

Ανθρώπινη τεχνογνωσία και μηχανική μάθηση: Ένα συνεργατικό μοντέλο για τη βιβλιογραφική επεξεργασία

Το παρόν κείμενο ανακοίνωσης παρουσιάστηκε στο:

31ο Πανελλήνιο Συνέδριο Ακαδημαϊκών Βιβλιοθηκών – 22-24/10/2025

Ιωάννινα

Βιβλιογραφική αναφορά ανακοίνωσης στα ελληνικά & αγγλικά:

- Γερόλιμος, Μ. & Φράγκου, Α. (2025). Ανθρώπινη τεχνογνωσία και μηχανική μάθηση: Ένα συνεργατικό μοντέλο για τη βιβλιογραφική επεξεργασία. *31ο Πανελλήνιο Συνέδριο Ακαδημαϊκών Βιβλιοθηκών*. Ακαδημαϊκή ευημερία, ελευθερία και ακεραιότητα: Από το ΑΒ στο ΑΙ, Ιωάννινα. <https://olympias.lib.uoi.gr/jspui/handle/123456789/39312>
- Gerolimos, M., & Fragkou, A. (2025). Human Expertise and Machine Learning: A Collaborative Model for Bibliographic Processing. *31st Panhellenic Academic Libraries Conference*. Academic Well-being, Freedom, and Integrity: From AB to AI, Ioannina. <https://olympias.lib.uoi.gr/jspui/handle/123456789/39312>

Περιεχόμενα ηλεκτρονικού αρχείου:

- Πλήρες κείμενο εργασίας / full-text in Greek
- Στοιχεία ανακοίνωσης στην ελληνική γλώσσα
- Conference paper's details in English

Για φιλικότερη πλοήγηση συνιστάται εμφάνιση του δυναμικού παραθύρου περιήγησης (επικεφαλίδων/σελιδοδεικτών).

Το περιεχόμενο διατίθεται **ελεύθερα** με άδεια χρήσης: Creative Commons Attribution CC **BY-NC-ND** 4.0 International License

<https://creativecommons.org/licenses/by-nc-nd/4.0/>



Πλήρες κείμενο εργασίας / full-text in Greek

Εισαγωγή

Η αυτοματοποίηση της βιβλιογραφικής επεξεργασίας (BE) αποτελεί διαρκή επιδίωξη των βιβλιοθηκών πάντα με τη χρήση προτύπων. Η έννοια του αυτοματισμού εμφανίστηκε στις βιβλιοθήκες από τη δεκαετία του 1960, με την εισαγωγή του MARC, και εδραιώθηκε με την εμφάνιση των ολοκληρωμένων συστημάτων βιβλιοθηκών, τη δεκαετία 1970-1980. Στο πλαίσιο αυτό, το Τμήμα Καταλόγων της Εθνικής Βιβλιοθήκης της Ελλάδος θέλησε να αναπτύξει ένα εργαλείο αυτόματης καταλογογράφησης σε MARC21 με χρήση τεχνητής νοημοσύνης (TN). Το εγχείρημα αποτελεί εξέλιξη του αυτοματισμού της BE, πάντα με τη χρήση τυποποίησης. Η εργασία συνοψίζει τη μέχρι τώρα πορεία, καταγράφοντας στόχους, στάδια ανάπτυξης της ροής εργασίας, τεχνικές επιλογές, αποτυχίες και βελτιώσεις, με στόχο μια ρεαλιστική, αναπαραγώγιμη και επεκτάσιμη πρακτική.

Βιβλιογραφική επισκόπηση

Σύμφωνα με την IFLA (2020) η εφαρμογή της TN στις βιβλιοθήκες δίνει τη δυνατότητα για ακόμα ουσιαστικότερη εκπλήρωση του βασικού τους στόχου, της προσφοράς υπηρεσιών στην κοινωνία. Ειδικότερα, μέσω της TN μπορεί να εισαχθούν νέες υπηρεσίες και να εξελιχθούν υπάρχουσες εφαρμογές, όπως το OCR και τα υπάρχοντα βιβλιοθηκονομικά συστήματα. Ειδικότερα, στην καταλογογράφηση, εργασία παραδοσιακά απαιτητική από άποψη μόχθου και διανοητικής προσπάθειας, ενισχύεται σταδιακά από αυτοματοποιημένα συστήματα που βασίζονται σε αλγορίθμους TN, μηχανικής μάθησης και επεξεργασία φυσικής γλώσσας (Natural Language Processing, NLP).

Από την δεκαετία του 1990, οι μελέτες στρέφονται στην αυτοματοποίηση της θεματικής ευρετηρίασης με απόδοση ταξινομικών αριθμών της Βιβλιοθήκης του

Κογκρέσου (Library of Congress Classification - LCC) και της Δεκαδικής Ταξινόμησης Dewey (Dewey Decimal Classification - DDC) (Golub, 2021).

Οι βιβλιοθήκες πειραματίζονται στη χρήση της TN ή/και τη μηχανική μάθηση, ώστε να βοηθηθούν στην αυτοματοποιημένη παραγωγή θεματικών επικεφαλίδων. Συγκεκριμένα, από το 2017 η Φινλανδική Εθνική Βιβλιοθήκη ανέπτυξε το Annpif, ένα εργαλείο ανοιχτού λογισμικού για την αυτοματοποιημένη θεματική ευρετηρίαση (Suominen, 2019). Η Γερμανική Εθνική Βιβλιοθήκη (DNB) από το 2022 το χρησιμοποιεί στο σύστημα EMa (Erschließungsmaschine, «μηχανή καταλογογράφησης»), το οποίο παράγει αυτόματα επικεφαλίδες θεμάτων από αρχεία καθιερωμένων ονομάτων (GND) και συνοπτικούς αριθμούς DDC για επιλεγμένους θεματικούς τομείς (Poley et al., 2025). Τέλος, οι Asula et al. (2022) περιγράφουν την ανάπτυξη του «Krat» ενός εργαλείου αυτόματης θεματικής ευρετηρίασης που αξιοποιεί μηχανική μάθηση.

