



ΕΘΝΙΚΟ ΜΕΤΣΟΒΙΟ ΠΟΛΥΤΕΧΝΕΙΟ
ΣΧΟΛΗ ΗΛΕΚΤΡΟΛΟΓΩΝ ΜΗΧΑΝΙΚΩΝ ΚΑΙ ΜΗΧΑΝΙΚΩΝ ΥΠΟΛΟΓΙΣΤΩΝ
ΤΟΜΕΑΣ ΤΕΧΝΟΛΟΓΙΑΣ ΠΛΗΡΟΦΟΡΙΚΗΣ ΚΑΙ ΥΠΟΛΟΓΙΣΤΩΝ

Επίδραση των Μεγάλων Γλωσσικών Μοντέλων στο Δίκαιο

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ

του

ΑΦΟΡΔΑΚΟΥ Κ. ΝΙΚΟΛΑΟΥ

Επιβλέπων: Παναγιώτης Τσανάκας
Καθηγητής Ε.Μ.Π.

Συνεπιβλέπων: Μάριος Κόνιαρης
Ε.ΔΙ.Π. Ε.Μ.Π.

Αθήνα, Ιούλιος 2025



ΕΘΝΙΚΟ ΜΕΤΕΩΡΙΟ ΠΟΛΥΤΕΧΝΕΙΟ
ΣΧΟΛΗ ΗΛΕΚΤΡΟΛΟΓΩΝ ΜΗΧΑΝΙΚΩΝ ΚΑΙ ΜΗΧΑΝΙΚΩΝ ΥΠΟΛΟΓΙΣΤΩΝ
ΤΟΜΕΑΣ ΤΕΧΝΟΛΟΓΙΑΣ ΠΛΗΡΟΦΟΡΙΚΗΣ ΚΑΙ ΥΠΟΛΟΓΙΣΤΩΝ

Επίδραση των Μεγάλων Γλωσσικών Μοντέλων στο Δίκαιο

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ

του

ΑΦΟΡΔΑΚΟΥ Κ. ΝΙΚΟΛΑΟΥ

Επιβλέπων: Παναγιώτης Τσανάκας
Καθηγητής Ε.Μ.Π.

Συνεπιβλέπων: Μάριος Κόνιαρης
Ε.ΔΙ.Π. Ε.Μ.Π.

Εγκρίθηκε από την τριμελή εξεταστική επιτροπή την 10η Ιουλίου 2025.

(Υπογραφή)

(Υπογραφή)

(Υπογραφή)

.....
Παναγιώτης Τσανάκας
Καθηγητής Ε.Μ.Π.

.....
Ανδρέας-Γεώργιος Σταφυλοπάτης
Ομότιμος Καθηγητής Ε.Μ.Π.

.....
Γεώργιος Ματσόπουλος
Καθηγητής Ε.Μ.Π.

Αθήνα, Ιούλιος 2025



Copyright © Νικόλαος Αφορδακός, 2025.

All rights reserved. Με την επιφύλαξη παντός δικαιώματος.

Απαγορεύεται η αντιγραφή, αποθήκευση και διανομή της παρούσας εργασίας, εξ ολοκλήρου ή τμήματος αυτής, για εμπορικό σκοπό. Επιτρέπεται η ανατύπωση, αποθήκευση και διανομή για σκοπό μη κερδοσκοπικό, εκπαιδευτικής ή ερευνητικής φύσης, υπό την προϋπόθεση να αναφέρεται η πηγή προέλευσης και να διατηρείται το παρόν μήνυμα.

Το περιεχόμενο αυτής της εργασίας δεν απηχεί απαραίτητα τις απόψεις του Τμήματος, του Επιβλέποντα, ή της επιτροπής που την ενέκρινε.

ΔΗΛΩΣΗ ΜΗ ΛΟΓΟΚΛΟΠΗΣ ΚΑΙ ΑΝΑΛΗΨΗΣ ΠΡΟΣΩΠΙΚΗΣ ΕΥΘΥΝΗΣ

Με πλήρη επίγνωση των συνεπειών του νόμου περί πνευματικών δικαιωμάτων, δηλώνω ενυπογράφως ότι είμαι αποκλειστικός συγγραφέας της παρούσας Πτυχιακής Εργασίας, για την ολοκλήρωση της οποίας κάθε βοήθεια είναι πλήρως αναγνωρισμένη και αναφέρεται λεπτομερώς στην εργασία αυτή. Έχω αναφέρει πλήρως και με σαφείς αναφορές, όλες τις πηγές χρήσης δεδομένων, απόψεων, θέσεων και προτάσεων, ιδεών και λεκτικών αναφορών, είτε κατά κυριολεξία είτε βάσει επιστημονικής παράφρασης. Αναλαμβάνω την προσωπική και ατομική ευθύνη ότι σε περίπτωση αποτυχίας στην υλοποίηση των ανωτέρω δηλωθέντων στοιχείων, είμαι υπόλογος έναντι λογοκλοπής, γεγονός που σημαίνει αποτυχία στην Πτυχιακή μου Εργασία και κατά συνέπεια αποτυχία απόκτησης του Τίτλου Σπουδών, πέραν των λοιπών συνεπειών του νόμου περί πνευματικών δικαιωμάτων. Δηλώνω, συνεπώς, ότι αυτή η Πτυχιακή Εργασία προετοιμάστηκε και ολοκληρώθηκε από εμένα προσωπικά και αποκλειστικά και ότι, αναλαμβάνω πλήρως όλες τις συνέπειες του νόμου στην περίπτωση κατά την οποία αποδειχθεί, διαχρονικά, ότι η εργασία αυτή ή τμήμα της δεν μου ανήκει διότι είναι προϊόν λογοκλοπής άλλης πνευματικής ιδιοκτησίας.

(Υπογραφή)

.....

Νικόλαος Αφορδακός

Διπλωματούχος Ηλεκτρολόγος Μηχανικός και Μηχανικός Υπολογιστών Ε.Μ.Π.

10 Ιουλίου 2025

στη γυναίκα της ζωής μου, Άρτεμις

Περίληψη

Η παρούσα διπλωματική εργασία διερευνά την επίδραση των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στο δίκαιο, με έμφαση στις πρακτικές εφαρμογές, τις τεχνικές προκλήσεις και τις θεσμικές προεκτάσεις της τεχνητής νοημοσύνης στον νομικό τομέα. Αρχικά, παρουσιάζεται το τεχνολογικό και κοινωνικό πλαίσιο που επιταχύνει την υιοθέτηση των LLMs, ενώ καθορίζονται οι βασικοί στόχοι και το ερευνητικό πεδίο της μελέτης.

Η εργασία εστιάζει σε τέσσερις βασικούς άξονες: την νομική ανάκτηση πληροφορίας και ανάλυση εγγράφων, την πρόβλεψη δικαστικών αποφάσεων μέσω LLMs, τη ραγδαία εξέλιξη των «έξυπνων» δικαστηρίων και των ψηφιακών νομικών εργαλείων, και τέλος, στη συγκριτική αξιολόγηση, ανίχνευση παραισθήσεων και μέτρηση της νομικής ακρίβειας των συστημάτων. Η μεθοδολογία που ακολουθείται συνδυάζει τη βιβλιογραφική επισκόπηση πρόσφατων ερευνητικών εργασιών με την αναλυτική αξιολόγηση διεθνών περιπτώσεων μελέτης (case studies), αξιοποιώντας εξειδικευμένα benchmarks, ανοικτά σύνολα δεδομένων (datasets) και πραγματικά παραδείγματα από την Ευρώπη, την Ασία και τις Ηνωμένες Πολιτείες.

Ιδιαίτερη βαρύτητα δίνεται στη μελέτη των τεχνικών αρχιτεκτονικών (semantic search, retrieval-augmented generation, hybrid pipelines), στην κριτική επισκόπηση των μεθόδων εξήγησης και ελέγχου αλλά και στη διαχείριση θεσμικών ζητημάτων όπως ο κανονισμός AI Act και η ανάγκη για θεσμικά διαφανή συστήματα. Τα αποτελέσματα αναδεικνύουν τόσο τις τεχνολογικές δυνατότητες όσο και τους κινδύνους των LLMs στη νομική πράξη, επισημαίνοντας την ανάγκη για διαλειτουργικά σημεία αναφοράς, προσαρμογή σε πολυγλωσσικά περιβάλλοντα, ενίσχυση της ερμηνευσιμότητας και διαρκή θεσμικό έλεγχο.

Η εργασία καταλήγει σε συγκεκριμένες προτάσεις για τη μελλοντική έρευνα, δίνοντας έμφαση στη διεπιστημονική συνεργασία μηχανικών και νομικών, στην ανάπτυξη διεθνών νομικών δεδομένων και στη διαμόρφωση πρακτικών οδηγών για την υπεύθυνη υιοθέτηση της τεχνητής νοημοσύνης στον χώρο του δικαίου.

Λέξεις Κλειδιά

Μεγάλα Γλωσσικά Μοντέλα, Επεξεργασία Φυσικής Γλώσσας, Τεχνητή Νοημοσύνη και Δίκαιο, Νομική Πληροφορική, Έξυπνα Δικαστήρια, Συγκριτική Αξιολόγηση, Large Language Models, Natural Language Processing, Legal Informatics, Legal Decision, Hallucination Detection, Smart Courts

Abstract

This thesis explores the impact of Large Language Models (LLMs) on the legal domain, emphasizing practical applications, technical challenges, and institutional implications of artificial intelligence within legal practice. Initially, the technological and social context accelerating the adoption of LLMs is presented, along with a definition of the primary objectives and research scope of the study. The thesis concentrates on four main areas: legal information retrieval and document analysis, judicial decision prediction through LLMs, the rapid evolution of smart courts and digital legal tools, and finally, the comparative evaluation, hallucination detection, and measurement of legal accuracy in these systems. The methodology employed integrates a literature review of recent research articles and a detailed assessment of international case studies, utilizing specialized benchmarks, open datasets, and practical examples from Europe, Asia, and the USA. Particular emphasis is placed on examining technical architectures (semantic search, retrieval-augmented generation, hybrid pipelines), providing a critical overview of methods for explainability and control (explainability, auditing, provenance tracking), as well as managing institutional issues such as the AI Act regulation and the need for institutionally transparent systems. The results highlight both the technological potential and the risks associated with LLMs in legal practice, underscoring the necessity for interoperable benchmarks, adaptation to multilingual environments, enhancement of interpretability, and continuous institutional oversight. The thesis concludes with specific recommendations for future research, emphasizing interdisciplinary collaboration between engineers and legal professionals, the development of international legal datasets, and the creation of practical guidelines for responsibly adopting artificial intelligence within the legal field.

Keywords

Language Models, Natural Language Processing, Legal Informatics, Legal Decision, Hallucination Detection, Smart Courts

Ευχαριστίες

Η παρούσα διπλωματική εργασία σηματοδοτεί το τέλος ενός απαιτητικού, αλλά εξαιρετικά γόνιμου και συναρπαστικού ταξιδιού. Κανένα όμως ταξίδι δεν είναι μοναχικό, και η ολοκλήρωσή του δεν θα ήταν εφικτή χωρίς τη στήριξη και την καθοδήγηση ορισμένων ανθρώπων, στους οποίους νιώθω την ανάγκη να εκφράσω τη βαθιά μου ευγνωμοσύνη.

Αρχικά, θα ήθελα να ευχαριστήσω θερμά τον καθηγητή κ. Παναγιώτη Τσάνακα για την πολύτιμη ευκαιρία που μου έδωσε να εκπονήσω την εργασία μου.

Ιδιαίτερες ευχαριστίες οφείλω στον Δρ. Μάριο Κόνιαρη, του οποίου η καθοδήγηση υπήρξε καταλυτική. Η διαρκής του στήριξη, η διάθεσή του να μοιραστεί τη γνώση του και η αφοσίωσή του σε μια συνεργασία γεμάτη έμπνευση, έκαναν την ερευνητική αυτή διαδρομή όχι απλώς ευκολότερη, αλλά και απολαυστική.

Και φυσικά, ένα ξεχωριστό ευχαριστώ από καρδιάς στη γυναίκα της ζωής μου, την Άρτεμις. Ήταν εκεί από την πρώτη στιγμή του ακαδημαϊκού μου ταξιδιού, στηρίζοντάς με σε κάθε δυσκολία, ενθαρρύνοντάς με στις αμφιβολίες και πανηγυρίζοντας μαζί μου κάθε μικρή ή μεγάλη νίκη. Ήταν ο ακλόνητος βράχος μου, η πυξίδα μου όταν έχανα τον δρόμο σε όλη αυτή την απαιτητική διαδρομή.

Τέλος, τίποτα από όλα αυτά δεν θα ήταν εφικτό χωρίς τους γονείς μου. Αυτοί που όχι μόνο μου έδωσαν τις βάσεις για να φτάσω ως εδώ, αλλά και με στήριξαν με κάθε τρόπο—με την αγάπη τους, την υπομονή τους, τις αμέτρητες θυσίες τους και την ακλόνητη πίστη τους σε εμένα.

Σε αυτούς οφείλω ένα «ευχαριστώ» που οι λέξεις δεν αρκούν για να περιγράψουν.

Αθήνα, Ιούλιος 2025

Νικόλαος Αφορδακός

Περιεχόμενα

Περίληψη	3
Abstract	5
Ευχαριστίες	7
1 Εισαγωγή	17
1.1 Τεχνολογικό και κοινωνικό πλαίσιο	18
1.2 Στόχοι, πεδίο και ερευνητικά ερωτήματα της έρευνας	18
1.3 Συμβολή της μελέτης	19
1.4 Δομή της εργασίας	20
2 Μεθοδολογία έρευνας	23
2.1 Προσέγγιση βιβλιογραφικής ανασκόπησης	23
2.2 Κριτήρια επιλογής εφαρμογών και περιπτώσιολογικών μελετών	25
2.3 Πλαίσιο αξιολόγησης για νομικές εφαρμογές LLM	26
2.4 Περιορισμοί της έρευνας	26
3 Θεωρητικό Υπόβαθρο	29
3.1 Μεγάλα Γλωσσικά Μοντέλα: Αρχιτεκτονική και Λειτουργία	29
3.1.1 Σύγχρονες Αρχιτεκτονικές Μεγάλων Γλωσσικών Μοντέλων	30
3.1.2 Μεθοδολογίες Εκπαίδευσης και Εξειδίκευσης	32
3.2 Τεχνητή Νοημοσύνη στη Νομική Πληροφορική	33
3.2.1 Εξέλιξη της Νομικής Επεξεργασίας Φυσικής Γλώσσας	34
3.2.2 Νομικά Γλωσσικά Μοντέλα Εξειδικευμένων Τομέων	35
3.2.3 Αναπαράσταση Νομικής Γνώσης	36
3.3 Θεμελιώδεις Προκλήσεις	37
3.3.1 Επεξηγησιμότητα και Ερμηνευσιμότητα	37
3.3.2 Προκατάληψη και Δικαιοσύνη στην Τεχνητή Νοημοσύνη του Δικαίου	38
3.3.3 Ποιότητα και Προσβασιμότητα Νομικών Δεδομένων	40
3.4 Συμπεράσματα	41
4 Πρακτικές εφαρμογές και τεχνική υλοποίηση	43
4.1 Νομική Έρευνα και Ανάκτηση Πληροφορίας	43
4.1.1 Συστήματα Ανάκτησης Νομικών Εγγράφων	44
4.1.2 Ανάλυση Ομοιότητας Νομικών Υποθέσεων	48

4.1.3 Τεχνικά Παραδείγματα Υλοποίησης	50
4.2 Ανάλυση Συμβολαίων και Επεξεργασία Νομικών Εγγράφων	52
4.2.1 Αυτοματοποιημένη Αξιολόγηση Συμβάσεων	53
4.2.2 Αξιολόγηση Απόδοσης	57
4.3 Πρόβλεψη Δικαστικών Αποφάσεων	61
4.3.1 Μεθοδολογία Προγνωστικών Μοντέλων	62
4.3.2 Ανάλυση Εκβάσεων Υποθέσεων	66
4.3.3 Μετρικές Αξιολόγησης και Αποτελέσματα	67
4.4 Έξυπνα Δικαστήρια και Συστήματα Ψηφιακής Δικαιοσύνης	69
4.4.1 Retrieval-Augmented Generation σε Νομικά Περιβάλλοντα	70
4.4.2 Μελέτη περίπτωσης: Εφαρμογή σε διαφορετικές δικαιοδοσίες	72
4.5 Εργαλεία Νομικής Υποβοήθησης	74
4.5.1 Νομικά Chatbots και Εικονικοί Βοηθοί	75
4.5.2 Ενσωμάτωση σε Συστήματα Διαχείρισης Νομικής Πρακτικής	76
4.6 Αξιολόγηση και Διασφάλιση Ποιότητας	78
4.6.1 Μεθοδολογίες σύγκρισης επιδόσεων	79
4.6.2 Ανίχνευση και Πρόληψη Παραισθήσεων	81
4.6.3 Συγκριτική Ανάλυση Απόδοσης	82
4.6.4 Μετρήσεις Νομικής Ακρίβειας	84
4.7 Συμπεράσματα	85
5 Ρυθμιστικό και ηθικό πλαίσιο	87
5.1 Νομική Ευθύνη και Διαφάνεια	87
5.1.1 Απαιτήσεις Επεξηγησιμότητας	89
5.1.2 Μοντέλα Λογοδοσίας	90
5.2 Ρυθμιστικό Τοπίο	92
5.2.1 Ο Κανονισμός Τεχνητής Νοημοσύνης της ΕΕ και η Νομική Τεχνητή Νοημοσύνη	93
5.2.2 Διεθνείς Ρυθμιστικές Προσεγγίσεις	95
5.2.3 Υλοποίηση Συμμόρφωσης για Προγραμματιστές	97
5.3 Ηθικές Παραμέτρους	99
5.3.1 Τεχνικές Μετριασμού Αλγοριθμικών Προκαταλήψεων	100
5.3.2 Προστασία της Ιδιωτικότητας στη Νομική ΤΝ	101
5.3.3 Δίκαιη Πρόσβαση και Πολιτισμικές Διαστάσεις	102
5.4 Συνεργασία Ανθρώπου-ΤΝ στο Νομικό Πλαίσιο	104
5.4.1 Υποστήριξη της Απόφασης έναντι Αυτοματοποίησης της Απόφασης	105
5.4.2 Επαγγελματική Ευθύνη	106
5.4.3 Μοντέλα Αλληλεπίδρασης Δικαστή-Αλγορίθμου	108
5.5 Προκλήσεις Εφαρμογής Μεγάλων Γλωσσικών Μοντέλων στο Ελληνικό Νομικό Πλαίσιο	110
5.6 Συμπεράσματα	111

6 Μελλοντικές Κατευθύνσεις	113
6.1 Αντίκτυπος στην νομική εκπαίδευση και τις επαγγελματικές δεξιότητες	113
6.1.1 Ανάπτυξη διεπιστημονικών δεξιοτήτων	115
6.1.2 Ενσωμάτωση τεχνικών και νομικών προγραμμάτων σπουδών	116
6.2 Τεχνικές Εξελίξεις	117
6.2.1 Βελτιώσεις στην ανθεκτικότητα	118
6.2.2 Βελτιώσεις στην εξηγήσιμη λειτουργία	120
6.2.3 Τεχνικές προσαρμογής στον τομέα	121
6.2.4 Πολυγλωσσία και Διασυννοριακές Δυνατότητες	122
6.3 Ανάπτυξη υβριδικών συστημάτων	123
6.3.1 Πλαίσια συνεργασίας μεταξύ ανθρώπων και τεχνητής νοημοσύνης	124
6.3.2 Εξέλιξη των έξυπνων δικαστηρίων	125
6.3.3 Αρχιτεκτονικές ενσωμάτωσης	126
6.4 Προσβασιμότητα και δικαιοσύνη	127
6.4.1 Ανάλυση κόστους-οφέλους	128
6.4.2 Εφαρμογές πρόσβασης στη δικαιοσύνη: Μελλοντικές κατευθύνσεις και προκλήσεις	129
6.4.3 Ζητήματα ψηφιακού χάσματος	130
6.5 Συμπεράσματα	131
7 Επίλογος	133
7.1 Περίληψη των ευρημάτων	133
7.2 Επιπτώσεις στη νομική πρακτική	134
7.3 Μελλοντικές κατευθύνσεις της έρευνας	136
Βιβλιογραφία	148

Κατάλογος Σχημάτων

3.1	Σχηματική απεικόνιση των σταδίων προεκπαίδευσης και εξειδίκευσης (fine-tuning) του BERT.	33
3.2	Ενδεικτικές κατηγορίες προκατάληψης (bias) που εμφανίζονται στα δεδομένα (data), στους αλγορίθμους (algorithm) και στην αλληλεπίδραση με τον χρήστη (user interaction), αποτυπώνοντας τον δυναμικό βρόχο αναπαραγωγής και ενίσχυσης προκαταλήψεων σε συστήματα τεχνητής νοημοσύνης [1].	40
4.1	Διαδικασία αξιολόγησης της φάσης ανάκτησης (retrieval step) σε συστήματα Retrieval-Augmented Generation (RAG). Τα ανακτηθέντα αποσπάσματα συγκρίνονται με τα χρυσά δεδομένα (gold data) για την αποτίμηση της συνάφειας και της ποιότητας της ανάκτησης.	46
4.2	Ανάλυση χρόνου επεξεργασίας ερωτημάτων στο σύστημα LawPal, σε συνάρτηση με το επίπεδο πολυπλοκότητας των queries. Το σύστημα διατηρεί χαμηλούς χρόνους απόκρισης ακόμη και σε απαιτητικά σενάρια, γεγονός που ενισχύει τη χρηστικότητα και την αποδοτικότητα σε νομικές εφαρμογές πραγματικού χρόνου [2].	48
4.3	Παράδειγμα ζεύγους ερωτήματος (query) και υποψήφιας υπόθεσης (candidate case) από το σύνολο δεδομένων LeCaRD.[3].	50
4.4	Αρχιτεκτονική συστήματος RAG για ανάλυση νομικών εγγράφων και συμβολαίων. Το pipeline ενσωματώνει διαδικασίες ελέγχου συνάφειας (Relevance Check) και βελτιστοποίησης ερωτήματος (Query Refinement), ενισχύοντας την ακρίβεια και τη διαφάνεια των απαντήσεων [4].	54
4.5	Αποτελέσματα αξιολόγησης διαφορετικών μεθόδων ανάλυσης ομοιότητας νομικών υποθέσεων στο σύνολο δεδομένων LeCaRD, με τις βασικές μετρικές P@5, P@10, MAP και nDCG.[3].	61
4.6	Διάγραμμα ροής που απεικονίζει τους στόχους και τις απαιτήσεις τις 3 εργασίες πρόβλεψης δικαστικών αποφάσεων	63
4.7	Αναλυτικά αποτελέσματα αξιολόγησης νομικής ακρίβειας σε πραγματικές δικαστικές αποφάσεις, όπως αποτυπώνονται στο GreekBarBench για 13 LLMs και ομάδες ανθρώπινων αξιολογητών, ανά τομέα δικαίου και τύπο ερωτήματος. [5].	85

Κατάλογος Πινάκων

Κεφάλαιο 1

Εισαγωγή

Η τελευταία δεκαετία χαρακτηρίζεται από την αλματώδη πρόοδο της τεχνητής νοημοσύνης (Artificial Intelligence – AI) και, ειδικότερα, των Μεγάλων Γλωσσικών Μοντέλων (Large Language Models – LLMs), τα οποία έχουν μετασχηματίσει ριζικά το τοπίο της πληροφορικής, της επικοινωνίας, αλλά και ευρύτερα, τον τρόπο οργάνωσης της κοινωνικής και θεσμικής ζωής. Τα μεγάλα γλωσσικά μοντέλα αποτελούν σήμερα τομή για την επεξεργασία φυσικής γλώσσας, με εφαρμογές που ξεκινούν από την αυτόματη μετάφραση και την ανάλυση συναισθήματος έως την παραγωγή περιεχομένου και την εξαγωγή δομημένης γνώσης από αδόμητα δεδομένα. Η εξέλιξη αυτή έχει άμεσο αποτύπωμα στην καθημερινότητα, τις θεσμικές πρακτικές και τη λειτουργία των δημοκρατικών δομών.

Στον τομέα του δικαίου, η τεχνητή νοημοσύνη και ειδικότερα τα LLMs έχουν ήδη αρχίσει να επηρεάζουν σημαντικά τόσο τον τρόπο άσκησης της νομικής επιστήμης όσο και τη διαμόρφωση του ίδιου του θεσμικού πλαισίου. Συστήματα νομικής ανάκτησης πληροφορίας, αυτόματης ανάλυσης συμβολαίων, πρόβλεψης δικαστικών αποφάσεων και παροχής αυτοματοποιημένης νομικής υποστήριξης εντάσσονται πλέον στην καθημερινή πρακτική των νομικών επαγγελματιών. Παράλληλα, η ενσωμάτωση της τεχνητής νοημοσύνης στις δικαστικές διαδικασίες με την υιοθέτηση έξυπνων δικαστηρίων και ψηφιακών συστημάτων δικαιοσύνης αναδιαμορφώνει ριζικά τη σχέση πολίτη και κράτους, επηρεάζοντας ζητήματα διαφάνειας, πρόσβασης και λογοδοσίας.

