Greece, September 2-4, 2020, (σελ. 101–9).

- Park, SH. (2023). Use of generative artificial intelligence, including large language models such as ChatGPT, in scientific publications: policies of KJR and prominent authorities. *Korean J Radiol.*, 24, 715–718.
- Radiological Society of North America (RSNA) (2023). *Editorial policies: guidelines for use of large language models*. Retrieved May 15, 2024 from <https://pubs.rsna.org/page/policies>
- Vagelatos, A., Mantzari, E., Pantazara, M., Tsalidis, C., & Kalamara C. (2011). Developing tools and resources for the biomedical domain of the Greek language. *Health Informatics Journal*, 17(2), 127–139. doi:10.1177/1460458211405007.
- Vagelatos, A., Mantzari, E., Pantazara, M., Tsalidis, C. & Kalamara, C. (2022). Healthcare NLP infrastructure for the Greek language. In O. Bojar, S. R. Dash, S. Parida, E. Villatoro-Tello, B. R. Acharya (Eds.), *Natural Language Processing in Healthcare: A Special Focus on Low Resource Language*. CRC Press
- Wikipedia, (2024). *Large Language Model*. Retrieved May 15, 2024 from https://en.wikipedia.org/wiki/Large_language_model
- Zervopoulos, A. D., E. Geramanis, A. Toulakis, A. Papamichail, D. Triantafylloy, T. Tasoulas, and K. Kermanidis. (2019). Language Processing for Predicting Suicidal Tendencies: A Case Study in Greek Poetry. In *Proceedings of the 15th IFIP International Conference on Artificial Intelligence Applications and Innovations (AIAI)*, Hersonissos, Crete, Greece, May 24–26, 2019 (pp. 173–183). New York: Springer.
- Γιαλούση, Ν. (2023). Φιλοσοφικές όψεις των Νευρωνικών Δικτύων. Ανακτήθηκε από <https://opensource.ellak.gr/2023/12/07/filosofikes-proektasis-nevronikon-diktion/>
- Σπινέλλης, Δ. (2024). Επιστολή προς τα μέλη της Εθνικής Επιτροπής Βιοηθικής & Τεχνοηθικής. Ανακτήθηκε 15 Μαΐου 2024 από <https://www2.dmst.aueb.gr/dds/pubs/tr/2024-ethics-hearing/html/response-2024-02-21.pdf>.