Παράλληλα, σύμφωνα με τη μελέτη των Luo, He & Zhang (2022), τα συστήματα AI μπορούν να αναγνωρίζουν βιβλιογραφικές πληροφορίες από εξώφυλλα, περιλήψεις ή πλήρη κείμενα, και να δημιουργούν αυτόματα MARC εγγραφές. Η χρήση τεχνικών NLP επιτρέπει στον υπολογιστή να κατανοεί τη δομή και το νόημα των κειμένων, διευκολύνοντας έτσι τη διαδικασία απόδοσης θεματικών επικεφαλίδων και ταξινομικών αριθμών. Στο άρθρο του Mwantimwa (2025), διερευνώνται οι τρόποι χρήσης της Παραγωγικής Τεχνητής Νοημοσύνης (Generative AI) στο χώρο των βιβλιοθηκών από τη δεκαετία του 1990 έως το 2023. Διαπιστώνει σαφή αύξηση στον πειραματισμό των βιβλιοθηκών στην εκτέλεση ποικίλων εργασιών με τη χρήση του. Εντοπίζει τόσο ευρύ πεδίο δυνατοτήτων όσο και προκλήσεων για τον κλάδο. Ακόμη, οι Saccucci & Potter (2025) από τη Βιβλιοθήκη του Κογκρέσου εξέτασαν τον τρόπο με τον οποίο η Τεχνητή Νοημοσύνη μπορεί να χρησιμοποιηθεί για την αυτόματη καταλογογράφηση μεγάλου όγκου ψηφιακών βιβλίων. Ωστόσο, επισημαίνουν ότι η ανθρώπινη επίβλεψη και ο έλεγχος παραμένουν απαραίτητα για τη διασφάλιση της ποιότητας και της ακρίβειας της διαδικασίας. Ακριβώς το ίδιο επισήμανε και ο Chow (2024), ο οποίος χρησιμοποίησε το ChatGPT για να δημιουργήσει

θεματικές επικεφαλίδες για ηλεκτρονικές διατριβές και εργασίες βάσει των τίτλων και των περιλήψεών τους. Την επίκριση του επιστημονικού κόσμου για την πλήρη αυτοματοποίηση τέτοιων εργασιών υπογραμμίζει και η Golub (2021).

Στόχοι και ιστορικό

Η αφετηρία της προσπάθειας αποτέλεσε ο Ιούλιος του 2024, οπότε για πρώτη φορά άρχισε να διερευνάται η προοπτική ανάπτυξης μιας ροής εργασίας που θα επέτρεπε τη δημιουργία ή/και τον εμπλουτισμό βιβλιογραφικών εγγραφών με μηχανικό τρόπο. Βασική επιδίωξη είναι η εξοικονόμηση χρόνου από την ΒΕ και επένδυσή του στην αμιγώς πνευματική εργασία της τεκμηρίωσης.

Σύμφωνα με τον αρχικό σχεδιασμό, η εφαρμογή θα έπρεπε να εξυπηρετεί τρεις διαφορετικές περιπτώσεις: (1) προσθήκη πεδίου περιεχομένων (505) και περίληψης (520) σε ήδη καταλογογραφημένα τεκμήρια, (2) καταλογογράφηση ακατάγραφων εισαγωγών, (3) καταλογογράφηση συλλογών από δωρεές. Επιπλέον, τέθηκαν μελλοντικοί στόχοι εμπλουτισμού με θεματικές επικεφαλίδες (6xx), ιδανικά προερχόμενες από το Αρχείο Καθιερωμένων Επικεφαλίδων της Εθνικής Βιβλιοθήκης της Ελλάδος (nlgaf).

Πρώτη φάση

Πρώτο στάδιο - Διερεύνηση

Στο πρώτο στάδιο ανάπτυξης έγινε προσπάθεια δημιουργίας βιβλιογραφικών εγγραφών από αρχείο PDF το οποίο περιείχε τις ακόλουθες σελίδες:

- σελίδα τίτλου
- περιεχόμενα (παραλείπεται αν δεν υπάρχουν ή αν περιλαμβάνουν αρίθμηση χωρίς τίτλους των επιμέρους περιεχομένων)
- οπισθόφυλλο (παραλείπεται αν δεν περιλαμβάνει περίληψη)
- σελίδα με ετικέτα RFID

Σε αυτό το στάδιο η επεξεργασία γινόταν με τη χρήση του ChatGPT (έκδοση 4.0) καθώς και με την εφαρμογή CATMELK του RCG Gamage, η οποία είναι

ενσωματωμένη στο ChatGpt και αξιοποιείται για δημιουργία βιβλιογραφικών εγγραφών.

Τα αποτελέσματα σε 790 τεκμήρια ήταν ικανοποιητικά, καθώς διαπιστώθηκε α) η σωστή αναγνώριση των τύπων δεδομένων, β) του σωστού πεδίου καταγραφής (ακόμη και σε παιδικά βιβλία που πηγή της καταλογογράφησης μπορεί να είναι σελίδες με πλούσια εικονογράφηση), γ) δυσκολία στη δημιουργία βιβλιογραφικής εγγραφής όταν η γλώσσα καταλογογράφησης ήταν η ελληνική, δ) δυσκολία στην απόδοση ταξινομικού αριθμού στα ελληνικά έργα και ε) περιορισμένη ακρίβεια στην ευρετηρίαση θεματικών όρων καθιερωμένων ή μη.