Η ανάπτυξη και η υιοθέτηση των LLMs δεν είναι μόνο τεχνολογικό φαινόμενο, αλλά και κατ' εξοχήν κοινωνικό. Η αυτοματοποίηση της ανάλυσης νομικών κειμένων, η επιτάχυνση της λήψης αποφάσεων και η διεύρυνση της πρόσβασης σε νομική γνώση δημιουργούν νέες ευκαιρίες, αλλά ταυτόχρονα εγείρουν σημαντικές προκλήσεις. Τα ζητήματα ερμηνευσιμότητας, δικαιοσύνης, προκατάληψης, ηθικής και προστασίας δεδομένων αποκτούν κεντρική σημασία, καθώς οι τεχνολογικές καινοτομίες τείνουν να μεταβάλλουν όχι μόνο τον τρόπο με τον οποίο αποδίδεται η δικαιοσύνη, αλλά και το ίδιο το περιεχόμενο του δικαίου.

Τέλος, το διεθνές και ευρωπαϊκό ρυθμιστικό πλαίσιο (όπως η πρόσδος του EU AI Act [6]) αντανakλά την αναγκαιότητα θεσμικού ελέγχου και οριοθέτησης της τεχνητής νοημοσύνης στον νομικό τομέα, διαμορφώνοντας το μέλλον της νομικής επιστήμης με γνώμονα την ισορροπία μεταξύ καινοτομίας, διαφάνειας και προστασίας των θεμελιωδών δικαιωμάτων.

1.1 Τεχνολογικό και κοινωνικό πλαίσιο

Η τεχνητή νοημοσύνη και ειδικότερα οι τεχνολογίες επεξεργασίας φυσικής γλώσσας (Natural Language Processing) έχουν μετασχηματίσει σημαντικά το τοπίο της πληροφορικής και της επικοινωνίας τα τελευταία χρόνια. Τα Μεγάλα Γλωσσικά Μοντέλα, όπως το GPT[7], το BERT[8] και οι παραλλαγές τους, επιτρέπουν πλέον την κατανόηση, ανάλυση και παραγωγή κειμένου σε βαθμό που μέχρι πρόσφατα θεωρούνταν ανέφικτος.

Η τεχνολογική αυτή εξέλιξη δεν περιορίζεται στον ακαδημαϊκό ή τον ερευνητικό χώρο, αλλά διαχέεται ραγδαία σε κρίσιμους τομείς της κοινωνίας και της οικονομίας. Από τη δημοσιογραφία, την εκπαίδευση και τις επιχειρήσεις, μέχρι τη δημόσια διοίκηση, την υγεία και ειδικότερα το δίκαιο, τα LLMs αναμένεται να μετασχηματίσουν τον τρόπο με τον οποίο οι άνθρωποι διαχειρίζονται, κατανοούν και αξιοποιούν τη γνώση[9].

Στον χώρο του δικαίου, η πίεση για διαχείριση τεράστιων όγκων πληροφορίας, η ανάγκη για άμεση, αξιόπιστη και τεκμηριωμένη λήψη αποφάσεων, καθώς και η αυξανόμενη πολυπλοκότητα των νομικών δεδομένων, δημιουργούν πρόσφορο έδαφος για την υιοθέτηση προηγμένων τεχνολογικών λύσεων. Τα LLMs ενσωματώνονται σε ένα ευρύ φάσμα εφαρμογών, από την αυτοματοποιημένη αναζήτηση και ανάλυση νομοθεσίας, τη δομημένη επεξεργασία και ερμηνεία συμβολαίων, την πρόβλεψη εκβάσεων δικαστικών υποθέσεων, μέχρι και την παροχή αυτοματοποιημένης νομικής υποστήριξης μέσω chatbots και εικονικών βοηθών.

Η ενσωμάτωση των LLMs στο δίκαιο επηρεάζει όχι μόνο τις τεχνικές διαδικασίες αλλά και τις θεσμικές, κοινωνικές και ηθικές ισορροπίες[10]. Από τη μία πλευρά, υπόσχονται μεγαλύτερη αποδοτικότητα, διαφάνεια, προσβασιμότητα και ταχύτητα στην απονομή της δικαιοσύνης. Από την άλλη, αναδεικνύουν νέους κινδύνους: ζητήματα ερμηνευσιμότητας (explainability), μεροληψίας (bias), αξιοπιστίας των παραγόμενων πληροφοριών, προστασίας προσωπικών δεδομένων και θεσμικής λογοδοσίας παραμένουν ανοιχτά.

Η διεθνής κοινότητα επιστημονική, επαγγελματική και ρυθμιστική αναζητεί τρόπους να ενσωματώσει τις νέες τεχνολογίες με ασφαλή και υπεύθυνο τρόπο. Ενδεικτικές είναι οι εξελίξεις στην ευρωπαϊκή νομοθεσία [6], οι δοκιμαστικές εφαρμογές έξυπνων δικαστηρίων σε κράτη όπως η Κίνα [11] και οι Φιλιππίνες [12], οι αυστηρότερες απαιτήσεις για auditing και εξηγήσιμη τεχνητή νοημοσύνη [13], αλλά και η διαρκής συζήτηση για τη θέση και τον ρόλο του ανθρώπου στη λήψη κρίσιμων νομικών αποφάσεων[14].

Συνολικά, το τεχνολογικό και κοινωνικό πλαίσιο εντός του οποίου εξετάζεται η επίδραση των LLMs στο δίκαιο χαρακτηρίζεται από δυναμισμό, διαρκή εξέλιξη και διαρκώς αυξανόμενες απαιτήσεις για τεχνική και θεσμική επάρκεια. Σε αυτό το πλαίσιο, η συστηματική μελέτη των πρακτικών εφαρμογών, των τεχνικών λύσεων και των προκλήσεων που συνοδεύουν τα LLMs καθίσταται επιτακτική για την επιστημονική κοινότητα, τους επαγγελματίες του χώρου και τους ρυθμιστικούς φορείς[15].

1.2 Στόχοι, πεδίο και ερευνητικά ερωτήματα της έρευνας

Η παρούσα διπλωματική εργασία έχει ως βασικό στόχο τη συστηματική διερεύνηση της επίδρασης των Μεγάλων Γλωσσικών Μοντέλων στον χώρο του δικαίου, εστιάζοντας τόσο στις τεχνικές όσο και στις θεσμικές, επιστημονικές και κοινωνικές διαστάσεις αυτής της αλλη-

λεπίδρασης. Η ανάλυση επικεντρώνεται στις πρακτικές εφαρμογές των LLMs σε διάφορους άξονες της νομικής επιστήμης και πράξης — όπως η ανάκτηση πληροφορίας, η ανάλυση συμβολαίων, η πρόβλεψη δικαστικών αποφάσεων, η αυτοματοποιημένη νομική υποστήριξη και η ενσωμάτωσή τους σε έξυπνα δικαστήρια και ψηφιακές δομές απονομής δικαιοσύνης.

Για την εκπλήρωση των παραπάνω στόχων, διατυπώνονται τα ακόλουθα βασικά ερευνητικά ερωτήματα, τα οποία καθοδηγούν τη δομή, τη μεθοδολογία και το περιεχόμενο της μελέτης:

1. Ποια είναι τα βασικά αρχιτεκτονικά και επιχειρησιακά χαρακτηριστικά των Μεγάλων Γλωσσικών Μοντέλων (LLMs) και οι εξειδικευμένες παραλλαγές τους, στις εφαρμογές του δικαίου;
2. Πώς διαφοροποιούνται τα νομικά LLMs από τα γενικής χρήσης γλωσσικά μοντέλα και ποιες είναι οι βασικές τεχνικές και θεσμικές προκλήσεις που προκύπτουν;
3. Πόσο αποτελεσματικά είναι τα LLMs σε βασικές νομικές εργασίες, όπως η νομική έρευνα, η ανάλυση συμβολαίων, η πρόβλεψη δικαστικών αποφάσεων και η αυτοματοποιημένη νομική υποστήριξη, και πώς αξιολογείται η τεχνική και θεσμική απόδοσή τους;
4. Ποιος είναι ο ρόλος και ο βαθμός εφαρμογής των LLMs σε συστήματα έξυπνων δικαστηρίων (smart courts) και ψηφιακής δικαιοσύνης σε διαφορετικές έννομες τάξεις;
5. Ποιες είναι οι σύγχρονες μεθοδολογίες αξιολόγησης της απόδοσης, της αξιοπιστίας, της ακρίβειας και της ερμηνευσιμότητας των LLMs στη νομική πράξη, με έμφαση στη μέτρηση παραισθήσεων, νομικής ακρίβειας και συνέπειας παραπομπών (citation consistency);
6. Ποιες προκλήσεις και όρια αναδεικνύουν οι ρυθμιστικές, ηθικές και πολιτισμικές παράμετροι (διαφάνεια, λογοδοσία, προστασία δεδομένων) στην ενσωμάτωση των LLMs στο δίκαιο;
7. Ποιες είναι οι τάσεις που διαμορφώνονται για τη μελλοντική εξέλιξη του πεδίου, την εκπαίδευση των νομικών επαγγελματιών, τη συνεργασία ανθρώπινης κρίσης και τεχνολογίας και την επίδραση των LLMs στην πρόσβαση στη δικαιοσύνη;

Τα ερωτήματα αυτά διατρέχουν οριζόντια τη συνολική δομή της εργασίας, με κάθε κεφάλαιο να εστιάζει σε συγκεκριμένες πτυχές τους, προσπαθώντας να προσφέρει τεκμηριωμένες και κριτικές απαντήσεις βάσει της πλέον σύγχρονης βιβλιογραφίας, των τεχνικών εξελίξεων και των θεσμικών μετασχηματισμών που παρατηρούνται διεθνώς.

1.3 Συμβολή της μελέτης

Η συμβολή της παρούσας διπλωματικής εργασίας έγκειται στη διεξοδική αποτίμηση της επίδρασης των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στο δίκαιο, με ολιστική προσέγγιση που καλύπτει τόσο τις τεχνολογικές και μηχανικές πτυχές όσο και τις θεσμικές, επιστημονικές και κοινωνικές παραμέτρους. Η εργασία προσφέρει μία συστηματική και συγκριτική

επισκόπηση των κυριότερων τεχνολογικών αρχιτεκτονικών, των τεχνικών υλοποιήσεων και των benchmarking frameworks που εφαρμόζονται σε νομικά περιβάλλοντα, συνδυάζοντας τη διεθνή βιβλιογραφία με πρακτικά παραδείγματα εφαρμογής.

Ιδιαίτερη έμφαση δίνεται στην τεχνική ανάλυση και την αξιολόγηση μεθοδολογιών που αφορούν τη νομική ανάκτηση πληροφορίας, την αυτοματοποιημένη επεξεργασία νομικών εγγράφων, καθώς και στην πρόβλεψη δικαστικών αποφάσεων με τη χρήση σύγχρονων LLMs. Εξετάζονται εκτενώς οι αρχιτεκτονικές και η λειτουργία των έξυπνων δικαστηρίων (smart courts), η πρακτική ενσωμάτωση καινοτόμων εργαλείων όπως τα νομικά chatbots και τα συστήματα υποβοήθησης διαχείρισης νομικής πρακτικής, ενώ αναλύονται τα κριτήρια συγκριτικής αξιολόγησης, η ανίχνευση παραισθήσεων και η μέτρηση της νομικής ακρίβειας των αποτελεσμάτων των LLMs.

Παράλληλα, η εργασία φωτίζει τις θεσμικές και ρυθμιστικές προκλήσεις που ανακύπτουν από την ευρεία χρήση των LLMs στο δίκαιο, τόσο σε ευρωπαϊκό επίπεδο (π.χ. EU AI Act) όσο και σε διεθνές περιβάλλον, αναδεικνύοντας τις εξελίξεις και τα διλήμματα που σχετίζονται με τη διαφάνεια, την υπευθυνότητα και τα δικαιώματα των διαδίκων. Η σύνδεση της τεχνικής ανάλυσης με τις πρακτικές καινοτομίες αποδεικνύεται καθοριστική, ενώ σκιαγραφούνται και μελλοντικές ερευνητικές κατευθύνσεις για τη βελτίωση τόσο της τεχνικής αξιοπιστίας όσο και των ηθικών και θεσμικών διαστάσεων της τεχνητής νοημοσύνης στο δίκαιο.

Συνολικά, η εργασία επιχειρεί να γεφυρώσει το χάσμα μεταξύ τεχνικής και νομικής επιστήμης, προσφέροντας ένα πλήρες και πρακτικά αξιοποιήσιμο πλαίσιο κατανόησης, αξιολόγησης και μελλοντικής εξέλιξης των LLMs στη δικαστική πράξη.

1.4 Δομή της εργασίας

Η παρούσα διπλωματική εργασία είναι οργανωμένη σε επτά διακριτά κεφάλαια, καθένα από τα οποία εστιάζει σε συγκεκριμένες πτυχές της επίδρασης των Μεγάλων Γλωσσικών Μοντέλων στον χώρο του δικαίου, υιοθετώντας μια διαθεματική, τεχνικά τεκμηριωμένη και θεσμικά προσανατολισμένη προσέγγιση.

- **Κεφάλαιο 1 – Εισαγωγή:** Θέτει το ευρύτερο τεχνολογικό και κοινωνικό πλαίσιο, παρουσιάζοντας την αλματώδη εξέλιξη των LLMs και τη σημασία τους για τον νομικό τομέα. Διατυπώνει με σαφήνεια τους στόχους, τα βασικά ερευνητικά ερωτήματα και τη συμβολή της μελέτης, ενώ δίνει συνοπτική εικόνα της διάρθρωσης του έργου.
- **Κεφάλαιο 2 – Μεθοδολογία Έρευνας:** Περιγράφει αναλυτικά τη μεθοδολογική προσέγγιση που ακολουθήθηκε, όπως τη συστηματική βιβλιογραφική ανασκόπηση, την επιλογή εξειδικευμένων τεχνικών frameworks (π.χ. LegalBench, LexGLUE, LeCaRD, LegalRAG) και τη χρήση benchmarking μετρικών για την αξιολόγηση LLMs. Επεξηγεί τα κριτήρια επιλογής case studies, datasets και εφαρμογών, το πλαίσιο αξιολόγησης για εφαρμογές τεχνητής νοημοσύνης στον χώρο του δικαίου, καθώς και τα μεθοδολογικά όρια και τις προκλήσεις της έρευνας.
- **Κεφάλαιο 3 – Θεωρητικές και Τεχνικές Βάσεις:** Αναλύει διεξοδικά τις αρχιτεκτονικές των LLMs (Transformers, attention mechanisms, fine-tuning, transfer learning),

εστιάζοντας τόσο στα γενικά μοντέλα (GPT, BERT, T5) όσο και στις νομικές παραλλαγές (LegalBERT, Lawformer, CaseLawBERT). Εξετάζει τις διαφορές με τα μοντέλα γενικής χρήσης, αναλύει τις απαιτήσεις για επεξεργασία νομικής γλώσσας (legal NLP), και παρουσιάζει τις σύγχρονες προκλήσεις της εξηγήσιμης τεχνητής νοημοσύνης, της μεροληψίας στα νομικά δεδομένα και των απαιτήσεων για διαφάνεια και ποιότητα επισημείωσης (annotation).

- **Κεφάλαιο 4 – Εφαρμογή και υλοποίηση:** Εστιάζει στις πρακτικές εφαρμογές και την τεχνική υλοποίηση των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στον χώρο του δικαίου. Παρουσιάζονται χαρακτηριστικά παραδείγματα συστημάτων νομικής έρευνας, ανάλυσης εγγράφων και πρόβλεψης δικαστικών αποφάσεων, ενώ αναλύονται οι τεχνικές, τα εργαλεία και οι μετρικές αξιολόγησης που χρησιμοποιούνται στη σύγχρονη νομική τεχνολογία. Ιδιαίτερη έμφαση δίνεται στην κριτική αποτίμηση των αποτελεσμάτων και στις προκλήσεις διαλειτουργικότητας μεταξύ συστημάτων.
- **Κεφάλαιο 5 – Ρυθμιστικό και Ηθικό Πλαίσιο:** Αξιολογεί το υπάρχον και αναδυόμενο θεσμικό περιβάλλον, εξετάζοντας ζητήματα διαφάνειας, λογοδοσίας, μεροληψίας, προστασίας προσωπικών δεδομένων, και compliance με ευρωπαϊκά και διεθνή στάνταρδς (EU AI Act, global regulatory trends). Αναδεικνύει τις ηθικές προκλήσεις και τις πρακτικές διασφάλισης του ανθρωποκεντρικού ελέγχου στη λήψη νομικών αποφάσεων.
- **Κεφάλαιο 6 – Μελλοντικές Κατευθύνσεις:** Εστιάζει στις αναδυόμενες τεχνολογικές τάσεις, τη μελλοντική εκπαίδευση των νομικών επαγγελματιών, τις προοπτικές για υβριδικά συστήματα ανθρώπου-τεχνητής νοημοσύνης, τις εξελίξεις στην ερμηνευσιμότητα, την ανθεκτικότητα και τα «έξυπνα» δικαστήρια, καθώς και στην εξισορρόπηση μεταξύ κόστους, προσβασιμότητας και ποιότητας στη δικαιοσύνη.
- **Κεφάλαιο 7 – Συμπεράσματα:** Συνοψίζει τα βασικά ευρήματα κάθε ενότητας, αξιολογεί τις επιπτώσεις των LLMs στην νομική πρακτική και την τεχνική ανάπτυξη, επισημαίνει τα όρια της υπάρχουσας γνώσης και προτείνει συγκεκριμένες μελλοντικές ερευνητικές κατευθύνσεις, τόσο για το διεθνές όσο και για το ελληνικό πλαίσιο.

Μεθοδολογία έρευνας

Η συστηματική διερεύνηση της επίδρασης των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στο δίκαιο προϋποθέτει την υιοθέτηση μιας πολυπαραγοντικής μεθοδολογικής προσέγγισης, η οποία συνδυάζει θεωρητική ανάλυση, βιβλιογραφική ανασκόπηση, τεχνική αποτίμηση και συγκριτική αξιολόγηση. Το παρόν κεφάλαιο αναλύει τη μεθοδολογική προσέγγιση που ακολουθήθηκε για τη διερεύνηση της επίδρασης των Μεγάλων Γλωσσικών Μοντέλων στον νομικό τομέα.

Η εργασία βασίζεται σε εκτενή βιβλιογραφική ανασκόπηση της διεθνούς επιστημονικής και τεχνικής βιβλιογραφίας, με έμφαση σε πρόσφατα peer-reviewed άρθρα, θεσμικά αποδεκτά πρότυπα αξιολόγησης, επίσημα σύνολα δεδομένων και εφαρμογές LLMs σε πραγματικά νομικά περιβάλλοντα. Ιδιαίτερη προσοχή δόθηκε στην επιλογή case studies που καλύπτουν διαφορετικά κράτη και νομικά πλαίσια, ώστε να αναδειχθούν τόσο οι τεχνικές όσο και οι θεσμικές και πολιτισμικές ιδιαιτερότητες.

Για την αποτίμηση της τεχνικής απόδοσης και αξιοπιστίας των LLMs αξιοποιήθηκαν εξειδικευμένα πλαίσια συγκριτικής αξιολόγησης (benchmarking frameworks) και μετρικές, όπως η ανάκληση (recall), ο κανονικοποιημένος βαθμός αθροιστικής κατάταξης (nDCG), η αληθοφάνεια (faithfulness), το ποσοστό παραισθήσεων (hallucination rate), η ακρίβεια παραπομπών (citation accuracy) και ο μακρο-μέσος όρος F1 (macro-F1). Παράλληλα, η πρακτική χρησιμότητα των συστημάτων εξετάστηκε σε συνάρτηση με τις ανάγκες του νομικού επαγγέλματος και τις ιδιαιτερότητες κάθε έννομης τάξης.

Αρχικά, παρουσιάζεται η προσέγγιση βιβλιογραφικής ανασκόπησης, με αναλυτική περιγραφή των εργαλείων αναζήτησης και αξιολόγησης της επιστημονικής πληροφορίας. Στη συνέχεια, εξειδικεύονται τα κριτήρια επιλογής των υπό εξέταση εφαρμογών και μελετών περίπτωσης, με έμφαση στην τεχνική και θεσμική τους συνάφεια. Ακολούθως περιγράφεται το πλαίσιο αξιολόγησης, αναλύοντας τα κύρια benchmarking frameworks, τις μετρικές και τις συγκριτικές μεθόδους που χρησιμοποιήθηκαν. Τέλος, επισημαίνονται οι βασικοί περιορισμοί της έρευνας που απορρέουν από τη φύση των δεδομένων, την επιλογή των case studies, τις τεχνικές μεθόδους αξιολόγησης και το ρυθμιστικό περιβάλλον.

2.1 Προσέγγιση βιβλιογραφικής ανασκόπησης

Βασικός στόχος είναι η εντοπισμένη κάλυψη του συνόλου της σύγχρονης επιστημονικής γνώσης σχετικά με τα Μεγάλα Γλωσσικά Μοντέλα και τις εφαρμογές τους στον χώρο του

δικαίου.

Αρχικά, ακολουθήθηκε μια πολυεπίπεδη στρατηγική αναζήτησης σε διεθνείς επιστημονικές βάσεις δεδομένων, όπως το IEEE Xplore, ScienceDirect, arXiv, Google Scholar, ResearchGate, καθώς και σε εξειδικευμένα repositories για νομική πληροφορική (Legal Informatics). Ειδική μέριμνα ελήφθη για τη συλλογή peer-reviewed άρθρων, επίσημων τεχνικών reports, διατριβών, και εγχειριδίων που εκδόθηκαν την τελευταία δεκαετία, με έμφαση στα έργα από το 2022 και μετά, λόγω της εκρηκτικής εξέλιξης των LLMs.

Η επιλογή των πηγών έγινε με βάση συγκεκριμένα keywords (π.χ. legal NLP, LLM benchmarks, contract analysis with transformers, legal outcome prediction, explainable AI in law).

Ιδιαίτερη προσοχή δόθηκε :

- **Στη διαθεματικότητα:** Ενσωματώθηκαν πηγές τόσο από το τεχνικό/μηχανικό πεδίο (computer science, machine learning, NLP) όσο και από το νομικό-θεσμικό (legal informatics, law & technology), διασφαλίζοντας πολυεπιστημονική κάλυψη.
- **Στη διασφάλιση ποιότητας:** Οι πηγές αξιολογήθηκαν ως προς το impact factor των περιοδικών, τη συχνότητα αναφορών, τη μεθοδολογική διαφάνεια και τη χρησιμότητα των δεδομένων.
- **Στην αντιπροσωπευτικότητα των εφαρμογών:** Εντάχθηκαν έργα από διαφορετικές έννομες τάξεις (αγγλοσαξονικό, ηπειρωτικό, ασιατικό, ελληνικό δίκαιο) καθώς και περιπτωσιολογικές μελέτες από δημόσια projects, εταιρικές εφαρμογές και πειραματικές πλατφόρμες.

Συνολικά, η βιβλιογραφική ανασκόπηση οργανώθηκε θεματικά και χρονολογικά, με κατηγοριοποίηση σε :

- Θεωρητικές βάσεις και αρχιτεκτονικές LLMs.
- Εξειδικευμένες νομικές εφαρμογές (legal research, contract analysis, outcome prediction).
- Μετρικές και πλαίσια αξιολόγησης απόδοσης (benchmarking metrics and frameworks)
- Ζητήματα αξιολόγησης, ερμηνευσιμότητας (explainability), και ρυθμιστικού πλαισίου.

Για τη διαχείριση και τεκμηρίωση των πηγών χρησιμοποιήθηκε εργαλείο βιβλιογραφικής οργάνωσης (Mendeley Reference Manager), επιτρέποντας συστηματική καταγραφή, επισήμανση σημείων ενδιαφέροντος και εξαγωγή των σχετικών παραπομπών.

Η συστηματική αυτή προσέγγιση διασφαλίζει ότι η εργασία καλύπτει σφαιρικά το διεθνές επιστημονικό πεδίο, ενσωματώνει τις πιο πρόσφατες και τεχνικά αξιόπιστες εξελίξεις και επιτρέπει την εξαγωγή ουσιαστικών συμπερασμάτων για την επίδραση των LLMs στον χώρο του δικαίου.

2.2 Κριτήρια επιλογής εφαρμογών και περιπτώσιολογικών μελετών

Η επιλογή των εφαρμογών, των μοντέλων και των περιπτώσιολογικών μελετών (case studies) που συμπεριλαμβάνονται στην εργασία δεν έγινε αυθαίρετα, αλλά στηρίχθηκε σε ένα σύνολο κριτηρίων που εξασφαλίζουν τόσο την τεχνική αρτιότητα όσο και τη θεσμική συνάφεια. Στόχος ήταν η αποτύπωση της πραγματικής πολυπλοκότητας, της λειτουργικής ποικιλίας και της επιστημονικής εγκυρότητας των σύγχρονων νομικών εφαρμογών των LLMs, μέσα από παραδείγματα που πληρούν αυστηρές προδιαγραφές επιλογής, όπως:

- **Τεχνική αρτιότητα και επικαιρότητα:** Προτιμήθηκαν εφαρμογές και μελέτες που βασίζονται σε state-of-the-art μοντέλα (GPT-4, BERT, LegalBERT, Lawformer, CaseLaw-BERT), αξιοποιώντας τις πιο πρόσφατες τεχνολογικές εξελίξεις στο πεδίο της μηχανικής μάθησης και της επεξεργασίας φυσικής γλώσσας.
- **Αντιπροσωπευτικότητα νομικών tasks:** Εντάχθηκαν παραδείγματα που καλύπτουν το πλήρες φάσμα των βασικών νομικών λειτουργιών, όπως η νομική έρευνα (legal research), η ανάλυση συμβολαίων (contract analysis), η κατηγοριοποίηση και εξαγωγή πληροφορίας (document classification, information extraction), η πρόβλεψη δικαστικών αποφάσεων (judicial decision prediction), τα αυτόνομα συστήματα νομικής υποστήριξης (legal assistance tools) και οι εφαρμογές σε «έξυπνα» δικαστήρια (smart courts).
- **Διαφορετικότητα δικαιοδοσιών και θεσμικών πλαισίων:** Συμπεριλήφθηκαν case studies και συστήματα από έννομες τάξεις με διαφορετικό θεσμικό υπόβαθρο, προκειμένου να αποτυπωθεί η επίδραση των LLMs σε ποικίλα ρυθμιστικά, πολιτισμικά και γλωσσικά περιβάλλοντα.
- **Προσβασιμότητα δεδομένων και benchmarking:** Επελέγησαν εφαρμογές που αξιοποιούν ανοιχτά, επισήμως τεκμηριωμένα και συχνά χρησιμοποιούμενα datasets (LexGLUE, LeCaRD, ECtHR, CUAD, LEDGAR, LegalBench-RAG), επιτρέποντας αξιόπιστη σύγκριση και επαλήθευση των αποτελεσμάτων.
- **Κριτήρια πρακτικής και θεσμικής συνάφειας:** Εντάχθηκαν εφαρμογές και μελέτες που συνδέονται με πραγματικές ανάγκες του νομικού επαγγέλματος, της δικαστικής πρακτικής και της δημόσιας διοίκησης, όπως π.χ. αυτοματοποιημένη νομική τεκμηρίωση, υποστήριξη λήψης απόφασης, ανάλυση νομολογίας ή διαχείριση νομικών ροών εργασίας.
- **Κριτική αποτίμηση:** Δόθηκε προτεραιότητα σε εφαρμογές και μελέτες που τεκμηριώνουν όχι μόνο τα πλεονεκτήματα, αλλά και τις τεχνικές, ρυθμιστικές ή ηθικές προκλήσεις (όπως bias, hallucinations, explainability, data privacy).

Στην εργασία συμπεριλαμβάνονται και παραδείγματα από την ελληνική πραγματικότητα ή από συστήματα που μπορούν να αξιοποιηθούν/προσαρμοστούν στο ελληνικό νομικό

περιβάλλον, ώστε να αναδειχθούν τα ειδικά ζητήματα εφαρμογής των LLMs σε χώρες με περιορισμένα datasets και ιδιαίτερες θεσμικές απαιτήσεις. Τέλος, σε κάθε περίπτωση, η επιλογή έγινε με στόχο τη διασφάλιση της τεχνικής πληρότητας, της θεσμικής εγκυρότητας και της δυνατότητας εξαγωγής συγκρίσιμων και επαληθεύσιμων συμπερασμάτων για την επίδραση των LLMs στον νομικό τομέα.

2.3 Πλαίσιο αξιολόγησης για νομικές εφαρμογές LLM

Η αξιολόγηση των νομικών εφαρμογών των Μεγάλων Γλωσσικών Μοντέλων στηρίζεται σε ένα πολυεπίπεδο και τεχνικά εστιασμένο πλαίσιο, το οποίο λαμβάνει υπόψη τόσο τις απαιτήσεις της μηχανικής μάθησης όσο και τις ιδιαιτερότητες του νομικού κλάδου. Η επιλογή των benchmarking μεθοδολογιών και των μετρικών έγινε με στόχο τη μέγιστη ακρίβεια, ερμηνευσιμότητα και επιστημονική εγκυρότητα των αποτελεσμάτων.

Για την αποτίμηση της απόδοσης των LLMs αξιοποιούνται διεθνώς αναγνωρισμένα benchmarking frameworks που έχουν σχεδιαστεί ειδικά για νομικά tasks, όπως το LegalBench-RAG (για Retrieval-Augmented Generation), το LexGLUE (για classification, outcome prediction και statute classification), το LeCaRD (για retrieval tasks), το ContractNLI (για entailment σε νομικά συμβόλαια), το CUAD (για clause extraction) και το LegalRAG (για αξιολόγηση faithfulness και citation accuracy σε real-world pipelines).

Σε εργασίες όπου η αυτόματη αξιολόγηση δεν αρκεί, ενσωματώνονται διαδικασίες αξιολόγησης από εξειδικευμένους νομικούς (expert review panels), τόσο για την ορθότητα όσο και για τη νομική πληρότητα των απαντήσεων. Η συνδυαστική χρήση automated και human-centered μετρικών ενισχύει την αξιοπιστία και τη ρεαλιστικότητα των αποτελεσμάτων. Το πλαίσιο αξιολόγησης περιλαμβάνει σύγκριση διαφορετικών μοντέλων (general-purpose έναντι legal-tuned LLMs), αρχιτεκτονικών (dense, sparse, hybrid retrievers), reranking μεθόδων (BM25, MMR, cross-encoders), αλλά και σύγκριση σε πολλαπλές γλώσσες και νομικά περιβάλλοντα (π.χ. σύγκριση επίδοσης σε αγγλικά, κινεζικά, ελληνικά datasets). Επιπλέον, λαμβάνονται υπόψη οι περιορισμοί των benchmarks, όπως ο περιορισμός στα αγγλικά δεδομένα, η διαθεσιμότητα πρότυπων σχολιασμένων δεδομένων gold standard annotations και η υποκειμενικότητα στην ανθρώπινη αξιολόγηση.

2.4 Περιορισμοί της έρευνας

Παρά τη συστηματική και πολυπαραγοντική μεθοδολογική προσέγγιση που υιοθετήθηκε στην παρούσα εργασία, εξακολουθούν να υφίστανται σημαντικοί περιορισμοί που επηρεάζουν τη γενικευσιμότητα, τη συγκρισιμότητα και την πρακτική αξιοποίηση των ευρημάτων.

Παρακάτω αναλύονται οι ακόλουθοι περιορισμοί:

1. **Περιορισμοί στα δεδομένα και στα διαθέσιμα πλαίσια αξιολόγησης (benchmarks):** Η πλειονότητα των δημοσίως διαθέσιμων συνόλων δεδομένων (datasets) και πλαισίων αξιολόγησης (benchmarks) έχει διαμορφωθεί με γνώμονα τα αγγλικά και τις αγγλοσαξονικές έννομες τάξεις (common law), παραμερίζοντας τη νομική, γλωσσική και θεσμική ποικιλομορφία που χαρακτηρίζει άλλες περιοχές (όπως ηπειρωτικές,

βαλκανικές ή ασιατικές δικαιοδοσίες). Αυτό έχει ως αποτέλεσμα να υπάρχει περιορισμένη αντιπροσωπευτικότητα για νομικά συστήματα με διαφορετική δομή, γλώσσα, νομοθεσία ή βαθμό ψηφιακής ωριμότητας. Επιπλέον, πολλά σύνολα δεδομένων περιλαμβάνουν μόνο επιφανειακή σήμανση (shallow annotation) ή παρουσιάζουν έλλειψη σταθερών και διαφανών κανόνων σχολιασμού, δυσχεραίνοντας την αξιόπιστη εκπαίδευση και αξιολόγηση μοντέλων σε πολύπλοκα tasks, όπως η επεξήγηση νομικής επιχειρηματολογίας ή η συσχέτιση απόφασης και πραγματικών περιστατικών.

2. **Υποκειμενικότητα, bias & consistency στην ανθρώπινη αξιολόγηση:** Αν και η ανθρώπινη αξιολόγηση (από νομικούς ή ερευνητές) παραμένει κρίσιμη για τη μέτρηση της νομικής πληρότητας και της αληθοφάνειας (faithfulness), εντούτοις χαρακτηρίζεται από υψηλό βαθμό υποκειμενικότητας, διαφορετικά στανταρδς, αλλά και τη δυνατότητα εμφάνισης προκαταλήψεων (bias) ανάλογα με το νομικό και πολιτισμικό υπόβαθρο του αξιολογητή. Η βιβλιογραφία καταγράφει σημαντική διακύμανση (μέχρι και 20–25%) στα αποτελέσματα μεταξύ αξιολογητών, ειδικά σε ερωτήματα με νομική ή πραγματική ασάφεια. Παράλληλα, τα υπάρχοντα εργαλεία αυτοματοποιημένης αξιολόγησης (π.χ. faithfulness checkers) έχουν περιορισμένη δυνατότητα ανίχνευσης λεπτών νομικών αποχρώσεων ή αμφισχημικών ερμηνειών.
3. **Τεχνικά και μεθοδολογικά όρια των πλαισίων αξιολόγησης και των μοντέλων:** Τα πιο εξελιγμένα πλαίσια αξιολόγησης (LegalBench-RAG, RAGAS, LeCoDe) εστιάζουν κυρίως σε συγκεκριμένους τύπους εργασιών (tasks) και συνήθως αγνοούν ή απλοποιούν τη ροή πραγματικών νομικών διαδικασιών, ενώ πολλά νομικά tasks (όπως το cross-jurisdictional legal reasoning ή η ανάλυση της εξέλιξης νομολογίας) δεν υποστηρίζονται ακόμη από επίσημα benchmarks. Η αδιαφάνεια των εξελιγμένων μοντέλων δυσκολεύει τον εντοπισμό λαθών, την ερμηνεία των αποτελεσμάτων (explainability) και την ανάλυση των failure modes, περιορίζοντας τη δυνατότητα επιστημονικής αναπαραγωγής και τεχνικής βελτίωσης.
4. **Εξάρτηση από εμπορικά APIs και περιορισμένη διαφάνεια:** Η ευρεία χρήση εμπορικών υπηρεσιών (commercial LLM APIs), με κλειστά δεδομένα εκπαίδευσης, αδιαφανείς μηχανισμούς αξιολόγησης και άγνωστα εσωτερικά κριτήρια (internal benchmarks), περιορίζει την επιστημονική διαφάνεια και τον ανοιχτό έλεγχο των αποτελεσμάτων. Η πρόσβαση σε λεπτομερή δεδομένα εκπαίδευσης, logs ή failure cases συνήθως δεν είναι δυνατή, γεγονός που δυσχεραίνει τόσο τον έλεγχο όσο και τη βελτιστοποίηση των συστημάτων σε πρακτικό επίπεδο.
5. **Ταχύτητα τεχνολογικής εξέλιξης, διαχρονικότητα και sustainability:** Το ερευνητικό πεδίο των Μεγάλων Γλωσσικών Μοντέλων (LLMs) βρίσκεται σε φάση εκθετικής ανάπτυξης. Νέα μοντέλα, τεχνικές και φραμεворκς εμφανίζονται διαρκώς, με αποτέλεσμα τα δεδομένα, τα benchmarks ή ακόμη και τα best practices που περιγράφονται στην παρούσα εργασία να κινδυνεύουν να καταστούν σύντομα παρωχημένα. Η διαχρονική εγκυρότητα των συμπερασμάτων απαιτεί διαρκή επικαιροποίηση, επαναξιολόγηση και προσαρμογή στα τρέχοντα επιστημονικά και θεσμικά δεδομένα.

6. **Περιορισμοί πρόσβασης, διαλειτουργικότητας και θεσμικά εμπόδια:** Η πρόσβαση σε πλήρη νομικά δεδομένα (όπως πλήρεις δικαστικές αποφάσεις, αναλυτικά συμβόλαια, πλήρη σώματα νομολογίας ή διοικητικά δεδομένα) παραμένει περιορισμένη σε πολλές έννομες τάξεις, είτε λόγω πνευματικών δικαιωμάτων, είτε λόγω προσωπικών δεδομένων, είτε εξαιτίας θεσμικών περιορισμών στην κυκλοφορία, αποθήκευση ή διαμοιρασμό των δεδομένων. Αυτό καθιστά δυσχερή την πρακτική αξιοποίηση, τη διαλειτουργικότητα και τη σύγκριση μοντέλων ή συστημάτων σε διεθνές επίπεδο.
7. **Ζητήματα διαφάνειας, αλγοριθμικής δικαιοσύνης (fairness) και ηθικής αξιολόγησης:** Η ανίχνευση και η αποτίμηση της αλγοριθμικής προκατάληψης (algorithmic bias), της ισότητας πρόσβασης, της επαληθευσιμότητας (verifiability) και της διαφάνειας (transparency) δεν έχουν ακόμη τυποποιηθεί πλήρως στη νομική πληροφορική. Τα υπάρχοντα εργαλεία και μετρικές συχνά αγνοούν τα διαφορετικά στάνταρδς, τους κοινωνικούς κινδύνους και τα ειδικά χαρακτηριστικά των επιμέρους δικαιοδοσιών, αφήνοντας σημαντικά ερωτήματα ανοιχτά για τη μελλοντική αξιολόγηση και αποδοχή των λύσεων LLM.

Συνοψίζοντας, οι παραπάνω περιορισμοί καταδεικνύουν μια επιτακτική ανάγκη, την ανάπτυξη ανοιχτών και διαλειτουργικών δεδομένων και πλαισίων αξιολόγησης που να σέβονται τη θεσμική πολυμορφία. Παράλληλα, ο εμπλουτισμός της τεχνικής αξιολόγησης με αρχές όπως η ανθρώπινη κριτική, η δικαιοσύνη και η διαφάνεια είναι απαραίτητος για τη βιώσιμη ενσωμάτωση αυτών των τεχνολογιών στο δίκαιο.

Κεφάλαιο 3

Θεωρητικό Υπόβαθρο

Στο παρόν κεφάλαιο παρέχεται το απαραίτητο θεωρητικό υπόβαθρο που αφορά τεχνολογίες και μεθοδολογίες οι οποίες αξιοποιούνται στην εκπόνηση της εργασίας.

Αρχικά, παρουσιάζονται τα Μεγάλα Γλωσσικά Μοντέλα (Large Language Models), αναλύεται η αρχιτεκτονική τους και η λειτουργία τους στο πλαίσιο της Επεξεργασίας Φυσικής Γλώσσας (Natural Language Processing). Στη συνέχεια, γίνεται αναφορά σε βασικά μοντέλα όπως οι Μετασχηματιστές (Transformers), το GPT[7] και το BERT [16] και εξετάζονται οι διαδικασίες εκπαίδευσης (training) και εξειδίκευσης (fine-tuning) που ακολουθούνται.

Η επόμενη ενότητα επικεντρώνεται στη χρήση της Τεχνητής Νοημοσύνης (Artificial Intelligence) στο πεδίο της Νομικής Πληροφορικής (Legal Informatics), παρουσιάζοντας την εξέλιξη της Νομικής Επεξεργασίας Φυσικής Γλώσσας (Legal Natural Language Processing) και τα εξειδικευμένα νομικά γλωσσικά μοντέλα (legal language models).

Τέλος, συζητούνται οι θεμελιώδεις προκλήσεις που αφορούν την επεξηγησιμότητα (explainability) των μοντέλων, τη διαχείριση της προκατάληψης (bias) και την ποιότητα των νομικών δεδομένων (legal data quality).

3.1 Μεγάλα Γλωσσικά Μοντέλα : Αρχιτεκτονική και Λειτουργία

Τα Μεγάλα Γλωσσικά Μοντέλα αποτελούν σήμερα την αιχμή της έρευνας στην Επεξεργασία Φυσικής Γλώσσας, έχοντας μετασχηματίσει την αυτόματη κατανόηση, παραγωγή και ανάλυση της ανθρώπινης γλώσσας σε πρωτοφανές επίπεδο. Η βασική τους αρχιτεκτονική στηρίζεται στους Μετασχηματιστές (Transformers)[16], οι οποίοι εισήγαγαν έναν ριζικά νέο τρόπο επεξεργασίας ακολουθιακών δεδομένων, ξεπερνώντας τα δομικά και υπολογιστικά όρια των προηγούμενων νευρωνικών προσεγγίσεων, όπως τα Επαναληπτικά Νευρωνικά Δίκτυα (Recurrent Neural Networks – RNNs) και τα Συνελικτικά Δίκτυα (Convolutional Neural Networks) [17].

Η δομή των Μετασχηματιστών επιτρέπει την παράλληλη επεξεργασία εισόδων αυθαίρετου μήκους μέσω του μηχανισμού αυτοπροσοχής (self-attention), καθιστώντας δυνατή την αποδοτική κωδικοποίηση και την ανίχνευση συσχετίσεων ακόμη και μεταξύ απομακρυσμένων λέξεων ή εννοιών μέσα σε ένα κείμενο [17]. Παράλληλα, παραλλαγές της αρχιτεκτονικής, όπως ο Transformer-XL[18], αντιμετωπίζουν αποτελεσματικά το ζήτημα της αποθήκευσης και επεξεργασίας μεγάλων ακολουθιών, το οποίο αποτελεί κρίσιμο παράγοντα για εφαρμογές στο δίκαιο, όπου τα έγγραφα παρουσιάζουν μεγάλο μήκος και περίπλοκη δομή.

Η πρακτική αξία των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στον νομικό τομέα αποδεικνύεται από πρόσφατες μελέτες που επισημαίνουν ότι τα μοντέλα αυτά μπορούν να παράγουν συνεκτικό, ορθό νομικό λόγο και να προσομοιώσουν νομικές γνωμοδοτήσεις σε επίπεδο που καθιστά δύσκολη ακόμη και για νομικούς την ανίχνευση της συνθετικής προέλευσης του κειμένου [19]. Ειδικότερα, η εκπαίδευση σε μεγάλες συλλογές νομικών δεδομένων, όπως δικαστικές αποφάσεις, κανονισμοί και συμβάσεις, επιτρέπει στα LLMs να μάθουν τις ιδιαιτερότητες της νομικής ορολογίας και τη λογική δόμηση των επιχειρημάτων [20].

Στην πράξη, τα LLMs χρησιμοποιούν το πλαίσιο της αυτοεποπτευόμενης μάθησης (self-supervised learning) κατά το στάδιο της προεκπαίδευσης, αξιοποιώντας τεχνικές όπως η κεκαλυμμένη μοντελοποίηση γλώσσας (Masked Language Modeling) ή η αιτιακή μοντελοποίηση (Causal Language Modeling) για να κατανοήσουν τη συντακτική και σημασιολογική πληροφορία που κρύβεται στα μεγάλα νομικά σώματα κειμένων [17, 19]. Στη συνέχεια, η εξειδίκευση των μοντέλων αυτών σε νομικά tasks γίνεται μέσω fine-tuning πάνω σε εξειδικευμένα νομικά δεδομένα, διαδικασία που οδηγεί σε μοντέλα όπως τα Legal-BERT και LegalPro-BERT, τα οποία υπερτερούν σημαντικά έναντι των γενικών LLMs σε tasks όπως η ταξινόμηση ρητρών, η εξαγωγή οντοτήτων και η αναγνώριση συσχετίσεων σε συμβατικά κείμενα [20][21].

Η σημασία της προσαρμογής στη νομική γλώσσα ενισχύεται από τα αποτελέσματα των πρόσφατων διαγωνισμών (όπως το COLIEE[21]), όπου τα εξειδικευμένα μοντέλα transformer επιτυγχάνουν σημαντικά καλύτερη απόδοση τόσο στην ανάκτηση νομικής πληροφορίας όσο και στη συμπερασματολογία (entailment). Τα μοντέλα αυτά αξιοποιούν τεχνικές όπως τα ενσωματώματα σε επίπεδο πρότασης (sentence-level embeddings), τη συσχέτιση μέσω συνημιτόνου (cosine similarity), την απόσταξη γνώσης (knowledge distillation) και προηγμένες μεθόδους ταξινόμησης. Επιπλέον, η αξιοποίηση τεχνικών μεταφοράς μάθησης (transfer learning) και απόσταξης γνώσης (knowledge distillation) έχει επιτρέψει την ανάπτυξη πιο αποδοτικών και ελαφριών μοντέλων, χωρίς να θυσιάζεται η ακρίβεια σε εξειδικευμένες νομικές εργασίες (tasks) [22].

Παράλληλα, μελέτες σε πραγματικές εφαρμογές, όπως η ανάκτηση απαντήσεων σε νομικά ερωτήματα, επιβεβαιώνουν ότι μοντέλα όπως το RoBERTa και το DeBERTa υπερτερούν σημαντικά των παραδοσιακών μεθόδων μηχανικής μάθησης (όπως τα Support Vector Machines), τόσο ως προς την ακρίβεια όσο και ως προς την κατάταξη των αποτελεσμάτων, ενώ παράλληλα αντιμετωπίζουν αποτελεσματικά τα σύνθετα χαρακτηριστικά της νομικής γλώσσας και της ρητορικής [22].

Τέλος, οι προκλήσεις που παραμένουν στο πεδίο των Μεγάλων Γλωσσικών Μοντέλων για νομικές εφαρμογές σχετίζονται κυρίως με τη διαχείριση της έλλειψης μεγάλων, αξιόπιστων επισημασμένων συνόλων δεδομένων, την ύπαρξη εξειδικευμένης ορολογίας, καθώς και τη βέλτιστη αξιοποίηση των τεχνικών fine-tuning και transfer learning για την προσαρμογή των γενικών LLMs στα ιδιαίτερα χαρακτηριστικά του δικαίου[19].

3.1.1 Σύγχρονες Αρχιτεκτονικές Μεγάλων Γλωσσικών Μοντέλων

Η ραγδαία πρόοδος στην ανάπτυξη Μεγάλων Γλωσσικών Μοντέλων (Large Language Models – LLMs) βασίζεται στην τεχνική ωρίμανση και τη διαρκή εξέλιξη των αρχιτεκτονικών

μετασχηματιστών (Transformers) [23]. Μετά την καθιέρωση των Transformers ως το κυρίαρχο υπόδειγμα για την Επεξεργασία Φυσικής Γλώσσας (Natural Language Processing – NLP), η ερευνητική κοινότητα έχει προτείνει και βελτιστοποιήσει μια ευρεία γκάμα παραλλαγών που ανταποκρίνονται στις προκλήσεις της κατανόησης, παραγωγής και ανάλυσης κειμένου σε πραγματικά σενάρια.

Κεντρικό ρόλο στην εξέλιξη αυτή κατέχουν τα μοντέλα μόνο κωδικοποιητή (encoder-only), με χαρακτηριστικότερο παράδειγμα το BERT (Bidirectional Encoder Representations from Transformers) [16]. Το BERT εισάγει τη στρατηγική της αμφίδρομης μάθησης (bidirectional learning) με την κεκαλυμμένη μοντελοποίηση γλώσσας (Masked Language Modeling – MLM) και την πρόβλεψη επόμενης πρότασης (Next Sentence Prediction – NSP), βελτιώνοντας την απόδοση σε εργασίες που απαιτούν κατανόηση σύνθετων γλωσσικών δομών. Μελέτες έδειξαν ότι η αφαίρεση ορισμένων περιορισμών, όπως το NSP, και η εκπαίδευση σε τεράστια datasets οδηγεί σε εντυπωσιακές βελτιώσεις, όπως καταδεικνύει το RoBERTa (Robustly Optimized BERT Approach) [24]. Το RoBERTa πέτυχε κορυφαίες επιδόσεις σε tasks αναγνώρισης οντοτήτων και ταξινόμησης ακόμα και σε νομικά δεδομένα [22], ενώ υποστηρίζει ταχύτερη και σταθερότερη εκπαίδευση μέσω καλύτερης διαχείρισης υπερπαραμέτρων (hyperparameters).