Δεύτερο στάδιο - Ανάπτυξη

Στο δεύτερο στάδιο, η EBE ξεκίνησε μια συνεργασία¹ με την Bibliothèque royale de Belgique (KBR), μέσω του Dr. Hannes Lowagie, στην οποία είχε ήδη αναπτυχθεί μια παραγωγική ροή εργασίας για την καταγραφή τεκμηρίων με χρήση TN. Σε αυτή τη φάση έγινε ανάπτυξη ενός προσαρμοσμένου μοντέλου στο εργαλείο Microsoft Power Apps και στη συνέχεια ένταξή του σε μια ροή στο Power Automate με δοκιμαστική άδεια.

Εκπαίδευση μοντέλου

Λόγω των περιορισμών του PowerApps, η EBE δεν μπορούσε να χρησιμοποιήσει το εκπαιδευμένο μοντέλο της KBR και ως εκ τούτου, έπρεπε να αναπτύξει και να εκπαιδεύσει το δικό της προσαρμοσμένο μοντέλο. Στο δεύτερο στάδιο ανάπτυξης, επελέγησαν επίσης τεκμήρια με λατινικό αλφάβητο, θεωρώντας ότι η ύπαρξη λατινικών χαρακτήρων θα διευκόλυνε το μοντέλο να αναγνωρίσει τον τύπο δεδομένων που αντλεί από το τεκμήριο.

Η εκπαίδευση του μοντέλου στην αναγνώριση του τύπου δεδομένων από τη σελίδα τίτλου, προϋπέθετε την ομαδοποίησή τους βάσει σελιδοποίησης, δηλαδή

¹ Στο πλαίσιο του Erasmus+, το προσωπικό της EBE συμμετείχε σε πρόγραμμα επιμόρφωσης που υλοποίησε δια ζώσης ο Dr. Hannes Lowagie, στο οποίο παρουσίασε τη ροή που υλοποιούσαν και συνεργάστηκε με το Τμήμα Καταλόγων ως καθοδηγητής για το σχεδιασμό μιας ροής που θα εξυπηρετούσε τις ανάγκες της EBE, όπως αυτές ορίζονταν από τους στόχους της πρωτοβουλίας αυτής.

τη δημιουργία ομάδων τεκμηρίων με συγκεκριμένη δομή σελίδας τίτλου. Για παράδειγμα, σελίδες με έναν συγγραφέα στο πάνω μέρος και κύριο τίτλο εκδότη στο κάτω μέρος εντάχθηκαν σε μία ομάδα, ενώ μια άλλη περιείχε σελίδες τίτλου όπου κάτω από τον τίτλο υπήρχε υπότιτλος ακολουθούμενος από δευτερεύουσα υπευθυνότητα και στη συνέχεια τα πλήρη στοιχεία δημοσίευσης-διάθεσης. Έτσι, στις πρώτες δοκιμές δημιουργήθηκαν οκτώ διαφορετικές ομάδες συνδυάζοντας τα παρακάτω στοιχεία: Συγγραφέα 1 και Συγγραφέα 2, Τίτλο, Υπότιτλο, Εκδότη σε διαφορετικές παραλλαγές ως προς τη σειρά με την οποία εμφανίζονταν.

Η ομάδα ανάπτυξης του μοντέλου ταυτοποιούσε πάνω στη σελίδα τους διαφορετικούς τύπους δεδομένων. Τα δεδομένα καταχωρίζονταν από το μοντέλο σε ένα πίνακα με τις ομώνυμες στήλες. Ο πίνακας περιείχε επιπλέον στήλες που θα συμπληρώνονταν μετά την εκπαίδευση άλλου μοντέλου ώστε να ενσωματώσει στοιχεία του verso της σελίδας τίτλου (χρονολογία copyright και ISBN), του οπισθόφυλλου (περίληψη) και των περιεχομένων.

Από τα δεδομένα του πίνακα δημιουργούνταν μια εγγραφή με αντιστοίχιση των στηλών του πίνακα σε πεδία MARC21 (διαδικασία παρόμοια με τη δημιουργία εγγραφής MARC από αρχείο csv).

Η εκπαίδευση με χρήση λατινικών χαρακτήρων ήταν ικανοποιητική με την προϋπόθεση η σελίδα τίτλου να ακολουθούσε σταθερή δομή.

Τρίτο στάδιο - Ανάπτυξη νέου μοντέλου

Στο στάδιο αυτό ξεκίνησε η εκπαίδευση του μοντέλου με τεκμήρια με ελληνικούς χαρακτήρες.

Έχοντας ως στόχο την παραγωγή εγγραφής MARC21, η διάκριση διαφορετικών τύπων τίτλων (παράλληλος, εναλλακτικός κ.λπ.) ή υπευθυνοτήτων (εικονογράφος, μεταφραστής κ.λπ.) ήταν θεμελιώδης για την επίτευξη της σωστής στίξης γεγονός που αύξησε τη δυσκολία εκπαίδευσης του μοντέλου.

Ενδεικτικό είναι ότι κατά την εκπαίδευση² με 184 σελίδες τίτλου, η ομάδα αναγνώρισε 92 διαφορετικές παραλλαγές διάταξης για τους εξής τύπους δεδομένων:

- κύριος τίτλος,
- υπότιτλος,
- συγγραφέας,
- συγγραφέας επίλογου, επιμέτρου, κ.λπ.,
- επιμελητής έκδοσης,
- επιστημονικός επιμελητής,
- εικονογράφος,
- μεταφραστής,
- συγγραφέας εισαγωγής,
- συντελεστής,
- δήλωση έκδοσης
- τόπος έκδοσης,
- εκδότης,
- ημερομηνία έκδοσης,
- αριθμός τόμου
- τίτλος τόμου
- τίτλος σειράς
- αριθμός σειράς.

Η εκπαίδευση του μοντέλου σε 92 ομάδες θα ήταν σχετικά δύσκολη λόγω του μεγάλου όγκου εργασίας που θα απαιτούσε. Ως εκ τούτου, δημιουργήθηκαν μόνο έντεκα ομάδες με τις συχνότερες διατάξεις σελίδας τίτλου. Κατά τη διαδικασία δοκιμής του μοντέλου, διαπιστώθηκε ότι υπήρξαν λάθη ακόμη και σε σελίδες με πανομοιότυπη διάταξη. Για παράδειγμα, υπήρχε δυσκολία αναγνώρισης του τύπου των δεδομένων και συγκεκριμένα δεν μπορούσε να διακρίνει αν η πληροφορία που βρισκόταν κάτω από τον τίτλο ήταν υπότιτλος ή υπευθυνότητα.

² Για την εκπαίδευση, το PowerApps, προϋποθέτει τη χρήση τουλάχιστον πέντε δειγμάτων, στην περίπτωση αυτή σελίδων τίτλου.

Αντίστοιχα προβλήματα διαπιστώθηκαν κατά την οπτική αναγνώριση χαρακτήρων (Optical Character Recognition, OCR).

Δεύτερη φάση: υλοποίηση

Τα αποτελέσματα του τρίτου σταδίου οδήγησαν σε αναθεώρηση της όλης διαδικασίας. Κατά το ίδιο χρονικό διάστημα, στο Power Apps ενσωματώθηκε η δυνατότητα χρήσης του ChatGpt, η οποία έκανε εφικτή την αλλαγή της ροής καταλογογράφησης μέσω TN, αυτή τη φορά όχι μέσω εκπαίδευσης, αλλά χρήσης αποτελεσματικών προτροπών (prompts).

Στη διαδικασία αναθεώρησης του τρόπου προσέγγισης αξιοποίησης της TN, συνέβαλε και η απαίτηση να χρησιμοποιηθεί και για τεκμήρια που δεν έχουν βιβλιογραφική εγγραφή. Δηλαδή, η απαίτηση ήταν όχι μόνο ο εμπλουτισμός υπάρχουσών εγγραφών, αλλά και η δημιουργία νέων. Επίσης, η ομάδα διαπίστωσε ότι ήταν ωφέλιμο και αποδοτικό να εξάγονται κατευθείαν εγγραφές με έτοιμα τα πεδία MARC21 αντί για δεδομένα σε πίνακα σε μορφή .csv, όπως συνέβαινε με το αρχικό μοντέλο. Έτσι, θα ήταν πιο εύκολη και η συγχώνευση με υπάρχουσες εγγραφές στον κατάλογο της ΕΒΕ, μια ακόμα συνθήκη αξιοποίησης της TN για καταλογογράφηση.

Αλλαγή στη ροή σάρωσης

Λόγω της νέας προσέγγισης στην αξιοποίηση της TN με χρήση προτροπών, κρίθηκε σκόπιμο να αλλάξει και η ροή της σάρωσης των τεκμηρίων. Καθώς η πλειοψηφία των τεκμηρίων είχε σελίδα τίτλου, verso και μια σελίδα με την ετικέτα RFID, η σάρωση ακολουθεί αυτή τη σειρά. Έπειτα, σαρώνεται το οπισθόφυλλο, το οποίο συνήθως περιέχει την περίληψη και τέλος σαρώνονται τα περιεχόμενα, με ποικίλη έκταση.

Σύνταξη προτροπής (prompt)

Μέσα από το AI hub του Power Apps δίνεται η δυνατότητα σύνταξης προτροπών. Πρόκειται για οδηγίες που καταγράφονται σε φυσική γλώσσα και με την ενσωματωμένη υλοποίηση του ChatGpt 4.1 mini εκτελούνται ως πρόγραμμα.

Στην υλοποίηση της EBE, οι προτροπές ενσωματώνουν τις οδηγίες καταλογογράφησης που ακολουθεί το Τμήμα Καταλόγων της EBE. Σημειώνεται ότι ο σκοπός της αξιοποίησης της TN είναι η υποστήριξη της BE του υλικού με την παραγωγή μιας βασικής εγγραφής με τα κύρια στοιχεία, όπως τίτλος, συγγραφέας και εκδότης, προκειμένου, στη συνέχεια, να εμπλουτιστεί από τον καταλογογράφο.

Η εφαρμογή Power Apps δίνει τη δυνατότητα χρήσης τριών διαφορετικών εκδόσεων του ChatGPT, συγκεκριμένα, των o3, 4.1 mini και 4.1.