Η ανάγκη για υλοποίηση σε υπολογιστικά περιορισμένα περιβάλλοντα και εφαρμογές πραγματικού χρόνου οδήγησε στην ανάπτυξη ελαφριών παραλλαγών όπως το DistilBERT [25]. Με τεχνική απόσταξης γνώσης (knowledge distillation), το DistilBERT διατηρεί το 97% της απόδοσης του BERT με σχεδόν το μισό μέγεθος και πολύ χαμηλότερο latency, ενώ ο σχεδιασμός του βασίζεται σε τριπλή απώλεια (triple loss) ώστε να μεταφέρει τις γνώσεις και τη συμπεριφορά του teacher μοντέλου στο συμπιεσμένο student μοντέλο.

Στον αντίποδα, τα μοντέλα αποκωδικοποιητή (decoder-only) επικεντρώνονται στη δημιουργία κειμένου, με σημείο αναφοράς το GPT (Generative Pre-trained Transformer) [26]. Η εκπαίδευση μέσω αιτιακής μοντελοποίησης (causal language modeling – CLM) τα καθιστά ικανά να παράγουν φυσικό, συνεκτικό λόγο, με κάθε επόμενη λέξη να εξαρτάται δυναμικά από το συνολικό ιστορικό της ακολουθίας. Η τεράστια ανάπτυξη των παραμέτρων σε μοντέλα όπως το GPT-3 ή τα σύγχρονα PaLM, Llama και Mistral [23], σε συνδυασμό με τεχνικές όπως το Mixture-of-Experts, έθεσε νέα πρότυπα για την πολυλειτουργικότητα και την προσαρμοστικότητα, προσφέροντας δυνατότητες few-shot, zero-shot και in-context learning. Η χρήση prompt engineering, η παραμετρική προσαρμογή (parameter-efficient tuning), αλλά και η υιοθέτηση τεχνικών όπως το Reinforcement Learning from Human Feedback – RLHF [27] ενισχύουν περαιτέρω την απόδοση, την ευθυγράμμιση με ανθρώπινες προσδοκίες και την ικανότητα για ασφαλή, ελεγχόμενη παραγωγή περιεχομένου.

Ειδική θέση στη σύγχρονη βιβλιογραφία καταλαμβάνουν τα μοντέλα κωδικοποιητή-αποκωδικοποιητή (encoder-decoder), όπως το T5 (Text-to-Text Transfer Transformer), που εισάγει μια ενοποιημένη προσέγγιση: όλα τα προβλήματα NLP διατυπώνονται ως μετατροπές κειμένου σε κείμενο. Αυτή η στρατηγική απλοποιεί δραστικά το transfer learning και επιτρέπει το fine-tuning σε πλήθος tasks (μετάφραση, περίληψη, ερώτηση-απάντηση, σύνθεση νομικών εγγράφων) [23].

Οι σύγχρονες εξελίξεις περιλαμβάνουν την ενσωμάτωση εξωτερικών πηγών γνώσης μέσω αρχιτεκτονικών retrieval-augmented generation, επιτρέποντας στα LLMs να συνδυάζουν ε-

πεξεργασία με ενημερωμένες απαντήσεις βασισμένες σε πραγματικό χρόνο [28]. Μελέτες έδειξαν ότι, ειδικά σε domains με υψηλή απαίτηση ακρίβειας και εξειδίκευσης, όπως η νομική πληροφορική, οι RAG προσεγγίσεις βελτιώνουν σημαντικά την αξιοπιστία και την εγκυρότητα των αποτελεσμάτων, αντιμετωπίζοντας αποτελεσματικά το φαινόμενο των «παραισθήσεων» (hallucinations)[27], όπως θα παρουσιαστεί αναλυτικά σε επόμενη ενότητα.

Σημαντικό στοιχείο της αξιολόγησης των αρχιτεκτονικών αποτελεί η συγκριτική τους απόδοση σε τυποποιημένα benchmarks, καθώς και σε εξειδικευμένα σύνολα δεδομένων που αντανakλούν τις απαιτήσεις πραγματικών εφαρμογών [21][20]. Η ανάλυση των αποτελεσμάτων σε διάφορες μετρικές, σε συνδυασμό με τη διερεύνηση της αποδοτικότητας και της προσαρμοστικότητας κάθε προσέγγισης, αποτελεί βασικό κριτήριο για την επιλογή της κατάλληλης αρχιτεκτονικής σε κάθε πεδίο χρήσης [17].

Συνολικά, η διαρκής εξέλιξη των αρχιτεκτονικών LLMs στηρίζεται στη συνδυαστική αξιοποίηση προόδων στον σχεδιασμό νευρωνικών δικτύων, στην αύξηση της υπολογιστικής ισχύος και στην καινοτομία στα μαθησιακά σχήματα. Οι προοπτικές για τα επόμενα χρόνια περιλαμβάνουν ακόμα πιο αποδοτικά, εξειδικευμένα και ερμηνεύσιμα μοντέλα, καθώς και βαθύτερη ενσωμάτωση σε εφαρμογές υψηλής κρισιμότητας όπως το δίκαιο, η υγεία και η δημόσια διοίκηση [23][28].

3.1.2 Μεθοδολογίες Εκπαίδευσης και Εξειδίκευσης

Η εκπαίδευση των Μεγάλων Γλωσσικών Μοντέλων (LLMs) αποτελείται τυπικά από δύο διακριτά στάδια: την προεκπαίδευση (pre-training) και την εξειδίκευση ή περαιτέρω εκπαίδευση (fine-tuning). Στο πρώτο στάδιο, το μοντέλο εκπαιδεύεται σε τεράστιες συλλογές μη επισημειωμένων δεδομένων, όπως βιβλία, άρθρα, ιστότοπους και φόρουμ. Για το BERT[29], χρησιμοποιείται η τεχνική της κεκαλυμμένης μοντελοποίησης γλώσσας (MLM), όπου ορισμένες λέξεις αντικαθίστανται από ειδικά tokens και το μοντέλο καλείται να προβλέψει το αρχικό περιεχόμενο, αξιοποιώντας τα συμφραζόμενα τόσο πριν όσο και μετά το κενό. Επίσης, εφαρμόζεται η διαδικασία Next Sentence Prediction, κατά την οποία το μοντέλο μαθαίνει να διακρίνει αν δύο προτάσεις ακολουθούν λογικά η μία την άλλη.

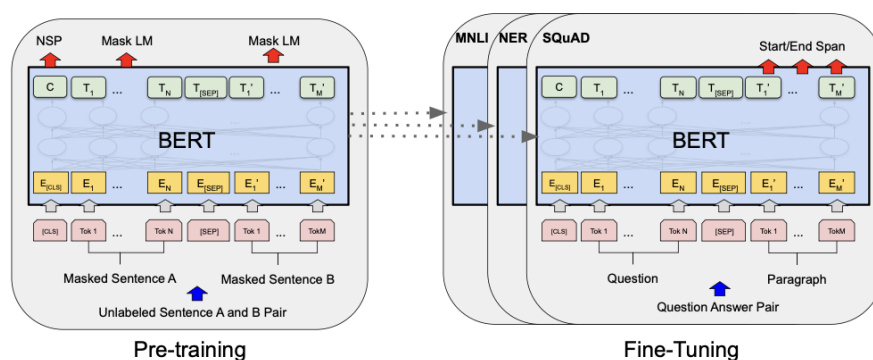
Στο στάδιο του fine-tuning, το ήδη προεκπαιδευμένο μοντέλο προσαρμόζεται σε συγκεκριμένες εργασίες μέσω εκπαίδευσης σε μικρότερα, εξειδικευμένα σύνολα δεδομένων. Για παράδειγμα, ένα μοντέλο μπορεί να εξειδικευθεί σε εργασίες ταξινόμησης νομικών εγγράφων, αναγνώρισης οντοτήτων ή απάντησης σε ερωτήματα νομικού περιεχομένου. Κατά το fine-tuning, όλες οι παράμετροι του μοντέλου αναπροσαρμόζονται με βάση το νέο task, γεγονός που εξασφαλίζει υψηλή απόδοση ακόμη και σε πεδία με έντονη εξειδίκευση, όπως το δίκαιο.

Η συνολική διαδικασία προεκπαίδευσης και εξειδίκευσης για το BERT αποτυπώνεται χαρακτηριστικά στο παρακάτω διάγραμμα, το οποίο παρουσιάζει την αρχιτεκτονική του μοντέλου και τα βασικά στάδια μάθησης, όπως δημοσιεύτηκαν στην εργασία των [16]

Επιπλέον, σύγχρονες τεχνικές, όπως η προσαρμογή προτροπών (prompt tuning) — δηλαδή η προσεκτική διαμόρφωση των εισόδων προς το μοντέλο — και η ενισχυτική μάθηση με ανθρώπινη ανατροφοδότηση (Reinforcement Learning from Human Feedback – RLHF), έχουν επιτρέψει τη δραστική βελτίωση των παραγόμενων αποτελεσμάτων. Τα μοντέλα αυτά,

μέσω fine-tuning σε νομικά corpus, προσφέρουν πλέον λύσεις για απαιτητικές εφαρμογές, όπως αυτόματη εξαγωγή ρητρών από συμβόλαια, νομική κατηγοριοποίηση και ερμηνεία σύνθετων νομικών εγγράφων.

Η παραπάνω μεθοδολογία εκπαίδευσης και εξειδίκευσης αποτελεί το θεμέλιο για τη μετάβαση από γενικούς γλωσσικούς επεξεργαστές σε εξειδικευμένα εργαλεία για τον νομικό τομέα, διασφαλίζοντας υψηλή ακρίβεια και αξιοπιστία στα παραγόμενα αποτελέσματα.



Σχήμα 3.1: Σχηματική απεικόνιση των σταδίων προεκπαίδευσης και εξειδίκευσης (fine-tuning) του BERT.

3.2 Τεχνητή Νοημοσύνη στη Νομική Πληροφορική

Η Νομική Πληροφορική αποτελεί το επιστημονικό πεδίο που εστιάζει στην ανάλυση, οργάνωση και επεξεργασία νομικών πληροφοριών μέσω προηγμένων τεχνολογικών μεθόδων. Με τη ραγδαία εξέλιξη της Τεχνητής Νοημοσύνης ο τρόπος διαχείρισης, ανάλυσης και αξιοποίησης νομικών δεδομένων έχει μετασχηματιστεί, αναβαθμίζοντας ριζικά τόσο την ερευνητική όσο και την επαγγελματική πρακτική [30].

Η αξιοποίηση της τεχνητής νοημοσύνης στη νομική πρακτική εκτείνεται σε ένα ευρύ φάσμα εφαρμογών: από τη νομική έρευνα και ανάκτηση πληροφορίας (Legal Research and Retrieval), τη σύνοψη και ανάλυση δικαστικών αποφάσεων, την αυτόματη σύνταξη νομικών εγγράφων, μέχρι τη νομική πρόβλεψη (Legal Judgment Prediction) και τη συμμόρφωση με κανονιστικά πρότυπα [31]. Τα σύγχρονα εργαλεία AI επιτρέπουν την ταχύτερη και πιο στοχευμένη εύρεση νομολογίας και νομοθεσίας, ακόμη και με ερωτήματα διατυπωμένα σε φυσική γλώσσα, παρέχοντας ουσιαστική υποστήριξη στη λήψη αποφάσεων τόσο για επαγγελματίες όσο και για φορείς της δικαιοσύνης [32].

Ιδιαίτερη πρόοδος έχει σημειωθεί στην ανάπτυξη εξειδικευμένων μοντέλων για την ανάλυση σύνθετων νομικών εγγράφων και τη σύνοψη πολυσελίδων δικαστικών αποφάσεων, συμβάλλοντας στην αποτελεσματική αποτύπωση της ουσίας του νομικού συλλογισμού [33]. Παράλληλα, εφαρμογές αυτόματης σύνταξης συμβολαίων (Legal Document Drafting) και συστήματα πρόβλεψης εκβάσεων υποθέσεων (Legal Judgment Prediction) αξιοποιούν τεχνικές μηχανικής μάθησης (Machine Learning) και επεξεργασίας φυσικής γλώσσας (NLP) για να προσφέρουν γρήγορες και τεκμηριωμένες λύσεις σε νομικά ερωτήματα [31][32].

Ωστόσο, η ενσωμάτωση της τεχνητής νοημοσύνης στη νομική πληροφορική παρουσι-

άζει αρκετές δυσκολίες. Σημαντικές προκλήσεις αφορούν την εναρμόνιση των έξυπνων συστημάτων με το υπάρχον ρυθμιστικό και δεοντολογικό πλαίσιο κάθε χώρας, τη δυσκολία μεταφοράς τεχνογνωσίας μεταξύ διαφορετικών δικαιοδοσιών και τη διατήρηση της λογοδοσίας (accountability) των τελικών αποφάσεων [34]. Η παραγωγή αξιόπιστων μοντέλων συχνά προσκρούει στην έλλειψη μεγάλων, ισορροπημένων και σωστά ανωνυμοποιημένων συνόλων δεδομένων, ενώ η σωστή ερμηνεία των αποτελεσμάτων απαιτεί την ενεργή συμμετοχή νομικών εμπειρογνομόνων στη διαδικασία ελέγχου και αξιολόγησης [31]. Επιπλέον, η συνεχής εξέλιξη της νομολογίας και των νομικών εννοιών καθιστά απαραίτητη τη συστηματική επικαιροποίηση τόσο των μοντέλων όσο και των ίδιων των μεθόδων αξιολόγησής τους, ώστε να διασφαλίζεται η συμμόρφωση με τις αρχές της δικαιοσύνης και της ασφάλειας δικαίου.

Τέλος, η Τεχνητή Νοημοσύνη δεν αντικαθιστά απλώς τον νομικό επαγγελματία, αλλά αναμένεται να αναδιαμορφώσει το πεδίο τα επόμενα χρόνια, καθώς προωθεί τη σταδιακή αλλαγή ρόλων και την ανάδειξη διεπιστημονικών ομάδων), όπου η τεχνική γνώση και η νομική εμπειρία συνεργάζονται στον σχεδιασμό και την αξιολόγηση έξυπνων συστημάτων [30]. Ο εν εξέλιξη αυτός μετασχηματισμός δημιουργεί νέες προοπτικές και θέτει στο επίκεντρο τον διάλογο για τα όρια, τις ευθύνες και τις αρμοδιότητες του νομικού επαγγέλματος στη νέα ψηφιακή εποχή, ανοίγοντας σημαντικά ερευνητικά και θεσμικά ερωτήματα για το μέλλον.

3.2.1 Εξέλιξη της Νομικής Επεξεργασίας Φυσικής Γλώσσας

Η Νομική Επεξεργασία Φυσικής Γλώσσας αποτελεί θεμέλιο της σύγχρονης νομικής πληροφορικής, με στόχο την αυτόματη ανάλυση, οργάνωση και παραγωγή νομικών κειμένων μέσω υπολογιστικών μεθόδων. Η εξέλιξη του πεδίου είναι άρρηκτα συνδεδεμένη με τη γενικότερη πρόοδο της επεξεργασίας φυσικής γλώσσας, αλλά χαρακτηρίζεται από ιδιαιτερότητες που απορρέουν από τη δομή, το ύφος και την πολυπλοκότητα των νομικών εγγράφων.

Στα πρώτα βήματα, η έρευνα επικεντρώθηκε σε συστήματα βασισμένα σε κανόνες (rule-based systems), όπου ομάδες νομικών και μηχανικών κατέγραφαν λεξικά, μοτίβα και λογικούς κανόνες για την αναγνώριση νομικών όρων, τύπων εγγράφων και βασικών οντοτήτων [33]. Η προσέγγιση αυτή ήταν αποτελεσματική σε περιορισμένα, αυστηρά ορισμένα tasks, όμως αποδείχθηκε ανεπαρκής μπροστά στην εκθετική αύξηση του όγκου και της ποικιλίας των νομικών δεδομένων και τη διαρκή εξέλιξη της γλώσσας του δικαίου.

Η επόμενη φάση χαρακτηρίστηκε από την είσοδο των μεθόδων μηχανικής μάθησης (machine learning) και τη σταδιακή αξιοποίηση επισημειωμένων συνόλων δεδομένων (annotated corpora), όπως τα EURLEX, LEDGAR και COLIEE [31]. Αλγόριθμοι όπως οι Υποστηρικτές Διανυσματικών Μηχανών (Support Vector Machines), τα δέντρα αποφάσεων (decision trees) και τα πρώτα νευρωνικά δίκτυα (neural networks) χρησιμοποιήθηκαν εκτεταμένα για την ταξινόμηση νομικών εγγράφων, την αναγνώριση νομικών οντοτήτων (Legal Named Entity Recognition), αλλά και την πρόβλεψη εκβάσεων (outcome prediction) σε συγκεκριμένες υποθέσεις. Η πρόοδος αυτή υποστηρίχθηκε από τη δημοσίευση και ανοιχτή διάθεση μεγάλων συνόλων δεδομένων με δομημένες πληροφορίες για νομολογία, αποφάσεις, κατηγορίες και θεματικές ενότητες, επιτρέποντας για πρώτη φορά τη συστηματική αξιολόγηση και σύγκριση αλγορίθμων.

Το πραγματικό άλμα στην επεξεργασία νομικού κειμένου σημειώθηκε με την έλευση

των Μετασχηματιστών (Transformers) και των Μεγάλων Γλωσσικών Μοντέλων. Οι νεότερες αρχιτεκτονικές, όπως το BERT και οι εξειδικευμένες εκδοχές του για το δίκαιο (π.χ. LegalBERT, Lawformer), εκπαιδεύτηκαν πάνω σε ογκώδη νομικά σώματα κειμένου και επέτρεψαν σαφώς πιο ακριβή κατανόηση, ανάλυση και διαχείριση της νομικής ορολογίας, της δομής και της σημασίας των εγγράφων [35]. Οι τεχνικές μεταφοράς μάθησης (transfer learning) και εξειδικευμένης προεκπαίδευσης (domain-adaptive pretraining) επιτρέπουν πλέον στα γλωσσικά μοντέλα να προσαρμόζονται εύκολα σε διαφορετικά νομικά συστήματα (legal systems) και κατηγορίες εγγράφων, αυξάνοντας μετρήσιμα την απόδοση σε απαιτητικά tasks όπως η αυτόματη σύνοψη (summarization), η εξόρυξη επιχειρημάτων (argument mining), η ανάλυση συμβολαίων (contract analytics), αλλά και η πρόβλεψη δικαστικών αποφάσεων (legal judgment prediction).

Αξίζει να σημειωθεί ότι οι σύγχρονες εφαρμογές νομικής επεξεργασίας φυσικής γλώσσας στηρίζονται συχνά σε πολυεπίπεδα μοντέλα με μηχανισμούς προσοχής (attention mechanisms), πολυεργασιακή μάθηση (multi-task learning) και ιεραρχική δόμηση (hierarchical modeling), αξιοποιώντας τόσο τη δομή των εγγράφων όσο και το σημασιολογικό τους βάθος [31]. Παράλληλα, η ανάπτυξη εξειδικευμένων μετρικών αξιολόγησης (evaluation metrics), που λαμβάνουν υπόψη τη συνοχή, την ερμηνευσιμότητα και τη δικανική αξιοπιστία των αποτελεσμάτων, ενισχύει τη διαφάνεια και την εμπιστοσύνη στα συστήματα αυτά.

Τέλος, οι πλέον εξελιγμένες πλατφόρμες νομικής τεχνολογίας παρέχουν ολοκληρωμένες λύσεις που αξιοποιούν ενσωματώματα λέξεων (embeddings), πιπελινες βασισμένα σε ανάκτηση πληροφορίας (retrieval-based pipelines) και μηχανισμούς εξατομίκευσης (personalization mechanisms), καθιστώντας δυνατή την ανάκτηση νομολογίας (case law retrieval), την αυτοματοποιημένη δημιουργία συνοπτικών αναφορών (automated summarization) και την υποστήριξη στη λήψη νομικών αποφάσεων σε πραγματικό χρόνο (real-time legal decision support) [32]. Όπως θα αναλυθεί στο επόμενο κεφάλαιο, η ραγδαία αυτή εξέλιξη διαμορφώνει νέα πρότυπα στην έρευνα, τη δικαστική πρακτική και τις νομικές υπηρεσίες, μετασχηματίζοντας το οικοσύστημα της δικαιοσύνης στη σύγχρονη ψηφιακή εποχή.

3.2.2 Νομικά Γλωσσικά Μοντέλα Εξειδικευμένων Τομέων

Τα εξειδικευμένα νομικά γλωσσικά μοντέλα έχουν αναπτυχθεί για να αντιμετωπίζουν τις ιδιαίτερες απαιτήσεις της νομικής γλώσσας και πληροφορίας. Εκπαιδεύονται σε μεγάλα, επισημειωμένα νομικά σύνολα δεδομένων, όπως νομοθεσία, δικαστικές αποφάσεις και συμβάσεις, ώστε να βελτιστοποιούνται σε εξειδικευμένες εργασίες (legal tasks) όπως η ταξινόμηση εγγράφων (document classification), η αναγνώριση νομικών οντοτήτων (Legal Named Entity Recognition – Legal NER), η εξαγωγή ρητρών (clause extraction) και η ερώτηση-απάντηση νομικού περιεχομένου (Legal Question Answering) [35][31].

Ενδεικτικά παραδείγματα αποτελούν τα μοντέλα LegalBERT [35], CaseLawBERT[36] και Lawformer [37]. Το LegalBERT βασίζεται στην αρχιτεκτονική του BERT και είναι προεκπαιδευμένο σε ποικιλία νομικών κειμένων, αποδίδοντας με ακρίβεια σε νομικές ταξινομήσεις και αναγνώριση οντοτήτων. Το Lawformer ενσωματώνει τεχνικές επεξεργασίας μακροσκελών κειμένων (Longformer Attention), επιτρέποντας την ανάλυση πλήρων δικαστικών αποφάσεων χωρίς τεχνητούς περιορισμούς στο μήκος εισόδου. Ταυτόχρονα, εφαρμόζονται προηγμένες

τεχνικές ενσωμάτωσης (entity embeddings), ιεραρχικές αρχιτεκτονικές (hierarchical architectures) και εξειδικευμένο λεξιλόγιο (domain-adaptive tokenization), ώστε το μοντέλο να “κατανοεί” τη δομή και τα συμφραζόμενα της νομικής γλώσσας [31].

Στη σύγχρονη έρευνα, τα νομικά LLMs υιοθετούν τεχνικές πολυ-εργασιακής μάθησης (multi-task learning), προεκπαίδευσης προσαρμοσμένης σε πεδίο (domain-adaptive pretraining), καθώς και μεθόδους ενισχυμένης ανάκτησης (retrieval-augmented generation – RAG) για την αύξηση της ακρίβειας, τη μείωση ψευδαισθήσεων (hallucinations) και τη βελτίωση της γενίκευσης (domain generalization) σε διαφορετικές δικαιοδοσίες (cross-jurisdictional analysis) [32]. Η αξιολόγηση αυτών των μοντέλων βασίζεται σε τυποποιημένα νομικά σύνολα δεδομένων (legal benchmarks), όπως τα COLIEE, EURLEX, LEDGAR, και σε μετρικές που αντανakλούν τις απαιτήσεις και την πολυπλοκότητα του δικαίου [31].

Ειδικότερα, οι πιο πρόσφατες εξελίξεις στα νομικά LLMs αφορούν την ενσωμάτωση μηχανισμών ελέγχου αξιοπιστίας (confidence estimation), την υλοποίηση δομών ανιχνεύσιμης λογικής (traceable reasoning), καθώς και τη διασταύρωση αποτελεσμάτων με εξωτερικές νομικές βάσεις δεδομένων (external legal knowledge sources). Παράλληλα, εφαρμόζονται μέθοδοι αυτόματης ανίχνευσης και διόρθωσης μεροληψίας (bias detection and correction), που βασίζονται σε μετα-μάθηση (meta-learning) και επαναληπτική επικύρωση (iterative validation) των εξαγόμενων συμπερασμάτων [33][31].

3.2.3 Αναπαράσταση Νομικής Γνώσης

Η αναπαράσταση της νομικής γνώσης (legal knowledge representation) αποτελεί κρίσιμο πεδίο για την ανάπτυξη συστημάτων τεχνητής νοημοσύνης που εφαρμόζονται στο δίκαιο. Αφορά τον τρόπο με τον οποίο η νομική πληροφορία μοντελοποιείται και οργανώνεται, ώστε να μπορεί να κατανοηθεί, να αναλυθεί και να αξιοποιηθεί αποτελεσματικά από αλγοριθμικά συστήματα.

Οι βασικές προσεγγίσεις στην αναπαράσταση της νομικής γνώσης περιλαμβάνουν τη χρήση οντολογιών (ontologies), οι οποίες δημιουργούν ιεραρχικές δομές εννοιών και σχέσεων στον νομικό χώρο, επιτρέποντας στα συστήματα να κατανοούν τις σχέσεις μεταξύ διαφορετικών νομικών όρων, όπως νόμοι, διατάξεις, προηγούμενα δικαστικών αποφάσεων και συμβάσεις [38]. Επίσης, αξιοποιούνται γραφήματα γνώσης (knowledge graphs), τα οποία οργανώνουν τις νομικές πληροφορίες ως δίκτυα κόμβων και σχέσεων, με κάθε κόμβο να αντιπροσωπεύει μία νομική οντότητα και κάθε σύνδεσμο να αποτυπώνει τη σχέση μεταξύ τους. Αυτή η μορφή αναπαράστασης διευκολύνει την αναζήτηση, την ανακάλυψη και την κατανόηση της νομικής πληροφορίας. Μια ακόμη σημαντική τεχνική είναι η επισήμανση σημασιολογικών ρόλων (Semantic Role Labeling – SRL)[38], η οποία βοηθά στη δόμηση των νομικών κειμένων με βάση τους ρόλους που διαδραματίζουν οι διάφορες έννοιες σε μια πρόταση ή σε ένα νομικό επιχείρημα. Η επισήμανση σημασιολογικών ρόλων είναι τεχνική ανάλυσης που προσδιορίζει τα νοηματικά συστατικά μιας πρότασης, αποδίδοντας ρόλους όπως υποκείμενο, αντικείμενο και ενέργεια. Για παράδειγμα, στη φράση ‘Ο δικαστής καταδίκασε τον κατηγορούμενο’, η ΣΡΛ αναγνωρίζει τον «δικαστή» ως agent, την καταδίκη ως action και τον «κατηγορούμενο» ως patient ή recipient.

Τέλος, τα λογικά πλαίσια και ο συμπερασμός (logical frameworks and reasoning) χρη-

σιμοποιούνται για τη διατύπωση κανόνων και την αυτόματη εκτέλεση ελέγχων ορθότητας σε νομικά επιχειρήματα ή διαδικασίες (legal reasoning).

Η σωστή αναπαράσταση της νομικής γνώσης είναι απαραίτητη για την ανάπτυξη εξειδικευμένων νομικών LLMs που μπορούν να χειριστούν σύνθετες νομικές εργασίες, τη διασφάλιση συνέπειας και ακρίβειας στα αποτελέσματα των νομικών εφαρμογών τεχνητής νοημοσύνης (AI legal applications), καθώς και για τη διευκόλυνση της επεξηγησιμότητας και ερμηνευσιμότητας (explainability and interpretability) των αποφάσεων που λαμβάνονται από τεχνητά συστήματα. Χωρίς σαφή και ορθή αναπαράσταση, ακόμη και τα πιο προηγμένα γλωσσικά μοντέλα κινδυνεύουν να παρανοήσουν τη σημασία και τη λειτουργία κρίσιμων νομικών όρων και θεσμικών δομών.

3.3 Θεμελιώδεις Προκλήσεις

Παρά την εντυπωσιακή τεχνολογική πρόοδο που αναλύθηκε παραπάνω, η ενσωμάτωση των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στη νομική επιστήμη, αν και προσφέρει σημαντικές δυνατότητες, συνοδεύεται από σοβαρές προκλήσεις που επηρεάζουν άμεσα τη βιωσιμότητα και την αξιοπιστία των εφαρμογών Τεχνητής Νοημοσύνης στο δίκαιο. Οι κύριες προκλήσεις αφορούν την επεξηγησιμότητα και ερμηνευσιμότητα των μοντέλων (explainability and interpretability), την προκατάληψη και τη δικαιοσύνη στην απόδοση των αποτελεσμάτων (bias and fairness), καθώς και την ποιότητα και την προσβασιμότητα των νομικών δεδομένων (legal data quality and accessibility). Η κατανόηση και η αντιμετώπιση αυτών των θεμελιωδών ζητημάτων είναι αναγκαία για την υπεύθυνη και ηθική χρήση των LLMs στο νομικό πλαίσιο. Στις επόμενες ενότητες αναλύονται αναλυτικά οι τρεις αυτοί βασικοί άξονες.

3.3.1 Επεξηγησιμότητα και Ερμηνευσιμότητα

Η επεξηγησιμότητα (explainability) και η ερμηνευσιμότητα (interpretability) αποτελούν αδιαπραγμάτευτες απαιτήσεις για την αξιόπιστη και θεσμικά αποδεκτή χρήση των LLMs σε νομικά πλαίσια. Η εφαρμογή της τεχνητής νοημοσύνης στη νομική πρακτική δημιουργεί ένα ιδιαίτερα απαιτητικό περιβάλλον, όπου κάθε απόφαση, κάθε πρόβλεψη και κάθε αυτόματη σύσταση οφείλει να συνοδεύεται από διαφανή και πλήρη αιτιολόγηση. Η ερμηνευσιμότητα (interpretability) σχετίζεται με τη δυνατότητα των συστημάτων να παρέχουν μια διαυγή και κατανοητή περιγραφή του εσωτερικού τους μηχανισμού λήψης αποφάσεων, ενώ η επεξηγησιμότητα (explainability) επικεντρώνεται στη δυνατότητα παροχής πειστικής και τεκμηριωμένης αιτίας για κάθε συγκεκριμένη έξοδο ή πρόβλεψη, ακόμα και όταν το μοντέλο παραμένει εσωτερικά πολύπλοκο ή αδιαφανές [33][39].

Η πρόκληση της επεξηγησιμότητας γίνεται εντονότερη λόγω της αρχιτεκτονικής φύσης των σύγχρονων LLMs, οι οποίοι βασίζονται σε βαθιά νευρωνικά δίκτυα (deep neural networks) με εκατομμύρια παραμέτρους και μη-γραμμικές αλληλεπιδράσεις. Στην πράξη, το μαύρο κουτί (black box) των μοντέλων αυτών δυσχεραίνει την εξαγωγή λογικής αλυσίδας που συνδέει τα δεδομένα εισόδου με το τελικό αποτέλεσμα, καθιστώντας την επεξήγηση του μοντέλου απολύτως απαραίτητη τόσο για την οικοδόμηση εμπιστοσύνης όσο και για τη διασφάλιση της νομιμότητας της διαδικασίας λήψης αποφάσεων. Ειδικά στο δίκαιο, η αδυνα-

μία αιτιολόγησης μιας απόφασης, είτε πρόκειται για δικαστική πρόβλεψη είτε για αυτόματη σύνοψη ή σύσταση συμβολαίου, εγείρει ηθικά και θεσμικά ζητήματα, καθώς παραβιάζει την αρχή της λογοδοσίας, δυσχεραίνει τον δικαστικό ή διοικητικό έλεγχο και ενδέχεται να οδηγήσει σε μη ανιχνεύσιμες προκαταλήψεις (bias) ή συστηματικά σφάλματα [33][34].

Τα σύγχρονα νομικά LLMs συχνά παράγουν αποτελέσματα χωρίς να συνοδεύονται από επαρκή αιτιολόγηση. Αυτό έχει ως συνέπεια οι νομικοί επαγγελματίες να αδυνατούν να ελέγξουν τη διαδρομή του συλλογισμού που οδήγησε σε μια σύσταση, απόφαση ή διάγνωση. Το πρόβλημα γίνεται οξύτερο σε ζητήματα που αφορούν τα δικαιώματα των διαδίκων, την αρχή της δίκαιης δίκης (fair trial) και την προστασία των προσωπικών δεδομένων, καθώς η αδιαφάνεια (opacity) στα συστήματα AI μπορεί να αποτελέσει εμπόδιο στη δικαστική επανεξέταση και να μειώσει τη θεσμική αποδοχή και την κοινωνική εμπιστοσύνη στη χρήση της τεχνητής νοημοσύνης στη δικαιοσύνη[30].

Για να καλυφθεί το χάσμα της επεξηγησιμότητας, η έρευνα εστιάζει σε πολυεπίπεδες τεχνικές που συνδυάζουν την ακριβή περιγραφή των εσωτερικών μηχανισμών με λειτουργικές εξηγήσεις για τον τελικό χρήστη. Μεθοδολογίες όπως οι ανεξάρτητες από το μοντέλο (model-agnostic) τεχνικές, π.χ. το LIME (Local Interpretable Model-agnostic Explanations)[39] και το SHAP (SHapley Additive exPlanations)[40], επιτρέπουν την τοπική προσομοίωση της συμπεριφοράς του μοντέλου, παρέχοντας κατανοητές εξηγήσεις για συγκεκριμένες προβλέψεις ή συστάσεις, χωρίς να απαιτείται η πλήρης αποκωδικοποίηση του ίδιου του LLM [39]. Ταυτόχρονα, αυτοεπεξηγηματικά μοντέλα (self-explaining models) επιχειρούν να παραγάγουν αναλυτικές εξηγήσεις ενσωματωμένες στην τελική έξοδο του συστήματος, ενώ η τεχνική της οπτικοποίησης προσοχής (attention visualization) επιτρέπει την εντοπισμό των λέξεων ή των τμημάτων του κειμένου που επηρέασαν καθοριστικά την τελική πρόβλεψη. Η προσέγγιση του prompt engineering και ιδιαίτερα τα επεξηγηματικά prompts αξιοποιούνται για να υποχρεώσουν τα LLMs να παράγουν απαντήσεις με δομημένη αιτιολόγηση, μειώνοντας την αδιαφάνεια και αυξάνοντας την τεκμηρίωση της κάθε πρόβλεψης. Επιπλέον, δοκιμάζονται ερμηνεύσιμες αρχιτεκτονικές (interpretable architectures), όπου η διασύνδεση συμβολικής λογικής (symbolic logic) με νευρωνικά δίκτυα διευκολύνει τη μετατροπή της αλληλουχίας των συλλογισμών του μοντέλου σε ευθέως ελέγξιμα και ερμηνεύσιμα μοτίβα [33].

Παρά τις σημαντικές προόδους, η επίτευξη πλήρους επεξηγησιμότητας και ερμηνευσιμότητας για τα σύγχρονα νομικά LLMs παραμένει ανοιχτή ερευνητική και τεχνολογική πρόκληση. Οι αυξανόμενες απαιτήσεις για διαφάνεια, δίκαιη χρήση και νομοθετική συμμόρφωση επιβάλλουν τη συνεχή ανάπτυξη μεθοδολογιών που γεφυρώνουν το χάσμα μεταξύ τεχνικής πολυπλοκότητας και πρακτικής ανάγκης για εξήγηση. Σε αυτήν τη συνάρτηση, η διεπιστημονική συνεργασία μεταξύ τεχνικών επιστημόνων, νομικών και ρυθμιστικών φορέων θεωρείται κρίσιμη για τον σχεδιασμό και την αποδοχή συστημάτων τεχνητής νοημοσύνης που λειτουργούν με διαφάνεια, λογοδοσία και θεσμική νομιμοποίηση στο νομικό οικοσύστημα [33] [34][30].

3.3.2 Προκατάληψη και Δικαιοσύνη στην Τεχνητή Νοημοσύνη του Δικαίου

Η προκατάληψη (bias) και η δικαιοσύνη (fairness) συνιστούν θεμελιώδη ζητήματα στη χρήση Μεγάλων Γλωσσικών Μοντέλων (LLMs) σε εφαρμογές δικαίου. Τα νομικά συστήματα

βασίζονται στις αρχές της ισότητας, της αντικειμενικότητας και της διαφάνειας, και, επομένως, οποιαδήποτε τεχνολογική λύση που εισάγεται στο νομικό πλαίσιο οφείλει να διασφαλίζει ότι δεν εισάγει ή αναπαράγει άδικες προκαταλήψεις. Οι πηγές προκατάληψης στα LLMs είναι ποικίλες. Καταρχάς, η προκατάληψη στα δεδομένα εκπαίδευσης (training data bias) είναι εκτεταμένη, καθώς τα μοντέλα αυτά εκπαιδεύονται σε τεράστιες συλλογές δεδομένων που συχνά περιέχουν λανθάνουσες κοινωνικές, πολιτισμικές ή νομικές προκαταλήψεις [1]. Επιπλέον, η ανισορροπία στις αναπαραστάσεις (representation imbalance) μπορεί να οδηγήσει σε συστηματικές ανισότητες, καθώς κάποιες κατηγορίες πληθυσμών ή νομικών εννοιών υποεκπροσωπούνται στα δεδομένα. Τέλος, η ασυνείδητη αναπαραγωγή ιστορικών αδικιών (historical injustice perpetuation) ενδέχεται να ενισχύσει υπάρχουσες προκαταλήψεις, π.χ. μέσω της άνισης πρόβλεψης εκβάσεων ανάλογα με την κοινωνική ή οικονομική τάξη δίκων.

Οι επιπτώσεις της προκατάληψης στη νομική πρακτική είναι ιδιαίτερα σημαντικές. Η διατάραξη της ισότητας ενώπιον του νόμου (disruption of legal equality) αποτελεί άμεση απειλή, καθώς η αθέλητη μεροληψία στα αποτελέσματα μπορεί να υπονομεύσει το θεμελιώδες δικαίωμα των πολιτών σε δίκαιη και ίση μεταχείριση. Επιπλέον, η απώλεια εμπιστοσύνης στα συστήματα δικαιοσύνης (trust erosion in legal systems) μπορεί να αποτελέσει σοβαρό εμπόδιο για την κοινωνική αποδοχή της τεχνολογίας στον χώρο της δικαιοσύνης, αν τα αυτόματα συστήματα νομικής υποστήριξης εμφανίζουν συστηματικά προκατειλημμένες προβλέψεις. Τέλος, υπάρχουν και νομικές και ρυθμιστικές συνέπειες (legal and regulatory implications), καθώς η χρήση συστημάτων με προκαταλήψεις μπορεί να οδηγήσει σε νομικές προσφυγές ή κυρώσεις για τους φορείς που τα χρησιμοποιούν.

Για την αντιμετώπιση της προκατάληψης έχουν προταθεί διάφορες προσεγγίσεις, οι οποίες περιλαμβάνουν την εξυγίανση δεδομένων (data sanitization), μέσω τεχνικών καθαρισμού και φιλτραρίσματος των δεδομένων πριν από την εκπαίδευση, τη χρήση μετρικών δικαιοσύνης (fairness metrics) όπως το Demographic Parity και το Equalized Odds για την αξιολόγηση της δικαιοσύνης στα αποτελέσματα, καθώς και αλγοριθμικές τεχνικές άμβλυνσης προκατάληψης (algorithmic debiasing) όπως η αφαίρεση της διαφοροποίησης (disparate impact removal) [41]. Επιπλέον, η διαφάνεια και η λογοδοσία (transparency and accountability) ενισχύονται μέσω της ανάπτυξης διαφανών μοντέλων και της δημιουργίας μηχανισμών ελέγχου αποφάσεων. Η διαχείριση των προκαταλήψεων αποτελεί κρίσιμο ζήτημα για την υπεύθυνη χρήση της τεχνητής νοημοσύνης στον χώρο του δικαίου και αναγνωρίζεται πλέον ως προτεραιότητα τόσο από την ερευνητική κοινότητα όσο και από θεσμικούς φορείς διεθνώς. Οι νομικές εφαρμογές της τεχνητής νοημοσύνης, ιδίως όσες βασίζονται σε LLMs, καλούνται να λειτουργούν σε περιβάλλοντα όπου η αμεροληψία, η ισότητα και η διαφάνεια είναι θεμελιώδεις αρχές. Ωστόσο, η ενσωμάτωση τέτοιων τεχνολογιών ενέχει σημαντικούς κινδύνους αναπαραγωγής ή και ενίσχυσης υπάρχουσων προκαταλήψεων, είτε αυτές προέρχονται από τα αρχικά δεδομένα, είτε από τους ίδιους τους αλγορίθμους, είτε από την αλληλεπίδραση με τους τελικούς χρήστες.

Στο Σχήμα που ακολουθεί παρουσιάζεται σχηματικά ο κύκλος ενίσχυσης της προκατάληψης στα συστήματα τεχνητής νοημοσύνης.

Το παραπάνω διάγραμμα (Σχήμα 3.2) καταδεικνύει πώς η προκατάληψη μπορεί να εισέλθει και να ανακυκλωθεί σε κάθε στάδιο του οικοσυστήματος: στα δεδομένα, μέσω συ-

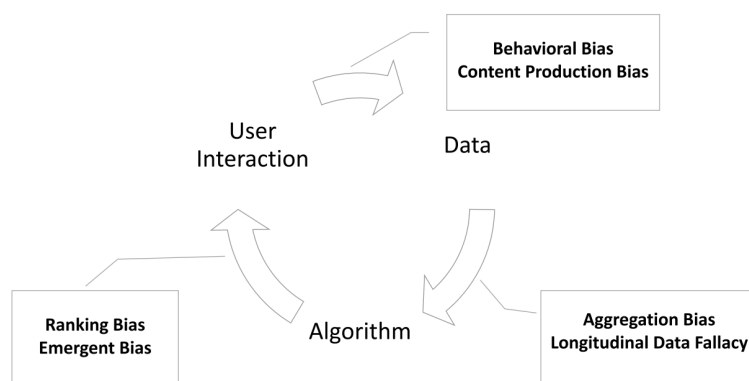


Fig. 1. Examples of bias definitions placed in the data, algorithm, and user interaction feedback loop.

Σχήμα 3.2: Ενδεικτικές κατηγορίες προκατάληψης (bias) που εμφανίζονται στα δεδομένα (data), στους αλγορίθμους (algorithm) και στην αλληλεπίδραση με τον χρήστη (user interaction), αποτυπώνοντας τον δυναμικό βρόχο αναπαραγωγής και ενίσχυσης προκαταλήψεων σε συστήματα τεχνητής νοημοσύνης [1].

μπεριφορικών ή θεματικών στρεβλώσεων (behavioral bias, content production bias), στους αλγορίθμους, με φαινόμενα όπως η προκατάληψη στην κατάταξη (ranking bias) ή η εμφάνιση νέων, απρόβλεπτων στρεβλώσεων (emergent bias), και στην αλληλεπίδραση με τον χρήστη, όπου οι επιλογές και οι αντιδράσεις των χρηστών επηρεάζουν περαιτέρω τη διαδικασία εκμάθησης και αξιολόγησης του συστήματος. Η κυκλική αυτή δυναμική δημιουργεί έναν βρόχο ανατροφοδότησης (feedback loop), όπου κάθε μεροληψία που εισάγεται σε ένα σημείο του συστήματος μπορεί να ενισχυθεί και να μεταφερθεί στα υπόλοιπα, οδηγώντας σε συστηματική αναπαραγωγή ανισοτήτων. Συνεπώς, η κατανόηση και ο έλεγχος αυτού του μηχανισμού είναι ζωτικής σημασίας για την ανάπτυξη νομικών εφαρμογών τεχνητής νοημοσύνης που προάγουν τη δικαιοσύνη και την ισότητα.

3.3.3 Ποιότητα και Προσβασιμότητα Νομικών Δεδομένων

Η ποιότητα και η προσβασιμότητα των νομικών δεδομένων (legal data quality and accessibility) αποτελούν κρίσιμες παραμέτρους για την επιτυχημένη ανάπτυξη, εκπαίδευση και λειτουργία των Μεγάλων Γλωσσικών Μοντέλων στον τομέα της νομικής τεχνητής νοημοσύνης (Legal AI). Η επάρκεια, η εγκυρότητα και η αντιπροσωπευτικότητα του διαθέσιμου νομικού υλικού διαμορφώνουν άμεσα τα όρια και τις δυνατότητες των συστημάτων αυτών – καθορίζοντας το βαθμό ακρίβειας, τη γενίκευση σε νέες περιπτώσεις και την αξιοπιστία των προβλέψεων και της ανάλυσης που παράγεται [31][32].

Η ακρίβεια (accuracy) και η ορθότητα (correctness) του νομικού σώματος κειμένων είναι θεμελιώδεις προϋποθέσεις, καθώς η παρουσία λαθών, ασαφειών ή ασυνεπειών μπορεί να οδηγήσει τα μοντέλα σε σφάλματα ερμηνείας και αδικαιολόγητες προκαταλήψεις (bias). Τα παραδοσιακά νομικά δεδομένα χαρακτηρίζονται συχνά από ετερογένεια ως προς τη μορφοποίηση (format heterogeneity), πολυπλοκότητα ως προς τη γλώσσα και συχνά την απουσία ενιαίων σημάνσεων (annotation standards), με αποτέλεσμα να απαιτείται εντατική προ-

επεξεργασία (preprocessing) και κανονικοποίηση (normalization) πριν τη χρήση τους στην εκπαίδευση των LLMs [31]. Επιπλέον, η αντιπροσωπευτικότητα (representativeness) του συνόλου δεδομένων, δηλαδή το κατά πόσο καλύπτονται διαφορετικές κατηγορίες δικαίου, γεωγραφικές περιοχές, χρονικές περίοδοι και τύποι υποθέσεων, καθορίζει την ικανότητα του μοντέλου να γενικεύει σωστά [31]. Ειδικά στις πολυγλωσσικές και πολυδικαιοδοτικές εφαρμογές, η ποιοτική ανισότητα ανάμεσα στα δεδομένα διαφόρων γλωσσών ή νομικών παραδόσεων μπορεί να επιτείνει τις ανισότητες στη νομική τεχνητή νοημοσύνη, περιορίζοντας την προσβασιμότητα και την αποτελεσματικότητα για χρήστες πέραν του αγγλόφωνου δικαίου [31].

Από την άλλη πλευρά, τα ζητήματα προσβασιμότητας (accessibility) αποκτούν διαρκώς μεγαλύτερη σημασία, καθώς μεγάλος όγκος νομικής πληροφορίας προστατεύεται από δικαιώματα πνευματικής ιδιοκτησίας (copyright), εμποδίζοντας ή περιορίζοντας τη χρήση τους σε ανοιχτά δεδομένα (open data) για εκπαίδευση LLMs. Σε αρκετές δικαιοδοσίες, η πλήρης πρόσβαση σε δικαστικές αποφάσεις ή νομικά σχόλια απαιτεί συνδρομή ή παρέχεται αποκλειστικά σε θεσμικούς εταίρους, γεγονός που δημιουργεί ανισότητες στην έρευνα, την εκπαίδευση και την ανάπτυξη τεχνολογιών δικαιοσύνης [30][34]. Ταυτόχρονα, θεσμικά και γλωσσικά εμπόδια (institutional and linguistic barriers) παραμένουν: η πλειονότητα των νομικών datasets υψηλής ποιότητας διατίθεται μόνο στα αγγλικά, ενώ για τις μικρότερες ή λιγότερο τεκμηριωμένες έννομες τάξεις, οι συλλογές είναι περιορισμένες, ελλιπείς ή εντελώς απύλετες, δημιουργώντας ανισότητα στην ανάπτυξη και εφαρμογή συστημάτων Legal AI [34].

Οι επιπτώσεις από τα παραπάνω προβλήματα είναι ουσιώδεις και πολυεπίπεδες: η υποβάθμιση της απόδοσης των LLMs είναι άμεση, καθώς τα μοντέλα χάνουν τη δυνατότητα γενίκευσης και πέφτουν θύματα οερφίτινγκ σε μη αντιπροσωπευτικά ή θορυβώδη δεδομένα, παράγοντας λανθασμένες ή μεροληπτικές νομικές αναλύσεις. Επιπλέον, η περιορισμένη προσβασιμότητα επιβραδύνει την έρευνα, αναστέλλει την καινοτομία και συντηρεί το ψηφιακό χάσμα ανάμεσα σε δικαιοδοσίες ή κοινωνικές ομάδες με διαφορετικό επίπεδο τεχνολογικής ωριμότητας [32][34]. Από ηθικής και νομικής πλευράς, η χρήση μη ελεγμένων, αδειοδοτημένων ή επιμελώς επισημασμένων νομικών δεδομένων ενέχει κινδύνους παραβίασης του απορρήτου, της προστασίας προσωπικών δεδομένων και της ακριβούς εφαρμογής της αρχής της ίσης μεταχείρισης (fairness) [33].

Για τη βελτίωση της ποιότητας και της προσβασιμότητας των νομικών δεδομένων, η σύγχρονη βιβλιογραφία προτείνει τη δημιουργία ανοικτών, διαλειτουργικών και θεσμικά ελεγμένων νομικών συλλογών καθώς και τη συστηματική επιμέλεια των δεδομένων μέσω τεχνικών κανονικοποίησης, εμπλουτισμού με μεταδεδομένα και θεματικής ποικιλομορφίας. Παράλληλα, αναδεικνύεται η σημασία της διεθνούς συνεργασίας, τόσο μεταξύ νομικών όσο και τεχνικών εταίρων, για την υιοθέτηση κοινών προτύπων και πρακτικών ανοικτών δεδομένων, που διασφαλίζουν την ποιοτική, ασφαλή και δίκαιη χρήση της τεχνητής νοημοσύνης στο δίκαιο [31][34].