Η χρήση της έκδοσης o3 του ChatGPT είχε την καλύτερη απόδοση, τόσο σε σύγκριση με τις άλλες εκδόσεις του ChatGPT, όσο και με το προηγούμενο μοντέλο, καθώς δημιουργήθηκαν εγγραφές σε MARC21 οι οποίες περιείχαν τα πεδία που ζητήθηκαν:

- LDR: προεπιλεγμένη τιμή (δεν έχει τη δυνατότητα να κατασκευάσει σωστά τις 5 πρώτες θέσεις).
- 020: αντλείται από το verso αλλά δεν δημιουργεί \$z όταν υπάρχει λανθασμένο ούτε βάζει προσδιορισμό τόμου στο \$q όταν πρόκειται για πολύτομο.
- 082: δεν αποδίδει σταθερά τον ίδιο ταξινομικό αριθμό σε κάθε δοκιμή
- 245: καταγράφει σωστά όλες τις πληροφορίες, με σωστή χρήση πεζών-κεφαλαίων και στίξη.
- 250: καταγράφει τη δήλωση έκδοσης σωστά και με κανονικοποίηση των τακτικών αριθμητικών σύμφωνα με το προφίλ εφαρμογής του RDA που έχει επιλέξει η EBE.
- 264: Δημιουργεί σωστά το 264, όμως, δεν έχει τη δυνατότητα δημιουργίας επαναλαμβανόμενου πεδίου όταν το έτος έκδοσης διαφοροποιείται από το έτος copyright.
- 490: Αναγνωρίζει και καταγράφει σωστά τον τίτλο της σειράς (όχι καθιερωμένη μορφή αλλά απλή καταγραφή της μορφής που συναντά στο τεκμήριο).

- 500:
 - ο Δημιουργεί μια σωστή σημείωση για τη γλώσσα του κειμένου, η οποία αξιοποιείται εκτός του μοντέλου για τη δόμηση του σωστού 008.
 - ο Δημιουργεί σημείωση με τον τίτλο πρωτοτύπου, η οποία μπορεί να αξιοποιηθεί από τον καταλογογράφο για τη δημιουργία του 240.
- 505: Δεν έχει σταθερή επίδοση στη χρήση της στίξης.
- 520: Όταν επιτυγχάνεται η χρήση κεφαλαίων στα άλλα πεδία αποτυγχάνει στην περίληψη.

Αντίθετα, παρατηρήθηκε ότι το μοντέλο 4.1 αφαιρεί ακρίβεια από το αποτέλεσμα., καθώς α) περιέχει λάθη στη χρήση των πεζών/κεφαλαίων σε όλα σχεδόν τα πεδία, β) δεν παράγει σωστά περιεχόμενα (τόσο προς τη στίξη όσο και προς το περιεχόμενο) και γ) δεν καταγράφει σωστά την ημερομηνία έκδοσης. Ωστόσο, διατηρεί σωστά τα σημεία στίξης και τα κεφαλαία/πεζά στην περίληψη.

Αντίστοιχα, το μοντέλο 4.1 mini συγκριτικά έχει τη χειρότερη επίδοση, καθώς α) αφαιρεί τόνους β) κάνει λανθασμένη χρήση κεφαλαίων σε όλα τα πεδία, γ) δεν σχηματίζει σωστά το πεδίο 264, δ) αποδίδει πολύ γενικό ταξινομικό αριθμό και ε) βάζει λάθος στίξη στα περιεχόμενα. Ωστόσο, μεταφέρει σωστά την περίληψη, διατηρώντας τα κεφαλαία όπου χρειάζεται.

Ως γενική παρατήρηση, διαπιστώθηκε ότι δεν υπάρχει συνέπεια στα αποτελέσματα της προτροπής, σε όλα τα μοντέλα.

Δημιουργία ροής στο Power Automate

Για την αυτοματοποίηση της εργασίας έπρεπε να δημιουργηθεί μια ροή εργασίας η οποία θα ενεργοποιείται με αυτόματο τρόπο. Σε αυτό το στάδιο, έγινε χρήση του Power Automate σε μία ροή επτά βημάτων:

Βήμα 1^ο: ορίστηκε το έναυσμα της διαδικασίας η οποία ενεργοποιείται όταν αποθηκεύεται σε συγκεκριμένο φάκελο του OneDrive το PDF με τις σαρωμένες σελίδες του τεκμηρίου που περιγράφηκαν παραπάνω.

Βήμα 2º: Αντλείται το περιεχόμενο του PDF και αποστέλλεται στην προτροπή.

Βήμα 3º: Εκτελείται η προτροπή η οποία παράγει την εγγραφή MARC21.

Βήμα 4º: Εξάγεται το αποτέλεσμα της προτροπής.

Βήμα 5º: Εντοπίζεται στην εγγραφή MARC η τιμή της ετικέτας RFID του τεκμηρίου

Βήμα 6º: Χρησιμοποιείται η τιμή της ετικέτας RFID για την ονοματοδοσία του εξαγομένου αρχείου με επέκταση txt³ (text file).

Βήμα 7º: Δημιουργία του αρχείου σε ορισμένο φάκελο του One Drive for Business με την ονομασία που ορίστηκε στο προηγούμενο βήμα.

Τρίτη φάση - Παραγωγή

Κατά τρίτη φάση προβλέπεται η ένταξη της ροής εργασίας στον λειτουργικό ιστό του Τμήματος Καταλόγων. Το πρώτο επιδιωκόμενο αποτέλεσμα είναι ο εμπλουτισμός των εγγραφών τεκμηρίων που λαμβάνονται με την Κατά Νόμον Κατάθεση (ΚΝΚ) με περιεχόμενα και περιλήψεις, πριν την καταλογογράφηση, αυξάνοντας σημαντικά την αξία των εγγραφών. Σε μεταγενέστερο στάδιο, θα επιχειρηθεί ο εμπλουτισμός υπάρχουσών εγγραφών του καταλόγου με περιεχόμενα και περιλήψεις.