3.4 Συμπεράσματα

Η παρουσίαση του θεωρητικού υπόβαθρου στο Κεφάλαιο 3 καταδεικνύει ότι η αποτελεσματική αξιοποίηση των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στον χώρο του δικαίου

απαιτεί πολύ περισσότερα από την απλή ενσωμάτωση μιας τεχνολογικά προηγμένης λύσης. Παρά τις εντυπωσιακές δυνατότητες που προσφέρουν στην επεξεργασία και ανάλυση νομικού κειμένου, η πολυπλοκότητα και η αυστηρότητα της νομικής πληροφορίας, σε συνδυασμό με την αυξημένη ευθύνη που συνεπάγεται κάθε νομική απόφαση, επιβάλλουν τη δημιουργία μοντέλων ρητά προσαρμοσμένων και ευθυγραμμισμένων με τις θεσμικές, κοινωνικές και ηθικές απαιτήσεις του κλάδου.

Το πραγματικό στοίχημα δεν έγκειται μόνο στην τεχνολογική πρόοδο, αλλά στην ικανότητα των ερευνητών και νομικών επαγγελματιών να διαμορφώσουν τις κατάλληλες προϋποθέσεις ώστε η τεχνολογία αυτή να λειτουργήσει ως ουσιαστικό εργαλείο υποστήριξης, αποφεύγοντας να γίνει πηγή κινδύνου ή αμφιβολιών. Η επιτυχία των LLMs στον νομικό τομέα θα κριθεί όχι μόνο από την τεχνική τους αρτιότητα, αλλά και από την ικανότητα παροχής διαφανών, επαληθεύσιμων και αμερόληπτων λύσεων που ενσωματώνονται αρμονικά στην καθημερινή νομική πρακτική, ενισχύοντας την εμπιστοσύνη τόσο των φορέων όσο και των πολιτών.

Επιπλέον, η πρόκληση αυτή απαιτεί διεπιστημονική συνεργασία, όπου νομικοί, τεχνικοί και ρυθμιστικοί φορείς συνεργάζονται στενά για τον σχεδιασμό, την αξιολόγηση και τη συνεχή βελτίωση των συστημάτων τεχνητής νοημοσύνης, εξασφαλίζοντας τη συμμόρφωση με τα ηθικά και θεσμικά πρότυπα που διέπουν το νομικό οικοσύστημα. Μόνο έτσι τα LLMs θα μπορέσουν να εξελιχθούν σε εργαλεία που πραγματικά υποστηρίζουν την απονομή δικαιοσύνης και την πρόσβαση σε αυτήν, συνεισφέροντας σε ένα δίκαιο, διαφανές και προσβάσιμο μέλλον.

Πρακτικές εφαρμογές και τεχνική υλοποίηση

Η είσοδος της τεχνητής νοημοσύνης στη νομική επιστήμη σηματοδοτεί μια νέα εποχή για την ανάλυση, επεξεργασία και αξιοποίηση νομικής πληροφορίας. Οι σύγχρονες τεχνολογικές λύσεις, με αιχμή του δόρατος τα Μεγάλα Γλωσσικά Μοντέλα (LLMs), επαναπροσδιορίζουν τόσο τις μεθόδους αναζήτησης και κατηγοριοποίησης δεδομένων όσο και τον τρόπο με τον οποίο προσεγγίζονται σύνθετα νομικά ερωτήματα και λαμβάνονται αποφάσεις.

Η ραγδαία αύξηση της διαθεσιμότητας ψηφιακών δεδομένων, η ανάγκη για γρήγορη και αξιόπιστη επεξεργασία πληροφοριών, αλλά και οι απαιτήσεις διαφάνειας και ακρίβειας, καθιστούν αναγκαία την υιοθέτηση προηγμένων εργαλείων σε όλο το φάσμα της νομικής πρακτικής. Από την ανάκτηση εγγράφων και την αυτόματη ανάλυση συμβολαίων έως την πρόβλεψη εκβάσεων υποθέσεων και την υποστήριξη του έργου των δικαστηρίων, η τεχνολογία μετασχηματίζει τις διαδικασίες και ενισχύει τον ρόλο του νομικού επαγγελματία.

Στο παρόν κεφάλαιο παρουσιάζονται οι βασικές πρακτικές εφαρμογές των LLMs και των συναφών τεχνολογιών στον τομέα του δικαίου, οι μέθοδοι αξιολόγησης των συστημάτων αυτών, καθώς και οι προκλήσεις που προκύπτουν στη διασφάλιση της αξιοπιστίας και της νομικής ακρίβειας.

4.1 Νομική Έρευνα και Ανάκτηση Πληροφορίας

Η νομική έρευνα και η ανάκτηση πληροφορίας (Legal Research and Information Retrieval) συνιστούν θεμέλιο λίθο για κάθε δικαιοϋική διαδικασία, είτε πρόκειται για δικαστική επίλυση διαφορών, σύνταξη συμβολαίων, ή γνωμοδοτήσεις σε σύνθετα νομικά ζητήματα. Παραδοσιακά, η έρευνα αυτή στηριζόταν στην αναζήτηση νομοθεσίας, νομολογίας και βιβλιογραφίας μέσα από φυσικά αρχεία ή απλά ηλεκτρονικά συστήματα αναζήτησης, με κύριο εργαλείο τη λεκτική ταύτιση λέξεων-κλειδίων. Ωστόσο, η ταχεία πρόοδος της τεχνητής νοημοσύνης και ειδικότερα των Μεγάλων Γλωσσικών Μοντέλων, έχει ανατρέψει ριζικά το τοπίο: τα σύγχρονα νομικά πληροφοριακά συστήματα διαθέτουν πλέον δυνατότητα σημασιολογικής αναζήτησης, κατηγοριοποίησης και αυτόματης ανάλυσης τεράστιου όγκου νομικών κειμένων, προσφέροντας στον νομικό επαγγελματία εργαλεία που υπερβαίνουν κατά πολύ τις παραδοσιακές μεθόδους.

Η μετάβαση από την απλή αναζήτηση βάσει λέξεων-κλειδίων (keyword-based retrieval) σε πολυεπίπεδες, υβριδικές αρχιτεκτονικές, όπως τα συστήματα με ενσωματωμένες νευρωνικές αναπαραστάσεις (embeddings) και pipelines τύπου Retrieval-Augmented Generation

(RAG), δίνει τη δυνατότητα εντοπισμού κρίσιμων νομικών τεκμηρίων ακόμα και όταν διατυπώνονται με διαφορετική ορολογία ή εντάσσονται σε διαφορετικά θεσμικά πλαίσια. Η εφαρμογή τεχνικών σημασιολογικής ομοιότητας (semantic similarity), η χρήση προεκπαιδευμένων γλωσσικών μοντέλων όπως τα LegalBERT, Lawformer, και η αξιοποίηση πολυγλωσσικών εργαλείων έχουν επεκτείνει τη λειτουργικότητα των legal research systems σε παγκόσμιο επίπεδο.

Επιπλέον, οι σύγχρονες πλατφόρμες ενσωματώνουν πλέον δυνατότητες αυτόματης κατηγοριοποίησης εγγράφων (document classification), ανάλυσης ομοιότητας υποθέσεων (case similarity analysis) και επεξηγήσιμης ανάκτησης (explainable retrieval), συνδυάζοντας προηγμένη αναζήτηση (advanced search) με νομικά αναλυτικά εργαλεία (legal analytics) και υπομονάδες συστάσεων (recommendation modules). Η εξέλιξη αυτή έχει οδηγήσει στην ανάπτυξη εξειδικευμένων συνόλων δεδομένων (datasets) και πλαισίων σύγκρισης επιδόσεων (benchmarking frameworks), όπως το LeCaRD[3] και το LexGLUE[42], που θέτουν τα επιστημονικά πρότυπα (standards) για την αποτίμηση της ακρίβειας, της πληρότητας, της επεξηγηματικότητας (explainability) και της πολυγλωσσικής επάρκειας (multilingual competence) των συστημάτων αυτών.

Ωστόσο, η ενσωμάτωση των LLMs στα νομικά συστήματα ανάκτησης πληροφορίας εγείρει και νέες προκλήσεις. Τα θέματα της ερμηνευσιμότητας (explainability), της ιχνηλασιμότητας παραπομπών (citation provenance), της διασφάλισης νομικής εγκυρότητας και της προσαρμογής σε διαφορετικά θεσμικά πλαίσια παραμένουν αντικείμενο έντονης επιστημονικής συζήτησης. Επιπλέον, η διαχείριση της ασάφειας, των διφορούμενων ερωτημάτων και των πολυεπίπεδων νομικών εννοιών απαιτεί συνεχή εξέλιξη των αλγοριθμικών και γνωσιακών μοντέλων.

Η ενότητα 4.1 εστιάζει ακριβώς σε αυτή τη νέα πραγματικότητα: παρουσιάζει τα τεχνολογικά θεμέλια, τις εξελιγμένες μεθοδολογίες και τις πρακτικές εφαρμογές των σύγχρονων συστημάτων νομικής ανάκτησης πληροφορίας, εστιάζοντας τόσο στις δυνατότητες όσο και στους περιορισμούς που ανακύπτουν στη μετάβαση από τα «παραδοσιακά» πλαίσια ανάκτησης πληροφορίας (IR frameworks) στα σύγχρονα, νομικά ευαισθητοποιημένα (legal-aware) και επεξηγήσιμα συστήματα τεχνητής νοημοσύνης (explainable AI pipelines). Μέσα από την ανάλυση διεθνών παραδειγμάτων, τεχνικών υλοποιήσεων και επιστημονικών μετρικών, επιδιώκεται να αποτυπωθεί η πολυπλοκότητα, το εύρος και οι προοπτικές του πεδίου αυτού στη σύγχρονη νομική πρακτική.

4.1.1 Συστήματα Ανάκτησης Νομικών Εγγράφων

Η ανάκτηση νομικών εγγράφων (document retrieval) αποτελεί θεμέλιο της σύγχρονης νομικής επιστήμης και πρακτικής, καθώς επιτρέπει την ταχεία, στοχευμένη και τεκμηριωμένη πρόσβαση σε νομοθεσία, νομολογία, διοικητικές πράξεις και ερμηνευτική βιβλιογραφία. Τα τελευταία χρόνια, η πρόοδος στις τεχνολογίες επεξεργασίας φυσικής γλώσσας (Natural Language Processing) και Μεγάλων Γλωσσικών Μοντέλων (LLMs) έχει αναβαθμίσει ριζικά τις δυνατότητες των συστημάτων ανάκτησης πληροφορίας, φέρνοντας τον νομικό ερευνητή πιο κοντά σε σημασιολογικά πλούσια και επεξηγηματικά αποτελέσματα [43], [2].

Παραδοσιακά, τα νομικά συστήματα αναζήτησης βασίζονταν στην αντιστοίχιση λέξεων-

κλειδιών (lexical retrieval), όπου ο χρήστης έπρεπε να χρησιμοποιήσει ακριβή νομική ορολογία, ενώ η αντιστοίχιση βασιζόταν κυρίως σε μοντέλα τύπου TF-IDF, λογικούς τελεστές (Boolean retrieval) και φίλτρα μεταδεδομένων. Αυτές οι μέθοδοι, αν και πρακτικές, αδυνατούσαν να συλλάβουν τη σημασιολογική συνάφεια ανάμεσα σε διαφορετικά διατυπωμένες αλλά ουσιαστικά ταυτόσημες νομικές έννοιες, οδηγώντας είτε σε χαμηλή ανάκληση (recall) είτε σε υπερφόρτωση με άσχετα αποτελέσματα [43].

Η έλευση των μοντέλων νευρωνικών ενθυλακώσεων (embeddings) όπως τα LegalBERT[44], Lawformer[37] ή JuriBERT[45], επέτρεψε τη μετατροπή τόσο των νομικών εγγράφων όσο και των ερωτημάτων σε διανυσματικούς χώρους υψηλής διάστασης. Μέσω των embeddings, το σύστημα είναι πλέον σε θέση να υπολογίζει τη σημασιολογική εγγύτητα (semantic similarity) μεταξύ κειμένων, ακόμη και αν δεν υπάρχουν κοινές λέξεις-κλειδιά [2], [35].

Στην πράξη, τα σύγχρονα συστήματα επεξεργασίας (pipelines) ακολουθούν μια πολυεπίπεδη διαδικασία. Αρχικά, όλα τα έγγραφα του σώματος κειμένων (corpus) μετατρέπονται σε διανύσματα (vector embeddings) με τη χρήση βελτιστοποιημένων γλωσσικών μοντέλων (fine-tuned language models). Τα διανύσματα αυτά αποθηκεύονται σε εξειδικευμένες βάσεις δεδομένων ταχείας αναζήτησης[2] (vector databases), όπως οι FAISS, Weaviate ή Milvus, επιτρέποντας την άμεση και αποδοτική ανάκτηση σχετικών εγγράφων.

Όταν ο χρήστης υποβάλλει ένα ερώτημα, αυτό μετατρέπεται επίσης σε διάνυσμα και το σύστημα εντοπίζει τα πιο συναφή αποτελέσματα, χρησιμοποιώντας μετρικές όπως η συνημιτονοειδής ομοιότητα (cosine similarity). Η τελική κατάταξη και επιλογή των εγγράφων γίνεται συνδυαστικά, βάσει νομικών μεταδεδομένων (όπως το αρμόδιο δικαστήριο, η ημερομηνία ή ο τύπος εγγράφου), αλλά και σημασιολογικών βαθμολογιών (semantic scores)[4]. Με αυτόν τον τρόπο διασφαλίζεται ότι το σύστημα επιστρέφει όχι μόνο σχετικά, αλλά και νομικά έγκυρα αποτελέσματα.

Η καινοτομία των συστημάτων ανάκτησης με ενισχυμένη δημιουργία (Retrieval-Augmented Generation, RAG) βασίζεται ακριβώς σε αυτή τη δυνατότητα. Σε ένα τέτοιο σύστημα, ο χρήστης υποβάλλει ένα ερώτημα σε φυσική γλώσσα, το οποίο στη συνέχεια επεξεργάζεται από τον μηχανισμό ανάκτησης (retriever), ο οποίος εντοπίζει τα πλέον σχετικά τεκμήρια μέσω σημασιολογικής ανάκτησης (semantic retrieval). Ακολούθως, ένα Μεγάλο Γλωσσικό Μοντέλο (LLM) παράγει την τελική απάντηση ή μια συνοπτική ανάλυση, ενσωματώνοντας αποσπάσματα από τα ανακτημένα έγγραφα [2], [46].

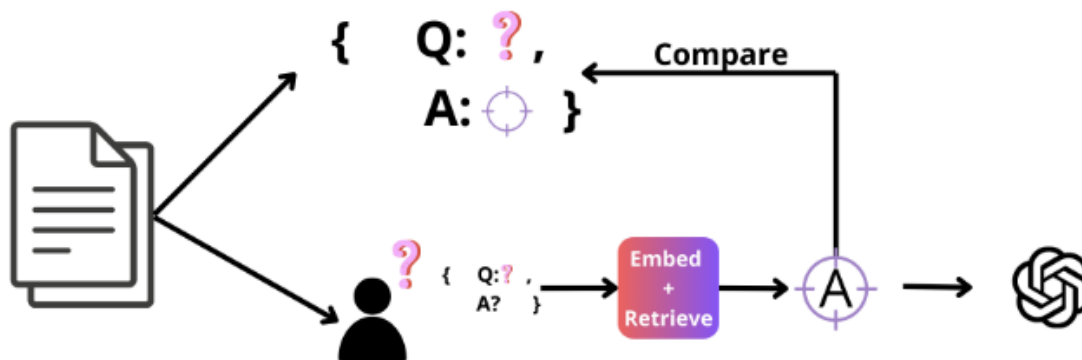
Τα σύγχρονα συστήματα ανάκτησης νομικής πληροφορίας με ενισχυμένη γενετική δυνατότητα (Retrieval-Augmented Generation – RAG) αξιολογούνται μέσω εξειδικευμένων πλαισίων σύγκρισης (benchmarking frameworks), όπως το LegalBench-RAG. Η αρχιτεκτονική RAG συνδυάζει έναν μηχανισμό ανάκτησης σχετικών εγγράφων (retriever) με ένα μεγάλο γλωσσικό μοντέλο παραγωγής (generator), ώστε οι απαντήσεις να βασίζονται όχι μόνο στην εκπαίδευση του μοντέλου αλλά και σε επίκαιρο, εξωτερικό υλικό. Τα συστήματα αυτά επιδιώκουν να αυξήσουν την ακρίβεια και τη διαφάνεια των παραγόμενων απαντήσεων. Τα πλαίσια αξιολόγησης, όπως το LegalBench-RAG, μετρούν με ακρίβεια τόσο την πληρότητα όσο και την πιθανότητα παραγωγής παραληρηματικών απαντήσεων (hallucinations), προσφέροντας αντικειμενικά metrics για κάθε φάση του pipeline [46].

Παράλληλα, σύνολα δεδομένων όπως το LeCaRD (Legal Case Retrieval Dataset) παρέχουν μεγάλης κλίμακας δεδομένα πραγματικών νομικών ερωτημάτων με επισημάνσεις

από ειδικούς, επιτρέποντας τη ρεαλιστική δοκιμή των retrieval models [3]. Εμπορικές πλατφόρμες, όπως οι Lexis+ AI[47] και Westlaw Edge[48], αξιοποιούν προηγμένους αλγορίθμους σημασιολογικής αναζήτησης σε πραγματικό χρόνο, αν και οι ακριβείς τεχνικές τους παραμένουν κλειστές (proprietary).

Η βασική ροή επεξεργασίας (pipeline) τέτοιων συστημάτων περιλαμβάνει αρχικά τον κατακερματισμό των νομικών εγγράφων (document chunking), όπου μεγάλα κείμενα διαχωρίζονται σε μικρότερα αποσπάσματα με επικαλύψεις. Ακολουθεί ο καθαρισμός του κειμένου για την απομάκρυνση άσχετων στοιχείων και μεταδεδομένων, και στη συνέχεια η μετατροπή των αποσπασμάτων σε διανυσματικές αναπαραστάσεις μέσω εξειδικευμένων μοντέλων ενσωμάτωσης (embedding models), όπως τα LegalBERT, all-MiniLM-L6-v2, LegalBERT-pt, ή BERTimbau [49]. Οι ενσωματώσεις αυτές αποθηκεύονται σε ειδικά ευρετήρια διανυσμάτων (vector stores), όπως το FAISS[2] ή το Weaviate[50], επιτρέποντας ταχύτατη αναζήτηση. Όταν εισάγεται ένα νέο ερώτημα, κωδικοποιείται αντίστοιχα σε διάνυσμα και το ευρετήριο επιστρέφει τα αποσπάσματα με τη μεγαλύτερη σημασιολογική συνάφεια.

Στο Σχήμα 4.1 απεικονίζεται σχηματικά η διαδικασία αξιολόγησης του σταδίου ανάκτησης (retrieval step) σε τέτοια συστήματα, όπου η απόδοση του μηχανισμού ανάκτησης συγκρίνεται με τα gold δεδομένα αναφοράς για τη μέτρηση της συνάφειας και της αξιοπιστίας των απαντήσεων.



Σχήμα 4.1: Διαδικασία αξιολόγησης της φάσης ανάκτησης (retrieval step) σε συστήματα Retrieval-Augmented Generation (RAG). Τα ανακτηθέντα αποσπάσματα συγκρίνονται με τα χρυσά δεδομένα (gold data) για την αποτίμηση της συνάφειας και της ποιότητας της ανάκτησης.

Η διαδικασία αυτή ενισχύει την αναζήτηση πέρα από απλές λέξεις-κλειδιά, αξιοποιώντας την πραγματική νομική συνάφεια. Ειδικά σε γλώσσες ή δικαιοδοσίες με περιορισμένους πόρους (low-resource systems), η προσαρμογή πολυγλωσσικών μοντέλων, όπως το XLM-Roberta[42], ή η χρήση εξειδικευμένων μοντέλων (LegalBERT-pt, BERTimbau), έχει συμβάλει σημαντικά στην ακρίβεια και τη λειτουργικότητα των συστημάτων αυτών [49][51].

Πέρα από τα κλασικά συστήματα ανάκτησης με πυκνές αναπαραστάσεις (dense retrieval pipelines), σήμερα αναπτύσσονται υβριδικά μοντέλα που συνδυάζουν αραιή (sparse, π.χ. TF-IDF/λέξεις-κλειδιά) και πυκνή (dense, δηλαδή σημασιολογική) ανάκτηση. Παράλ-

ληλα, προχωρημένα συστήματα ενσωματώνουν γνωσιακά γραφήματα (knowledge graphs, KG-RAG), όπου οι οντολογίες και οι σχέσεις του δικαίου κατευθύνουν την αναζήτηση σε θεματικούς κόμβους, επιτρέποντας περισσότερο επεξηγήσιμα (explainable) αποτελέσματα [52]. Χαρακτηριστικό παράδειγμα σύγχρονης αρχιτεκτονικής RAG αποτελεί το σύστημα LawPal, το οποίο έχει υλοποιηθεί για το ινδικό νομικό σύστημα[2]. Το LawPal ενσωματώνει τεχνικές προεπεξεργασίας, κατακερματισμού και σημασιολογικής ενσωμάτωσης εγγράφων, ενώ αξιοποιεί το FAISS για ιεραρχικό indexing και ταχύτατη ανάκτηση σχετικών αποσπασμάτων.

Η απόδοση των συστημάτων ανάκτησης νομικών εγγράφων μετριέται συνήθως με recall@k, μέσο όρο ακρίβειας (mean average precision, MAP), κανονικοποιημένο αθροιστικό βαθμό κατάταξης (nDCG), αλλά και με νομικά εξειδικευμένες μετρικές, όπως το ποσοστό παραισθήσεων (hallucination rate) και η συνέπεια παραπομπών (citation consistency), δηλαδή το κατά πόσο οι παραπομπές ενσωματώνονται ορθά στις τελικές απαντήσεις [46], [53].

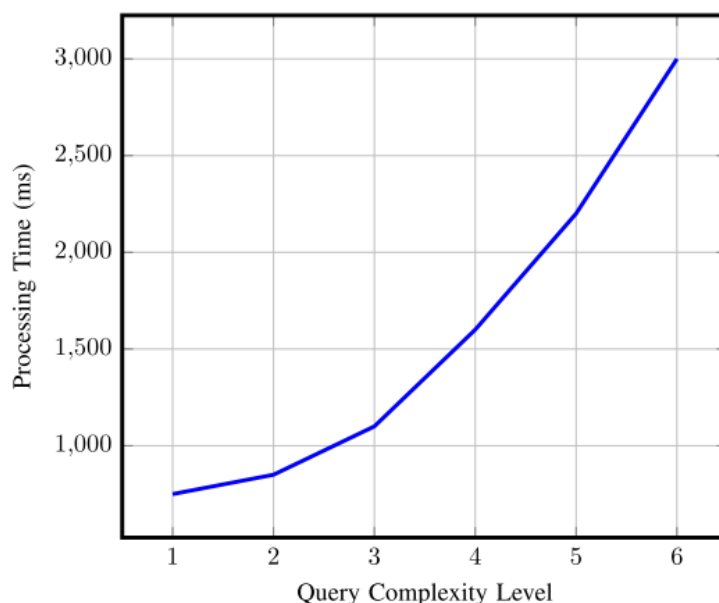
Ταυτόχρονα, εξακολουθούν να υπάρχουν σημαντικές προκλήσεις. Πρώτον, η επαληθευσσιμότητα (verifiability) των απαντήσεων αποτελεί διαρκές ζητούμενο, ιδιαίτερα στα συστήματα RAG, όπου τα Μεγάλα Γλωσσικά Μοντέλα (LLMs) μπορεί να παράγουν μη τεκμηριωμένες πληροφορίες [2], [46]. Επιπλέον, παρατηρείται περιορισμένη διαφάνεια στα εμπορικά προγραμματιστικά περιβάλλοντα (commercial APIs), γεγονός που δυσχεραίνει τον έλεγχο και την αξιολόγηση των αποτελεσμάτων. Τέλος, παραμένει χαμηλή η κάλυψη σε γλώσσες ή έννομες τάξεις με περιορισμένα διαθέσιμα δεδομένα.

Η αξιολόγηση της απόδοσης των συστημάτων RAG στη νομική έρευνα πραγματοποιείται συνδυάζοντας μετρικές ανάκτησης (π.χ. Precision@K, MRR, nDCG) και γενετικής συνέπειας (π.χ. BLEU, ROUGE, Legal Consistency Score). Όπως αναδεικνύεται στην πρόσφατη μελέτη για το LawPal[2], η χρήση του FAISS προσφέρει σημαντικά πλεονεκτήματα σε recall και ταχύτητα έναντι άλλων επιλογών (π.χ. Chroma), ενώ το σύστημα επιτυγχάνει συνολική ακρίβεια άνω του 90% στην αναζήτηση και ερμηνεία νομικών δεδομένων. Τα παραπάνω ποσοστά προκύπτουν από συνδυασμό αυτοματοποιημένων μετρικών (π.χ. Legal Consistency Score, BLEU, ROUGE) και ποιοτικής αξιολόγησης από εξειδικευμένους νομικούς εμπειρογνώμονες. Παράλληλα, η ανάλυση ικανοποίησης χρηστών που συμμετείχαν στη δοκιμή της πλατφόρμας καταγράφει ιδιαίτερα υψηλό ποσοστό (85%), επιβεβαιώνοντας τη χρηστικότητα και αξιοπιστία του συστήματος στην πράξη [2].