Εφόσον ολοκληρωθεί η διαδικασία εμπλουτισμού των εγγραφών με τα δύο προαναφερόμενα πεδία, απώτερος στόχος είναι: η δημιουργία νέων εγγραφών στην πληρέστερη δυνατή μορφή όλων των πεδίων MARC που μπορεί να δημιουργηθούν με τη χρήση της TN, ώστε να καταγραφούν και να καταλογογραφηθούν σε μια βασική μορφή α) βιβλία που έχουν ενταχθεί στις συλλογές της ΕΒΕ αλλά δεν έχουν ακόμα καταγραφεί. Σε αυτή την περίπτωση, η

³ Ο λόγος που επελέγη το txt αντί του mrk (marc mnemonic file) ήταν ότι για τα txt λειτουργεί τόσο η προεπισκόπηση στο Microsoft365 όσο και στο Windows File Explorer, χωρίς να απαιτείται άλλο λογισμικό ενώ η επεξεργασία τους μπορεί να γίνει με κάθε επεξεργαστή κειμένου. Όσο η επέκταση του αρχείου παραμένει txt, η μετατροπή τους σε mrk γίνεται με απλή αλλαγή της επέκτασης κατά τη μετονομασία του αρχείου χωρίς απώλεια δεδομένων.

διαδικασία θα εκκινεί από τις φωτογραφίες των βιβλίων, και στη συνέχεια θα γίνεται η εισαγωγή τους στα συστήματα της EBE και β) εισαγωγή νέων τεκμηρίων που η EBE παραλαμβάνει στο μέλλον και δεν υπάρχουν καταγεγραμμένες κάπου ως εγγραφές MARC.

Οι ελάχιστες εργασίες για τους υπαλλήλους του Τμήματος Καταλόγων θα είναι:

1. Η φωτογράφιση των απαραίτητων στοιχείων του βιβλίου,
2. ο έλεγχος του αποτελέσματος που παράγει η TN πριν την εισαγωγή στον κατάλογο και σχετική επιμέλεια,
3. ο έλεγχος των πληροφοριών που θα ενσωματωθούν στις εγγραφές του καταλόγου μετά την εισαγωγή, και σε μεταγενέστερο χρόνο,
4. ο έλεγχος των νέων εγγραφών που θα δημιουργούνται αυτόματα από την TN και θα εισάγονται στον κατάλογο

Προκλήσεις

Από τις έως τώρα δοκιμές που έχουν γίνει, φαίνεται ότι ο εμπλουτισμός των εγγραφών με περιεχόμενα και περίληψη θα μπορέσει να υλοποιηθεί με τη χρήση της TN, καθώς τα αποτελέσματα είναι ικανοποιητικά. Όταν στο μοντέλο η προτροπή που εκτελείται είναι αυτή της εξαγωγής μόνο αυτών των 2 πεδίων και του RFID, τότε η απόδοση γίνεται σωστά με ελάχιστα σφάλματα. Ωστόσο, έχουν καταγραφεί ορισμένα ζητήματα και προσκλήσεις, κοινά για όλες τις προσπάθειες που έχουν αναφερθεί παραπάνω, κατά την εκτέλεση των μοντέλων που αναπτύσσει η EBE.

Ελληνική γλώσσα, OCR και πολυτονικό

Η αναγνώριση της ελληνικής γλώσσας αποτελεί τη σημαντικότερη, ίσως, πρόκληση. Από τις δοκιμές φαίνεται να υπάρχει δυσκολία στην επεξεργασία της ελληνικής γλώσσας, τουλάχιστον με τα μοντέλα του ChatGPT που είναι ενσωματωμένα στο PowerApps (4.1-mini, 4.1 και o3), ειδικά στην εννοιολογική ανάλυση των ελληνικών, κάτι το οποίο γίνεται εμφανές στην αστάθεια που παρουσιάζει η απόδοση ταξινομικού αριθμού. Η ταξινόμηση από τη φύση της

ενέχει το στοιχείο της υποκειμενικότητας και είναι αναμενόμενο μια μηχανή να μην μπορεί να χαρακτηρίσει μονοσήμαντα ένα τεκμήριο, μπορεί, όμως, να λειτουργήσει ως οδηγός για τους καταλογογράφους που θα κληθούν να αποδώσουν ταξινομικό αριθμό, εργασία ούτως ή άλλως πνευματική.

Πρόκληση παραμένει, επίσης, και η αναγνώριση των σημείων του πολυτονικού, τα οποία φαίνεται να μην αναγνωρίζει καθόλου και συνεπώς τα αγνοεί. Η λύση που εξετάζεται είναι η εκ των υστέρων διόρθωση από τον άνθρωπο, έως ότου, τουλάχιστον, το μοντέλο να αποκτήσει την απαιτούμενη εμπειρία για την αντιμετώπισή του.

Leader

Ο Leader αποτελεί την ταυτότητα εγγραφής. Το γεγονός ότι το μοντέλο δεν δύναται να παράγει ορθά τις πέντε πρώτες θέσεις που ουσιαστικά διαφοροποιούν τη μια εγγραφή από την άλλη, οδήγησε στην απόφαση να δίνεται μια προεπιλεγμένη τιμή στο μοντέλο, ώστε να μη μένει κενό το πεδίο. Όμως, η χρήση του ίδιου Leader σε όλες τις εγγραφές προκαλεί συστημικά προβλήματα καθώς πρόκειται για την ταυτότητα εγγραφής και δεν είναι ορθό δύο εγγραφές να έχουν την ίδια, με αποτέλεσμα την αποτυχία εισαγωγής σε Koha. Η επίλυση του ζητήματος αυτού εκκρεμεί.

Πολυσέλιδα περιεχόμενα

Όπως έχει αναφερθεί, γίνεται χρήση της εφαρμογής OneDrive, καθώς επιτρέπει τη συνεργασία μεταξύ των υπαλλήλων, την αυτόματη αποθήκευση και διάθεση μέσω cloud των φωτογραφιών που παράγονται για τη χρήση από την εφαρμογή Power Apps. Ωστόσο, έχει τον περιορισμό ότι επιτρέπει έως δέκα σελίδες ανά παραγόμενο PDF, κάτι που θα δημιουργήσει προβλήματα εφόσον ένα τεκμήριο περιέχει περισσότερες από έξι σελίδες με περιεχόμενα⁴.

⁴ Η λύση που επεξεργάζεται το Τμήμα για την αντιμετώπιση του συγκεκριμένου ζητήματος είναι η εκ των υστέρων συμπλήρωση εκτός της ροής του PowerAutomate. Διερευνάται το ενδεχόμενο σάρωσης του επιπλέον κειμένου με

Τρέχουσα κατάσταση και επόμενα βήματα

Τη στιγμή που γράφεται η εργασία η EBE είναι έτοιμη να παράγει εγγραφές MARC21 με ικανοποιητική απόδοση σε περιεχόμενα και περίληψη, με την έμφαση για την ακρίβεια των αποτελεσμάτων να εστιάζει σε αυτά τα δύο πεδία, καθώς εγγραφές σε στάδιο προκαταλογογράφησης ήδη παράγονται στο Τμήμα Εισαγωγής και προωθούνται στο Τμήμα Καταλόγων. Στα ζητούμενα αυτή τη στιγμή εντάσσονται η ανάγκη για εξοικείωση με την αυτοματοποίηση εργασιών και τη δημιουργία ροών εργασίας αλλά και τον σχεδιασμό μιας ροής ελέγχου των παραγομένων αυτής της διαδικασίας.

Στο αμέσως επόμενο διάστημα θα ξεκινήσει η εκπαίδευση του προσωπικού που θα σαρώνει τεκμήρια και στη συνέχεια θα ελέγχει τα αποτελέσματα της ροής. Όταν κατακτηθεί η απαιτούμενη τεχνογνωσία και φτάσει το μοντέλο να παράγει ικανοποιητικά πρωτότυπες βιβλιογραφικές εγγραφές, θα ξεκινήσει η υλοποίηση σε μεγαλύτερη κλίμακα και με περισσότερα δεδομένα.

Σε επίπεδο προτύπων, σχεδιάζεται επιμερισμός της δημιουργίας MARC σε επιμέρους prompts ανά περιοχή ISBD, με στόχο βελτίωση της ακρίβειας και μείωση σφαλμάτων. Το έργο παραμένει υπό ανάπτυξη σε ό,τι αφορά την περιγραφική καταλογογράφηση όσο και την θεματική ευρετηρίαση.

Σχετικά με την καθημερινή ροή εργασιών, στόχος είναι η εκπαίδευση του προσωπικού τόσο του Τμήματος Καταλόγων όσο και των τμημάτων που εμπλέκονται στη διαδικασία εισαγωγής και επεξεργασίας τεκμηρίων με στόχο την ενσωμάτωση της καταλογογράφησης με χρήση TN στις ροές καταγραφής και επεξεργασίας τεκμηρίων.

Google Lens και επεξεργασία του παραχθέντος txt ώστε να συμπληρώνεται η σημείωση περιεχομένων με ενσωμάτωση της κατάλληλης στίξης στο ChatGpt εκτός του περιβάλλοντος της Microsoft.

Συμπεράσματα

Η εμπειρία της ΕΒΕ δείχνει ότι το PowerApps και η ΤΝ μπορούν να επιταχύνουν τη ΒΕ, αλλά εξακολουθούν να απαιτούν ανθρώπινη εποπτεία για κρίσιμα πεδία. Η μελλοντική εργασία θα επικεντρωθεί στη βελτίωση OCR πολυτονικού και στη σωστή δημιουργία Leader.

Βιβλιογραφία

- Chow, E. H. C., Kao, T. J., & Li, X. (2024). An experiment with the use of ChatGPT for LCSH subject assignment on electronic theses and dissertations. *Cataloging & Classification Quarterly*, 62 (7), 574–588.
<https://doi.org/10.48550/arXiv.2403.16424>
- Golub, K. (2021). Automated Subject Indexing: An Overview. *Cataloging & Classification Quarterly*, 59(8), 702–719.
<https://doi.org/10.1080/01639374.2021.2012311>
- IFLA FAIFE (Committee on Freedom of Access to Information and Freedom of Expression) (2020). IFLA Statement on Libraries and Artificial Intelligence. International Federation of Library Associations and Institutions (IFLA). <https://repository.ifla.org/handle/20.500.14598/1646>
- Mwantimwa, K., & Msoffe, G. (2025). Application of generative artificial intelligence in library operations and service delivery: A scoping review. *Technical Services Quarterly*, 42(2), 139–168.
<https://doi.org/10.1080/07317131.2025.2467574>
- **Poley, C., Uhlmann, S., Busse, F., Jacobs, J.-H., Kähler, M., Nagelschmidt, M., & Schumacher, M.** (2025). Automatic subject cataloguing at the German National Library. *LIBER Quarterly: The Journal of the Association of European Research Libraries*, 35(1), 1–29.
<https://doi.org/10.53377/lq.19422>