Πέρα από τις παραδοσιακές μετρικές ανάκτησης, ιδιαίτερη σημασία αποδίδεται στην υπολογιστική αποδοτικότητα των συστημάτων RAG, ειδικά όταν εφαρμόζονται σε μεγάλα, ετερογενή νομικά σύνολα δεδομένων. Στο Σχήμα 4.2 παρουσιάζεται η ανάλυση χρόνου επεξεργασίας ερωτημάτων στο σύστημα LawPal, ανάλογα με το επίπεδο πολυπλοκότητας των queries. Παρατηρείται ότι, παρά την αυξανόμενη πολυπλοκότητα, το σύστημα διατηρεί ικανοποιητικούς χρόνους απόκρισης, διασφαλίζοντας real-time legal assistance. Η αποτελεσματική χρήση του FAISS για retrieval και η αρχιτεκτονική με DeepSeek-R1:5B επιτρέπουν υψηλή κλιμάκωση και σταθερότητα απόδοσης, ακόμη και υπό έντονα φορτία [2].

Συμπερασματικά, τα σύγχρονα συστήματα ανάκτησης νομικών εγγράφων έχουν αλλάξει ριζικά το τοπίο της νομικής έρευνας. Από τα άκαμπτα συστήματα βασισμένα σε λέξεις-κλειδιά (keyword-based search), περάσαμε σε ιδιαίτερα ευέλικτα, πολυγλωσσικά και επεξηγήσιμα (explainable) συστήματα επεξεργασίας (pipelines), με τεράστιες προοπτικές αλλά και ουσιαστικές τεχνικές προκλήσεις. Ο διαρκής συνδυασμός πρακτικής μηχανικής υλοποίησης

και επιστημονικής αξιολόγησης θα παραμείνει στο επίκεντρο της εξέλιξης των συστημάτων νομικής ανάκτησης πληροφορίας (legal IR systems) τα επόμενα χρόνια.



Σχήμα 4.2: Ανάλυση χρόνου επεξεργασίας ερωτημάτων στο σύστημα LawPal, σε συνάρτηση με το επίπεδο πολυπλοκότητας των queries. Το σύστημα διατηρεί χαμηλούς χρόνους απόκρισης ακόμη και σε απαιτητικά σενάρια, γεγονός που ενισχύει τη χρηστικότητα και την αποδοτικότητα σε νομικές εφαρμογές πραγματικού χρόνου [2].

4.1.2 Ανάλυση Ομοιότητας Νομικών Υποθέσεων

Η ανάλυση ομοιότητας νομικών υποθέσεων (legal case similarity analysis) αποτελεί μία από τις πιο κρίσιμες λειτουργίες της νομικής πληροφορικής, καθώς ενισχύει τη δυνατότητα εντοπισμού σχετικών προηγούμενων, τη συγκριτική εξέταση αποφάσεων και τη διαμόρφωση συνεκτικών επιχειρημάτων σε πραγματικά δικαστικά ή ερευνητικά σενάρια. Η εξέλιξη των τεχνολογιών επεξεργασίας φυσικής γλώσσας και των LLMs έχει επιτρέψει την ανάπτυξη εργαλείων που υπερβαίνουν τα όρια της λεκτικής ομοιότητας και αξιοποιούν τη σημασιολογική, θεματική και δομική συνάφεια μεταξύ νομικών κειμένων [43], [2].

Η έννοια της ομοιότητας στη νομική πρακτική είναι πολυδιάστατη. Δεν αφορά απλώς την ταύτιση όρων, αλλά εκτείνεται στη σύγκριση νομικών ζητημάτων, πραγματικών περιστατικών, επιχειρηματολογίας και αποτελέσματος. Για παράδειγμα, δύο αποφάσεις μπορεί να θεωρούνται συναφείς αν αφορούν παρόμοια νομικά ζητήματα, ακόμα κι αν διαφέρουν οι επιμέρους διατυπώσεις ή τα πραγματικά περιστατικά. Αντίστοιχα, μπορεί να εμφανίζεται ομοιότητα σε επίπεδο εφαρμοστέων διατάξεων ή νομικών συλλογισμών, ακόμα και αν οι λέξεις διαφέρουν σημαντικά. Το γεγονός αυτό αναδεικνύει τη σημασία της σημασιολογικής εγγύτητας (semantic similarity) ως πρωτεύοντος κριτηρίου σε σύγχρονα πληροφοριακά συστήματα [43], [35].

Στην καρδιά αυτής της μετάβασης βρίσκονται τα τεχνολογικά εργαλεία των ενσωματώσεων (embeddings), όπου κάθε νομικό κείμενο μετατρέπεται σε πολυδιάστατο διάνυσμα που αποτυπώνει όχι μόνο το λεκτικό περιεχόμενο, αλλά και τις σημασιολογικές, δομικές και

θεματικές αποχρώσεις της υπόθεσης. Μοντέλα όπως το LegalBERT, το Lawformer[37], το all-MiniLM-L6-v2 προσφέρουν τη δυνατότητα αναπαράστασης περίπλοκων νομικών δεδομένων με μεγάλη πιστότητα [35], [4], [51].

Η ανάλυση ομοιότητας υλοποιείται συνήθως μέσω μιας αλληλουχίας βημάτων. Αρχικά, τα νομικά έγγραφα διασπώνται σε σημασιολογικά αποσπάσματα και υφίστανται προεπεξεργασία (καθαρισμό, κανονικοποίηση, αφαίρεση θορύβου). Στη συνέχεια, κάθε απόσπασμα μετατρέπεται σε διανυσματική αναπαράσταση με εκπαίδευση προσαρμοσμένων μετασχηματιστών (fine-tuned transformers). Το ίδιο γίνεται και με το ερώτημα ή τη νέα υπόθεση που εξετάζεται. Ο υπολογισμός της ομοιότητας πραγματοποιείται με μέτρα όπως η γωνιακή ομοιότητα (cosine similarity) ή η ευκλείδεια απόσταση, και στη συνέχεια επιλέγονται τα αποσπάσματα ή οι υποθέσεις με το μεγαλύτερο βαθμό συνάφειας [46]. Προχωρημένα συστήματα υλοποιούν πολυεπίπεδες ή υβριδικές αναπαραστάσεις, όπου κάθε υπόθεση αποκτά ξεχωριστά εμβεδδινγς για κρίσιμες διαστάσεις, όπως τα πραγματικά γεγονότα, τα νομικά ζητήματα, και το αποτέλεσμα της δίκης. Αυτή η προσέγγιση, γνωστή ως hierarchical embedding, επιτρέπει μεγαλύτερη ακρίβεια και ευελιξία στην ανάλυση και έχει αποδειχθεί αποτελεσματική σε μεγάλης κλίμακας νομικά datasets, όπως το LeCaRD [3].

Παράλληλα, η αξιοποίηση γράφων γνώσης (knowledge graphs) και αλγορίθμων νευρωνικών δικτύων πάνω σε γράφους (Graph Neural Networks – GNN) έχει δώσει νέα ώθηση στην ανάλυση ομοιότητας, επιτρέποντας την αποκάλυψη κρυφών θεματικών ή ρητορικών συσχετίσεων μεταξύ υποθέσεων. Τα πλαίσια ανάκτησης βάσει γραφημάτων (graph-based retrieval frameworks) αξιοποιούν πληροφορίες από δίκτυα παραπομπών (citation networks) και ιεραρχικές σχέσεις (hierarchical relations), προσφέροντας τη δυνατότητα εντοπισμού συσσωματώσεων (clusters) αποφάσεων που παρουσιάζουν παρόμοιο νομικό σκεπτικό [49], [52].

Η επιτυχία αυτών των συστημάτων αξιολογείται με χρήση εξειδικευμένων συνόλων δεδομένων (datasets) και προτύπων αξιολόγησης (benchmarks). Το LeCaRD (Legal Case Retrieval Dataset) αποτελεί χαρακτηριστικό παράδειγμα, καθώς παρέχει χιλιάδες ερωτήματα (queries) και εκατοντάδες χιλιάδες νομικές υποθέσεις με επισήμανση συνάφειας από ειδικούς (expert labeling). Ένα ενδεικτικό παράδειγμα ζεύγους ερωτήματος και υποψήφιας υπόθεσης παρουσιάζεται στο Σχήμα 4.3, προσφέροντας αξιόπιστη βάση για την αξιολόγηση της απόδοσης συστημάτων ανάκτησης νομικών εγγράφων.

Ένα από τα σημαντικότερα προβλήματα που αναδεικνύονται από τη βιβλιογραφία είναι το γεγονός ότι η συνάφεια δεν είναι πάντα αντικειμενική. Μηχανές και άνθρωποι διαφωνούν συχνά για το ποια υπόθεση είναι πραγματικά όμοια, ενώ ακόμα και μεταξύ νομικών η αξιολόγηση μπορεί να ποικίλει ανάλογα με το πλαίσιο και το είδος της νομικής ερώτησης [53]. Έτσι, η αξιόπιστη ανάλυση ομοιότητας απαιτεί συνδυασμό αυτόματων μετρικών, αξιολόγησης από ειδικούς (expert review), αλλά και επεξηγησιμότητας (explainability) — δηλαδή τη δυνατότητα του συστήματος να εξηγήει πώς και γιατί κατέληξε σε μια συγκεκριμένη συσχέτιση. Σύγχρονες προσεγγίσεις, όπως τα εργαλεία οπτικοποίησης προσοχής (attention visualization tools) και η ανίχνευση παραπομπών (citation tracing), επιτρέπουν την ερμηνεία της πορείας λήψης αποφάσεων από το σύστημα, προσφέροντας διαφάνεια και τεκμηρίωση, απαραίτητη για τη νομική πρακτική [52].

Συνολικά, η ανάλυση ομοιότητας νομικών υποθέσεων αποτελεί σήμερα πεδίο ταχείας

επιστημονικής και τεχνολογικής εξέλιξης, όπου τα νομικά ευαισθητοποιημένα εμβεδδινγς (legal-aware embeddings), τα γνωσιακά γραφήματα (knowledge graphs) και τα πρότυπα αξιολόγησης (benchmarking frameworks) διαμορφώνουν νέα πρότυπα για τη νομική έρευνα, προσφέροντας μεγαλύτερη ακρίβεια, ταχύτητα και διαφάνεια στην αναζήτηση συναφών υποθέσεων.

Query	Candidate
Case Description: From January 3, 2017 to March 12, 2019, the defendant A has illegally sold Mark Six through WeChat and bank card transfers ...	Case Name: The case of B illegally opening a casino Case Description: On September 19, 2019, the defendant B used WeChat to sale the Mark Six Lottery in a supermarket illegally ... Judgment: Crime of opening a casino ... Label: 1 (Very relevant)

Σχήμα 4.3: Παράδειγμα ζεύγους ερωτήματος (query) και υποψήφιας υπόθεσης (candidate case) από το σύνολο δεδομένων LeCaRD.[3].

4.1.3 Τεχνικά Παραδείγματα Υλοποίησης

Η πρακτική υλοποίηση συστημάτων νομικής ανάκτησης πληροφορίας με αξιοποίηση Μεγάλων Γλωσσικών Μοντέλων (LLMs) απαιτεί συνδυαστική αξιοποίηση τεχνολογιών επεξεργασίας φυσικής γλώσσας, διανυσματικών αναπαραστάσεων (embeddings), indexing υψηλής απόδοσης, και αρχιτεκτονικών τύπου Retrieval-Augmented Generation (RAG). Τα παρακάτω παραδείγματα αναδεικνύουν τη διαδρομή από την αρχική ανάλυση και τμηματοποίηση των εγγράφων έως την αξιολόγηση της απόδοσης του συστήματος σε πραγματικές νομικές ροές εργασίας, στηριγμένα στη διεθνή βιβλιογραφία.

Ένα τυπικό σύστημα σημασιολογικής αναζήτησης (semantic search pipeline) για νομικά έγγραφα ξεκινά με τη συλλογή και προκαταρκτική επεξεργασία μεγάλου αριθμού νομικών κειμένων, τα οποία συνήθως τεμαχίζονται σε μικρότερα αποσπάσματα μεγέθους 512 tokens με επικάλυψη (overlapping), ώστε να διατηρείται η θεματική συνοχή και να αποφεύγονται κοπές νοηματικών σημαντικών εννοιών. Ακολουθεί η μετατροπή κάθε αποσπάσματος σε διανυσματική αναπαράσταση (vector representation) μέσω προεκπαιδευμένων μοντέλων όπως το LegalBERT, που έχει αναπτυχθεί ειδικά για νομικά κείμενα και προσφέρει βελτιωμένη σημασιολογική απόδοση συγκριτικά με τα γενικά μοντέλα [35]. Τα προκύπτοντα διανύσματα, συνήθως διαστάσεων 768, αποθηκεύονται σε ειδικό ευρετήριο υψηλής απόδοσης (high-performance index), όπως το FAISS, που παρέχει γραμμική ή παραμετρική αναζήτηση (μέσω δομών IVF ή HNSW) για ταχύτατη ανάκτηση ακόμη και σε σώματα εκατοντάδων χιλιάδων εγγράφων [54].

Όταν ο χρήστης θέτει ένα ερώτημα, αυτό μετατρέπεται σε ενσωμάτωση (embedding) μέσω του ίδιου μοντέλου, και η αναζήτηση πραγματοποιείται με βάση τη συνημιτονοειδή

ομοιότητα (cosine similarity) μεταξύ του ερωτήματος και των αποθηκευμένων διανυσμάτων. Η διαδικασία αυτή επιτρέπει την ανάκτηση νομολογίας και συμβολαίων όχι μόνο με βάση λέξεις-κλειδιά, αλλά αναγνωρίζοντας την εννοιολογική συνάφεια, γεγονός ιδιαίτερα σημαντικό για το νομικό λεξιλόγιο, όπου πολλοί όροι αποδίδονται με διαφορετικές εκφράσεις. Η συγκεκριμένη αρχιτεκτονική εξηγεί και τα αποτελέσματα που είδαμε στο 4.5 σύμφωνα με τα ευρήματα του LeCaRD, η χρήση μοντέλων μετασχηματιστών (transformer-based models) με fine-tuning σε νομικά δεδομένα βελτίωσε το Recall@10 από 47% (βασική αντιστοίχιση λέξεων-κλειδιών) σε 67%, αναδεικνύοντας την πρακτική αξία των εξειδικευμένων LLMs στον νομικό τομέα [19]. Αυτό δείχνει ότι η μετάβαση από τις παραδοσιακές μεθόδους (π.χ. BM25, TF-IDF) προς legal-aware embeddings και advanced indexing οδηγεί όχι μόνο σε βελτίωση της ακρίβειας, αλλά και σε ουσιαστική ενίσχυση της χρηστικότητας για πραγματικά νομικά use cases.

Ωστόσο, η ανάγκη για ερμηνευσιμότητα και νομική τεκμηρίωση οδήγησε στην ανάπτυξη αρχιτεκτονικών τύπου ανάκτησης με ενισχυμένη δημιουργία (Retrieval-Augmented Generation, RAG). Στα συστήματα αυτά, το ερώτημα του χρήστη οδηγεί σε αναζήτηση σχετικών αποσπασμάτων μέσω ανάκτησης βάσει διανυσμάτων (vector-based retrieval), τα οποία στη συνέχεια παρέχονται ως συμπραζόμενα (context) σε ένα μεγάλο γλωσσικό μοντέλο (LLM), όπως το GPT-4 ή το LLaMA 2. Το γλωσσικό μοντέλο δεν απαντά απλώς γενικά, αλλά βασίζεται στα συγκεκριμένα ανακτημένα αποσπάσματα για να διαμορφώσει τεκμηριωμένες απαντήσεις, συχνά με ανίχνευση παραπομπών (citation tracing) προς τις πραγματικές πηγές, όπως εφαρμόζεται στο LegalBench-RAG [46]. Η διαδικασία αξιολογείται με μετρικές όπως η ακρίβεια παραπομπών (citation accuracy), το ποσοστό παραισθήσεων (hallucination rate), καθώς και μέσω επικύρωσης με συμμετοχή νομικών (human-in-the-loop validation), όπου οι ειδικοί ελέγχουν αν τα παραγόμενα αποτελέσματα στηρίζονται επαρκώς σε πραγματικές πηγές ή περιλαμβάνουν εφευρεμένες πληροφορίες. Πειράματα σε σύγχρονα RAG συστήματα (pipelines) έχουν καταδείξει ακρίβεια παραπομπών άνω του 85% και ποσοστά παραισθήσεων κάτω του 7% για κλασικά νομικά ερωτήματα, ενώ τα ποσοστά σφαλμάτων αυξάνονται σε ερωτήματα με ασάφεια ή πολλαπλές εν δυνάμει σωστές πηγές [46], [55].

Πέραν των «παραδοσιακών» συστημάτων σημασιολογικής αναζήτησης (semantic search pipelines) και RAG, η ανάλυση νομικών εγγράφων σε επίπεδο συμβολαίου απαιτεί εξειδικευμένες τεχνικές, όπως η εξαγωγή και σημασιολογική ερμηνεία ρητρών (clause extraction), η αναγνώριση οντοτήτων (named entity recognition) και ο εντοπισμός σχέσεων μεταξύ μερών, όρων και υποχρεώσεων. Σε τέτοια σενάρια, συστήματα που συνδυάζουν αναγνώριση οντοτήτων, κατακερματισμό ρητρών (clause segmentation) και ακολουθιακή επισήμανση (sequence labeling) μέσω μοντέλων όπως το Lawformer ή το RoBERTa-Large, αξιοποιούν σύνολα δεδομένων όπως το CUAD ή το LEDGAR, τα οποία περιλαμβάνουν εκατοντάδες χιλιάδες χειροκίνητα αντλημένες και σχολιασμένες ρήτρες [43], [53]. Τα αποτελέσματα των συστημάτων αυτών αξιολογούνται τόσο με κλασικές μετρικές (F1-score ανά τύπο ρήτρας) όσο και με σύγχρονες προσεγγίσεις ερμηνευσιμότητας (interpretability), όπως η οπτικοποίηση των βαρών προσοχής (attention weights visualization) για να τεκμηριώνεται ποιο σημείο του συμβολαίου συνέβαλε στην εξαγωγή, ή η παροχή βαθμολογίας νομικής βεβαιότητας (legal confidence scoring) για κάθε νομική πρόβλεψη.

Σε πιο προχωρημένες υλοποιήσεις, όπως το Krag ή το RainbowRAG, τα παραπάνω

συστήματα ενισχύονται με τη δημιουργία γνωσιακών γραφημάτων (knowledge graphs) που δομούν τις νομικές σχέσεις μεταξύ ρητρών και επιτρέπουν όχι μόνο τη σημασιολογική αναζήτηση, αλλά και τη λογική ανάλυση των αλληλεξαρτήσεων και των συμβατικών επιπτώσεων (RainbowRAG) [56]. Έτσι, ένα νομικό ερώτημα, π.χ. «ποιος φέρει την ευθύνη σε περίπτωση ανωτέρας βίας» δεν απαντάται μόνο με βάση την πλησιέστερη ρήτρα, αλλά και με εξερεύνηση των συνδεδεμένων οντοτήτων και των λογικών διαδρομών εντός του κνοωλεδγε γραφτη. Το τελικό αποτέλεσμα συνδυάζει τη σημασιολογική συνάφεια (embeddings) με τη λογική εγκυρότητα (ιεράρχηση βάσει γραφήματος, graph-based reranking), προσφέροντας τόσο ακρίβεια όσο και ερμηνευσιμότητα — στοιχεία κρίσιμα για τη συμμόρφωση με το δίκαιο (legal compliance) και τη δημιουργία audit trails.

Η πρακτική εφαρμογή όλων των παραπάνω απαιτεί εργαλεία ανοικτού κώδικα και εμπορικά πλαίσια όπως τα Transformers, FAISS, Haystack, LangChain, αλλά και συστήματα ενεργούς μάθησης (active learning) με ανατροφοδότηση από δικηγόρους (π.χ. ποσοστό διορθώσεων χρήστη, user correction rate, και κύκλους ενεργούς μάθησης, active learning loops). Οι πραγματικές υλοποιήσεις σε εταιρείες legal tech, όπως η Luminance και η Kira Systems, ενσωματώνουν οπτικοποίηση των αποτελεσμάτων εξαγωγής (visualization of extraction results), βαθμολόγηση βεβαιότητας (confidence scoring) και διαδραστική επικύρωση από νομικούς (interactive validation), επιτρέποντας συνεχή βελτίωση των μοντέλων και έλεγχο της ποιότητας σε πραγματικό χρόνο [55].

Τέλος, αξίζει να σημειωθεί ότι η αξιολόγηση της απόδοσης τέτοιων τεχνολογικών λύσεων δεν βασίζεται αποκλειστικά σε τεχνικές μετρικές, αλλά και στη λειτουργικότητα εντός πραγματικών νομικών ροών εργασίας (legal workflows), όπως το πόσα σφάλματα απαιτούν διόρθωση από τον τελικό χρήστη, η εξοικονόμηση χρόνου σε διαδικασίες ελέγχου (review procedures) και η συμβολή στην ανίχνευση κινδύνων συμμόρφωσης (compliance risk detection). Έρευνες σε υβριδικά συστήματα εξαγωγής ρητρών έδειξαν μείωση του χρόνου αναθεώρησης κατά 60% και αύξηση της ανάκλησης για κρίσιμους κινδύνους (recall in critical risks) πάνω από 90% όταν συνδυάζονται LLMs με ανατροφοδότηση ανθρώπινου παράγοντα (human-in-the-loop feedback) [43], [55].

Συνολικά, η τεχνική υλοποίηση των συστημάτων ανάκτησης και ανάλυσης νομικής πληροφορίας συνιστά ένα διαρκώς εξελισσόμενο επιστημονικό και μηχανικό πεδίο, με την έμφαση να μετατοπίζεται πλέον στη συνδυαστική χρήση σημασιολογικής ανάκτησης (semantic retrieval), επεξηγήσιμων συστημάτων RAG (explainable RAG), συλλογιστικής με γραφήματα (graph-based reasoning) και πρακτικών αξιολόγησης που ξεπερνούν τις απλές μετρήσεις ακρίβειας/ανάκλησης (precision/recall), ώστε να εξασφαλίζεται πραγματική αξιοπιστία και τεκμηριωμένη χρήση στη νομική πράξη.

4.2 Ανάλυση Συμβολαίων και Επεξεργασία Νομικών Εγγράφων

Η ανάλυση συμβολαίων και η αυτοματοποιημένη επεξεργασία νομικών εγγράφων αποτελούν κρίσιμες τεχνολογικές καινοτομίες στη σύγχρονη νομική επιστήμη, καθώς διευκολύνουν δραστηριότητες τον εντοπισμό, την αξιολόγηση και τη διαχείριση συμβατικών κινδύνων και υποχρεώσεων. Με την ταχεία πρόοδο των συστημάτων τεχνητής νοημοσύνης και ειδικότερα των Μεγάλων Γλωσσικών Μοντέλων (LLMs), αναδύονται νέα εργαλεία που αυτοματοποιούν

διαδικασίες όπως η κατηγοριοποίηση νομικών εγγράφων, η εξαγωγή κρίσιμων ρητρών, η ανίχνευση μη-συμμορφώσεων (non-compliance), καθώς και η αξιολόγηση της αποτελεσματικότητας των νομικών εγγράφων. Οι εξελίξεις αυτές όχι μόνο αυξάνουν την ταχύτητα και την ακρίβεια της νομικής εργασίας, αλλά επιτρέπουν και τη διείσδυση της τεχνητής νοημοσύνης σε τομείς όπως η συμβουλευτική, η νομική τεκμηρίωση και η πρόληψη ρίσκου.

Στο πλαίσιο αυτό, η τρέχουσα βιβλιογραφία και η βιομηχανική πρακτική εστιάζουν σε τρεις βασικούς άξονες:

(α) στην αυτοματοποιημένη αξιολόγηση και ερμηνεία συμβάσεων (automated contract review)

(β) στην κατηγοριοποίηση και εξαγωγή πληροφορίας από νομικά έγγραφα (document classification and information extraction),

(γ) στην αξιολόγηση της απόδοσης των συστημάτων αυτών (performance evaluation), τόσο σε ερευνητικό όσο και σε πρακτικό επίπεδο.

Στις υποενότητες που ακολουθούν, παρουσιάζονται οι κυριότερες μεθοδολογίες, τεχνικές αρχιτεκτονικές και σύγχρονα συστήματα που εφαρμόζονται στην ανάλυση συμβολαίων και την αυτοματοποιημένη νομική τεκμηρίωση, με αναφορά σε πρόσφατα πρότυπα αξιολόγησης (benchmarks), σύνολα δεδομένων (datasets) και μετρικές αξιολόγησης (evaluation metrics), βασιζόμενοι στην πλέον έγκυρη βιβλιογραφία του χώρου.