- Saccucci, C., & Potter, A. (2025). Assessing machine learning for cataloging at the Library of Congress. In E. Balnaves, L. Bultrini, A. Cox, & R. Uzwysyn (Eds.), *New horizons in artificial intelligence in libraries* (pp. 227–238). Berlin/Boston: De Gruyter Saur.
<https://doi.org/10.1515/9783111336435-017>
- **Suominen, O.** (2019). Annif: DIY automated subject indexing using multiple algorithms. *LIBER Quarterly: The Journal of the Association of European Research Libraries*, 29(1), 1–25.
<https://doi.org/10.18352/lq.10285>

Στοιχεία ανακοίνωσης στην ελληνική γλώσσα

Τίτλος εργασίας

Ανθρώπινη τεχνογνωσία και μηχανική μάθηση: Ένα συνεργατικό μοντέλο για τη βιβλιογραφική επεξεργασία

Στοιχεία συγγραφέα / συγγραφέων

Μιχάλης Γερόλιμος, Εθνική Βιβλιοθήκη της Ελλάδος, m.gerolimos@nlg.gr

Αφροδίτη Φράγκου, Εθνική Βιβλιοθήκη της Ελλάδος, afragkou@nlg.gr

Λέξεις-κλειδιά

1. Εθνική Βιβλιοθήκη
2. Τεχνητή Νοημοσύνη
3. Καταλογογράφηση
4. Βιβλιογραφική Επεξεργασία
5. PowerApps

Περίληψη

Από τη δεκαετία του 1960 έως και σήμερα, οι βιβλιοθηκονόμοι αναζητούν διαρκώς νέους τρόπους αξιοποίησης της τεχνολογίας της πληροφορικής σε διάφορες υπηρεσίες που προσφέρουν οι βιβλιοθήκες. Η βιβλιογραφική επεξεργασία των συλλογών αποτελεί διαχρονικά μία από τις εργασίες στις οποίες, λόγω της φύσης της, επιδιώκεται η μέγιστη δυνατή ενσωμάτωση τεχνολογικών εφαρμογών τόσο για τη διεκπεραίωση της διαδικασίας, όσο και για τον διαμοιρασμό των βιβλιογραφικών εγγραφών. Τα τελευταία χρόνια, η δυναμική εμφάνιση εφαρμογών τεχνητής νοημοσύνης (TN) έχει ανοίξει νέους δρόμους πειραματισμού για την αξιοποίησή της σε ποικίλες υπηρεσίες και δραστηριότητες των βιβλιοθηκών, μεταξύ των οποίων και η βιβλιογραφική επεξεργασία τεκμηρίων.

Στην εργασία παρουσιάζεται η έως τώρα πρόοδος της Εθνικής Βιβλιοθήκης της Ελλάδος (ΕΒΕ) στις πρώτες της προσπάθειες για τη βιβλιογραφική επεξεργασία υλικού των συλλογών της με τη χρήση γνωστών και ευρέως διαδεδομένων εργαλείων TN. Αναλύονται οι επιλογές που έγιναν από το καλοκαίρι του 2024 και τα αποτελέσματά τους, πριν καταλήξει η ΕΒΕ στην τρέχουσα υλοποίηση. Παράλληλα, παρουσιάζονται αποτελέσματα αυτοματοποιημένης καταλογογράφησης χωρίς ανθρώπινη παρέμβαση, καθώς και το μοντέλο που έχει επιλέξει η ΕΒΕ για τη βιβλιογραφική επεξεργασία του υλικού μέσω TN.

Στόχος είναι να αναδειχθούν οι δυνατότητες και οι προοπτικές της TN ως εργαλείου αυτοματοποιημένης επεξεργασίας υλικού, τα πλεονεκτήματα και τα προβλήματα που ενδέχεται να προκύψουν, καθώς και οι τρόποι συνεργασίας για τη μέγιστη δυνατή διάδοση ανάλογων προσπαθειών στις ελληνικές βιβλιοθήκες.

Conference paper's details in English

The paper was presented at the **31st Panhellenic Academic Libraries Conference**, which took place at Ioannina (Greece) during 22-24 of October 2025. Conference website: <https://en.31palc.lib.uoi.gr/>

Title

Human Expertise and Machine Learning: A Collaborative Model for Bibliographic Processing

Details of author / authors

Michalis Gerolimos, National Library of Greece, m.gerolimos@nlg.gr

Afroditi Fragkou, National Library of Greece, afragkou@nlg.gr

Keywords

1. National Library
2. Artificial Intelligence
3. Cataloguing
4. Bibliographic Processing
5. PowerApps

Abstract

Since the 1960s, librarians have continuously sought new ways to leverage information technology across the various services offered by libraries. Bibliographic processing of collections has consistently been one of the key areas where, due to its nature, maximum integration of technological applications is pursued—both for the execution of the process and for the sharing of bibliographic records. In recent years, the rapid emergence of artificial intelligence (AI) applications has opened up new avenues for experimentation in a wide range of library services and activities, including bibliographic processing.

This paper presents the progress made thus far by the National Library of Greece (NLG) in its initial efforts to process its collections using well-known and widely used AI tools. It analyzes the decisions made since the summer of 2024 and their outcomes, which have led to the Library's current implementation. The paper also showcases results from fully automated cataloging without human

intervention, as well as the model selected by the NLG for bibliographic processing through AI.

The aim is to highlight the potential and prospects of AI as a tool for automated content processing, advantages and challenges that may arise, and options of fostering collaboration to promote similar initiatives across Greek libraries.