4.2.1 Αυτοματοποιημένη Αξιολόγηση Συμβάσεων

Η αυτοματοποιημένη αξιολόγηση συμβάσεων (automated contract review) αποτελεί έναν από τους πλέον αναπτυσσόμενους τομείς της τεχνολογικής καινοτομίας στη νομική πρακτική. Η παραδοσιακή διαδικασία ανάλυσης και ελέγχου συμβολαίων είναι ιδιαίτερα απαιτητική σε χρόνο και ανθρώπινο δυναμικό, καθώς απαιτεί σχολαστική επισκόπηση μεγάλων και πολύπλοκων κειμένων, εντοπισμό κρίσιμων ρητρών και αξιολόγηση των συμβατικών κινδύνων. Η χρήση προηγμένων συστημάτων τεχνητής νοημοσύνης, και ειδικότερα των Μεγάλων Γλωσσικών Μοντέλων, έχει τη δυνατότητα να αυτοματοποιήσει σε μεγάλο βαθμό τη διαδικασία αυτή, βελτιώνοντας σημαντικά τόσο την ακρίβεια όσο και την αποδοτικότητα της αξιολόγησης συμβολαίων.

Η βασική μέθοδος αυτοματοποιημένης αξιολόγησης συμβολαίων στηρίζεται στη χρήση αλγορίθμων επεξεργασίας φυσικής γλώσσας (Natural Language Processing, NLP) και ειδικότερα βαθιών νευρωνικών μοντέλων (deep learning models), όπως τα μοντέλα μετασχηματιστών μεγάλων γλωσσικών μοντέλων (transformer-based LLMs) — π.χ. LegalBERT, Lawformer, GPT-4 [35]. Αυτά τα μοντέλα εκπαιδεύονται ή προσαρμόζονται (fine-tuned) σε εξειδικευμένα σώματα κειμένων που αποτελούνται αποκλειστικά από συμβατικές ρήτρες και νομική ορολογία, με στόχο την ακριβή κατανόηση και ερμηνεία των περιεχομένων τους [54].

Τα συμβόλαια αρχικά διαχωρίζονται σε θεματικές ρήτρες (clauses segmentation), όπως «ρήτρα καταγγελίας», «ρήτρα ευθύνης», «ρήτρα ανωτέρας βίας» κ.ά., και στη συνέχεια αναλύονται για την εξαγωγή συγκεκριμένων οντοτήτων και πληροφοριών μέσω αναγνώρισης οντοτήτων (Named Entity Recognition, NER) [35], [4], [54]. Η χρήση μοντέλων βασισμένων σε μετασχηματιστές (transformer-based models) όπως το LegalBERT ή το Lawformer επιτρέπει την καλύτερη κατανόηση της ιδιαίτερης νομικής φρασεολογίας, καθώς αυτά τα μο-

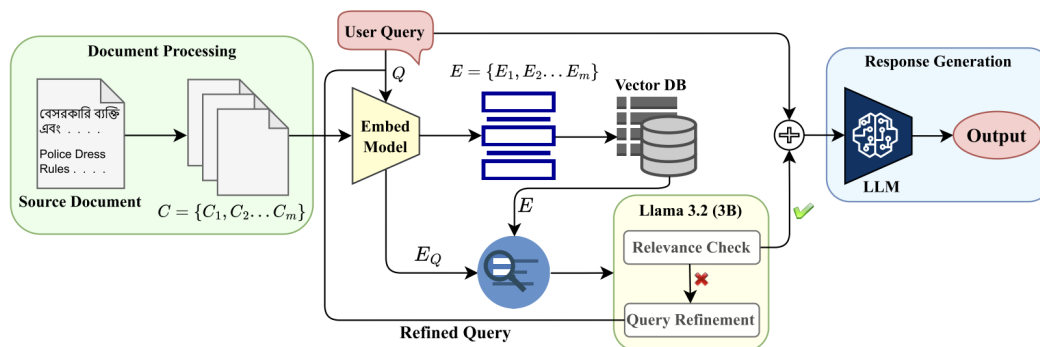
ντέλα έχουν εκπαιδευτεί ειδικά σε εκτενή σώματα νομικών εγγράφων, επιτυγχάνοντας υψηλή σημασιολογική ακρίβεια [35].

Η αρχιτεκτονική Ανάκτησης-Ενισχυμένης Παραγωγής (Retrieval-Augmented Generation – RAG) έχει πρόσφατα αναδειχθεί ως μία από τις πιο υποσχόμενες προσεγγίσεις για την αυτοματοποιημένη ανάλυση συμβολαίων. Τα συστήματα RAG αξιοποιούν τη σημασιολογική αναζήτηση (semantic search) και την παραγωγική ικανότητα των μεγάλων γλωσσικών μοντέλων (LLMs) για την παροχή ακριβών και τεκμηριωμένων απαντήσεων σε ερωτήματα σχετικά με συμβατικά έγγραφα.

Ένα χαρακτηριστικό παράδειγμα χρήσης της αρχιτεκτονικής RAG είναι η εφαρμογή Contract Law using RAG, η οποία επιτρέπει τη δυναμική ανάλυση συμβατικών ρητρών, εντοπίζοντας πιθανούς συμβατικούς κινδύνους ή σημεία μη συμμόρφωσης μέσω διαδραστικών ερωτημάτων [4]. Το pipeline της αρχιτεκτονικής αυτής περιλαμβάνει:

- Μετατροπή των συμβολαίων σε διανυσματικές αναπαραστάσεις (vector embeddings) μέσω εξειδικευμένων μοντέλων (π.χ. LegalBERT).
- Αποθήκευση των ενσωματώσεων (embeddings) σε εξειδικευμένες βάσεις δεδομένων (FAISS, Weaviate) για ταχεία ανάκτηση.
- Χρήση μεγάλου γλωσσικού μοντέλου (LLM) για παραγωγή απαντήσεων που τεκμηριώνονται από τις ανακτημένες πληροφορίες.

Η συνολική διαδικασία απεικονίζεται στο Σχήμα 4.4, όπου παρουσιάζεται το πλήρες pipeline μιας εφαρμογής RAG για ανάλυση νομικών εγγράφων. Το pipeline αυτό περιλαμβάνει μετατροπή των πηγών σε διανυσματικές αναπαραστάσεις, αποθήκευση και ανάκτηση σχετικών αποσπασμάτων από βάση δεδομένων διανυσμάτων, και τελική παραγωγή απάντησης μέσω LLM. Ιδιαίτερα σημαντικό είναι ότι η αρχιτεκτονική ενσωματώνει στάδια ελέγχου συνάφειας (relevance check) και βελτιστοποίησης ερωτήματος (query refinement), διασφαλίζοντας ότι οι ανακτημένες πληροφορίες είναι ουσιαστικά σχετικές με το ερώτημα του χρήστη, γεγονός που μειώνει τα ποσοστά παραίσθησης και αυξάνει την ακρίβεια των απαντήσεων [4].



Σχήμα 4.4: Αρχιτεκτονική συστήματος RAG για ανάλυση νομικών εγγράφων και συμβολαίων. Το pipeline ενσωματώνει διαδικασίες ελέγχου συνάφειας (Relevance Check) και βελτιστοποίησης ερωτήματος (Query Refinement), ενισχύοντας την ακρίβεια και τη διαφάνεια των απαντήσεων [4].

Η αξιολόγηση της αποτελεσματικότητας των συστημάτων ανάλυσης συμβολαίων πραγματοποιείται συνήθως με τη χρήση εξειδικευμένων datasets και benchmarking frameworks. Τα πλέον αναγνωρισμένα σύνολα δεδομένων περιλαμβάνουν το CUAD (Contract Understanding Atticus Dataset), το LexGLUE, και το LEDGAR, τα οποία παρέχουν annotated δεδομένα για tasks όπως η αναγνώριση συγκεκριμένων ρητρών, η ανίχνευση συμβατικών υποχρεώσεων και ο εντοπισμός νομικών κινδύνων [43], [3].

Συγκεκριμένα, το dataset CUAD, το οποίο αποτελείται από συμβόλαια με annotated κρίσιμες ρήτρες, χρησιμοποιείται εκτενώς ως πρότυπο για την αξιολόγηση της απόδοσης των συστημάτων extraction και classification συμβατικών ρητρών. Μετρικές όπως η ακρίβεια (precision), η ανάκληση (recall), το F1-score, αλλά και η ακρίβεια παραπομπών (citation accuracy) χρησιμοποιούνται για τη συγκριτική αξιολόγηση της απόδοσης των μοντέλων αυτών [3], [55].

Η πρακτική εφαρμογή των μεθόδων αυτοματοποιημένης αξιολόγησης συμβολαίων έχει ήδη αρχίσει να διεισδύει σε βιομηχανικούς κλάδους μέσω εργαλείων όπως το Kira Systems, το Luminance, και το Ironclad που αξιοποιούν τεχνολογίες βαθιάς μηχανικής μάθησης για την αυτοματοποίηση διαδικασιών due diligence, compliance και διαχείρισης συμβατικών κινδύνων [55].

Ένα χαρακτηριστικό παράδειγμα αποτελεί η χρήση του Luminance από μεγάλες δικηγορικές εταιρείες, όπου οι αυτοματοποιημένες αναλύσεις συμβολαίων επιτρέπουν τη γρήγορη επισκόπηση εκατοντάδων σελίδων συμβατικού κειμένου σε ελάχιστο χρόνο, με ακρίβεια που πλησιάζει ή υπερβαίνει τις επιδόσεις έμπειρων δικηγόρων. Συγκριτικές μελέτες έχουν δείξει ότι τέτοια συστήματα μπορούν να μειώσουν τον απαιτούμενο χρόνο αναθεώρησης συμβολαίων έως και 70% και να αυξήσουν την ακρίβεια αναγνώρισης κρίσιμων συμβατικών όρων σε σχέση με μανυαλ αναλύσεις [54].

Η αυτοματοποιημένη αξιολόγηση συμβολαίων μέσω προηγμένων NLP και LLM αρχιτεκτονικών έχει αποδειχθεί αποτελεσματική σε τεχνικό επίπεδο. Ωστόσο, παραμένουν προκλήσεις όπως η εξασφάλιση απόλυτης διαφάνειας, η δυνατότητα αιτιολόγησης των παραγόμενων αποτελεσμάτων (explainability), και η ανάγκη για συνεχή ενημέρωση και προσαρμογή των μοντέλων στις εξελίξεις του νομικού λόγου.

Η περαιτέρω ανάπτυξη των τεχνολογιών αυτών απαιτεί τη στενή συνεργασία τεχνικών και νομικών εμπειρογνομόνων για την προσαρμογή των εργαλείων στις ειδικές απαιτήσεις κάθε έννομης τάξης και την καθιέρωση τυποποιημένων και διαφανών διαδικασιών αξιολόγησης της απόδοσης και των αποτελεσμάτων.

Η κατηγοριοποίηση νομικών εγγράφων (document classification) και η εξαγωγή πληροφορίας (information extraction) αποτελούν θεμελιώδεις λειτουργίες στη σύγχρονη νομική πληροφορική. Οι τεχνολογίες αυτές επιτρέπουν τον αυτόματο χαρακτηρισμό μεγάλου όγκου νομικών κειμένων, όπως αποφάσεις, συμβόλαια, διατάξεις, καθώς και την άντληση κρίσιμων πληροφοριών και νομικών οντοτήτων, διευκολύνοντας ουσιαστικά την ανάλυση, την αρχειοθέτηση και τη λήψη τεκμηριωμένων αποφάσεων.

Στο πεδίο της κατηγοριοποίησης εγγράφων, οι πρώτες προσεγγίσεις βασίζονταν σε παραδοσιακές μεθόδους μηχανικής μάθησης, όπως Support Vector Machines (SVMs), Random Forests και bag-of-words, που όμως αδυνατούν να αποδώσουν τη σημασιολογική πολυπλοκότητα και τα ιδιαίτερα χαρακτηριστικά της νομικής γλώσσας [3], [4]. Με την εισαγωγή των

Μεγάλων Γλωσσικών Μοντέλων και των τεχνικών μετασχηματιστών (transformers), η ακρίβεια και η σημασιολογική πληρότητα των αποτελεσμάτων αυξήθηκε δραστικά.

Εξειδικευμένα μοντέλα, όπως το LegalBERT, το Lawformer και το LegalSearchLM, εκπαιδεύονται σε μεγάλες συλλογές νομικών κειμένων και καταφέρνουν να συλλάβουν τη θεματική, τη δομή και την ορολογία κάθε εγγράφου. Σε πρακτικό επίπεδο, το πρόβλημα κατηγοριοποίησης μπορεί να υλοποιηθεί είτε ως απλή ταξινόμηση (π.χ. συμβόλαιο, απόφαση, γνωμοδότηση) είτε ως πολυκατηγορική ταξινόμηση, όπου τα έγγραφα μπορούν να ανήκουν ταυτόχρονα σε πολλαπλές νομικές κατηγορίες (multi-label classification) [4], [55].

Το LexGLUE benchmark συγκεντρώνει διάφορα tasks κατηγοριοποίησης για διαφορετικά είδη νομικών εγγράφων, επιτρέποντας τη συγκριτική αξιολόγηση διαφορετικών μοντέλων σε πρακτικές συνθήκες [3]. Για παράδειγμα, σε tasks πρόβλεψης εφετειακής απόφασης, το LegalBERT πέτυχε σημαντικά υψηλότερη ακρίβεια από παραδοσιακά baseline μοντέλα, με F1-scores που αγγίζουν το 82% [4].

Η εξαγωγή πληροφορίας περιλαμβάνει ένα φάσμα τεχνικών: από την αναγνώριση οντοτήτων (Named Entity Recognition – NER), όπως πρόσωπα, νομικές οντότητες, ημερομηνίες και τοποθεσίες, μέχρι την εξαγωγή συγκεκριμένων ρητρών από συμβόλαια (contract clause extraction) ή τη δομική ανάλυση αποφάσεων (decision structure extraction) [35], [4], [55], [54].

Τα σύγχρονα συστήματα επεξεργασίας (pipelines) βασίζονται σε μοντέλα μετασχηματιστών (transformers) που υλοποιούν ακολουθιακή επισήμανση (sequence labeling), ταξινόμηση τοκενς (token classification) και εξαγωγή σχέσεων (relation extraction), με ιδιαίτερη έμφαση στη σημασιολογική ορθότητα των εξαγόμενων στοιχείων. Πλαίσια όπως το CUAD και το LEDGAR έχουν καθιερωθεί ως πρότυπα για την αξιολόγηση συστημάτων εξαγωγής ρητρών (clause extraction). Το CUAD, για παράδειγμα, περιλαμβάνει συμβόλαια με πάνω από 13.000 νομικές ρήτρες, επισημασμένες από έμπειρους νομικούς, ενώ το LEDGAR προσφέρει εκτενές σώμα δεδομένων με επισημασμένες συμβατικές διατάξεις (annotated contractual provisions) [43], [3].

Η αξιολόγηση των συστημάτων εξαγωγής πληροφορίας γίνεται με χρήση μετρικών όπως η ακρίβεια (precision), η ανάκληση (recall) και το F1-score, με στόχο τη μέγιστη εξισορρόπηση μεταξύ ψευδώς θετικών και ψευδώς αρνητικών εξαγωγών. Σύγχρονες προσεγγίσεις αξιοποιούν επίσης μηχανισμούς προσοχής (attention mechanisms) για την ερμηνευσιμότητα (explainability) των αποτελεσμάτων και τη δυνατότητα ανίχνευσης παραπομπών (citation tracing), ειδικά σε περιβάλλοντα RAG [4], [46].

Ένα ιδιαίτερο πρόβλημα στην εξαγωγή πληροφορίας από νομικά έγγραφα είναι η πολυγλωσσικότητα και η εφαρμογή σε δικαιοτικά συστήματα χαμηλών πόρων (low-resource legal systems). Τα περισσότερα εξελιγμένα μοντέλα (state-of-the-art models) εκπαιδεύονται κυρίως σε αγγλικά ή κινεζικά σώματα κειμένων (corpora), περιορίζοντας την απόδοσή τους σε χώρες όπως η Ελλάδα ή η Πορτογαλία [4]. Οι σύγχρονες ερευνητικές προσπάθειες εστιάζουν στην προσαρμογή πολυγλωσσικών μετασχηματιστών (fine-tuning multilingual transformers, π.χ. mBERT, XLM-R), στην ενίσχυση δεδομένων μέσω μεταφράσεων (data augmentation with translations), καθώς και στον συνδυασμό παραδοσιακών και μεθόδων βαθιάς μάθησης (deep learning methods) για βέλτιστη κάλυψη των νομικών εννοιών και στις λιγότερο διαδεδομένες γλώσσες [4], [55], [51].

Η τεχνολογία κατηγοριοποίησης και εξαγωγής πληροφορίας έχει βρει πρακτική εφαρμογή σε σύγχρονα εργαλεία νομικής τεχνολογίας (legal tech), όπως πλατφόρμες διαχείρισης εγγράφων (document management), συμμόρφωσης (compliance), νομικού ελέγχου (due diligence) και αυτόματης σύνταξης ή αναθεώρησης συμβολαίων (automated contract drafting/review). Παραδείγματα περιλαμβάνουν πιπελίνες που ενσωματώνουν το LegalBERT ή Lawformer για το tagging και την ταξινόμηση (classification), σε συνδυασμό με υπομονάδες εξαγωγής ρητρών (clause extraction modules) που χρησιμοποιούν επισημασμένα δεδομένα (annotated data) από τα CUAD ή LEDGAR.

Σε πραγματικές συνθήκες, η χρήση τέτοιων συστημάτων επιτρέπει σε μεγάλες εταιρείες και οργανισμούς να επεξεργάζονται αυτόματα χιλιάδες έγγραφα, να εντοπίζουν κρίσιμες ρήτρες, να παρακολουθούν ρομπλιανζε και να αυτοματοποιούν νομικές ροές εργασίας με ακρίβεια που, σύμφωνα με συγκριτικές μελέτες, πλησιάζει ή και υπερβαίνει το ανθρώπινο επίπεδο [55], [54].

Παρά τα εντυπωσιακά αποτελέσματα, παραμένουν ανοιχτά ζητήματα που αφορούν την ερμηνευσιμότητα των αποτελεσμάτων, την ανάγκη για συνεχή επέκταση των αννοτατεδ δατα-σετς, την προσαρμογή των μοντέλων στις εξελίξεις της νομολογίας και τη διασφάλιση της ακρίβειας σε πολυγλωσσικά περιβάλλοντα. Η διασύνδεση των αποτελεσμάτων εξαγωγής (extraction results) με εφαρμογές επόμενων σταδίων (downstream applications), όπως η ανασκόπηση εγγράφων (document review) ή τα συστήματα νομικού συλλογισμού (legal reasoning pipelines), αποτελεί αντικείμενο ενεργού έρευνας.

4.2.2 Αξιολόγηση Απόδοσης

Η συστηματική αξιολόγηση των συστημάτων αυτόματης ανάλυσης συμβολαίων και εξαγωγής πληροφορίας έχει κεντρική σημασία για την πρακτική και επιστημονική αξιοπιστία τους. Εδώ, η αξιολόγηση δεν εξαντλείται στα συνηθισμένα αριθμητικά scores, αλλά περιλαμβάνει λεπτομερή ανάλυση της συμπεριφοράς κάθε συστήματος επεξεργασίας (pipeline) σε διαφορετικού τύπου δεδομένα, συνθήκες και νομικές χρήσεις.

1. Αξιολόγηση Εξαγωγής Ρητρών Συμβολαίων (Clause Extraction)

Στα έργα (tasks) εξαγωγής ρητρών, όπως αυτά που αξιολογούνται με το CUAD, η μέτρηση γίνεται με:

- Ακρίβεια (precision): το ποσοστό των αποσπασμάτων που εξήγαγε το σύστημα και ήταν πράγματι σωστά.
- Ανάκληση (recall): το ποσοστό των συνολικών ρητρών που εντοπίστηκαν.
- F1-score: ο αρμονικός μέσος των παραπάνω.

Στα επίσημα benchmarks, μοντέλα όπως το RoBERTa-Large (με περαιτέρω εκπαίδευση στο CUAD) έφτασαν F1-score 91% σε κλασικές ρήτρες (π.χ. εχεμύθεια), αλλά μειώθηκε στο 62–75% σε σπάνιες ή ιδιαίτερα εξαρτώμενες από το συμφραζόμενο ρήτρες (context-dependent, π.χ. μη-τυποποιημένη ευθύνη) [43]. Η αξιολόγηση γίνεται ανά κατηγορία (per-class), καθώς κάθε ρήτρα διαφέρει σε συχνότητα και δυσκολία.

Η ανάλυση σφαλμάτων (error analysis) αποκαλύπτει ότι τα περισσότερα ψευδώς αρνητικά (false negatives) αφορούν «ένθετες» ρήτρες (nested clauses) ή περιπτώσεις όπου η νομική

σημασία εξαρτάται από την προηγούμενη παράγραφο — κάτι που ακόμα και άνθρωποι συχνά παραβλέπουν.

Για την πληρέστερη αποτίμηση της απόδοσης των μοντέλων σε νομικά δεδομένα, χρησιμοποιούνται εξειδικευμένα σχήματα αξιολόγησης που βασίζονται στις ακόλουθες τρεις κατηγορίες:

1. Σύγκριση Ανθρώπου–Μοντέλων (Human vs LLMs)

Σε μελέτη της ίδιας πλατφόρμας, νεότεροι νομικοί πέτυχαν ανάκληση 85–95% σε συχνές ρήτρες αλλά 60–70% σε σπάνιες, δηλαδή τα καλά βελτιστοποιημένα (fine-tuned) LLMs αρχίζουν να πλησιάζουν (ή και να ξεπερνούν) την ανθρώπινη απόδοση σε συγκεκριμένα έργα [43], [54].

2. Αξιολόγηση Κατηγοριοποίησης και Επισήμανσης (Classification & Tagging)

Η ταξινόμηση εγγράφων σε νομικές κατηγορίες (LexGLUE) γίνεται με τις κλασικές μετρικές ταξινόμησης (accuracy, macro/micro F1).

Για παράδειγμα, στο έργο πρόβλεψης αποφάσεων της ΕΕ (EU judgment prediction), το LegalBERT πέτυχε ακρίβεια (accuracy) 79% και macro-F1 77%, όταν το βασικό μοντέλο (baseline) ήταν 62% [4]. Σε έργα επισήμανσης (tagging, multi-label), χρησιμοποιούνται hamming loss, exact match ratio και ακρίβεια ανά ετικέτα (per-label precision), αφού ένα έγγραφο μπορεί να ανήκει ταυτόχρονα σε πολλαπλές νομικές κατηγορίες.

3. Αξιολόγηση Εξαγωγής Πληροφορίας (Information Extraction: NER, Relation Extraction)

Εδώ η μέτρηση περιλαμβάνει:

- Ακρίβεια/ανάκληση/F1 σε επίπεδο οντότητας (entity-level precision/recall/F1) π.χ. για εμπλεκόμενα μέρη, ημερομηνίες, υποχρεώσεις.
- Αξιολόγηση κατά τμήμα (span-based) και κατά αλληλεπικάλυψη (overlap-based evaluation): αν το μοντέλο αναγνωρίζει ολόκληρη ή μέρος της οντότητας.

Στο LEDGAR, το καλύτερο μοντέλο έφτασε entity F1 πάνω από 89% για νομικές οντότητες, αλλά κάτω από 70% για σύνθετες υποχρεώσεις (ένθετες ή διαδοχικές) [53]. Η ανάλυση σφαλμάτων κατά τμήμα (span-based error analysis) δείχνει ότι τα περισσότερα λάθη εμφανίζονται σε μακροσκελείς, μεικτές ρήτρες (δηλαδή με πολλαπλές πληροφορίες μαζί).

4. Πολυγλωσσική και Low-Resource Αξιολόγηση (Multilingual & Low-Resource Evaluation)

Τα περισσότερα benchmarks έχουν γίνει σε αγγλικά και κινεζικά. Σε γλώσσες όπως τα πορτογαλικά (BERTimbau) ή τα ελληνικά (μεταφορά από mBERT/XLM-R), το F1-score πέφτει έως και 15 μονάδες σε εξαγωγή ρητρών ή οντοτήτων [4], [51]. Η κύρια

πρόκληση εδώ είναι η έλλειψη επισημασμένων δεδομένων (annotated data), οι γλωσσικές ιδιομορφίες (π.χ. συντακτικό, σειρά λέξεων) και τα λάθη μετάφρασης κατά το fine-tuning.

5. Αξιολόγηση RAG: Παραπομπές, Παισιθήσεις, Downstream

Στα συστήματα (RAG), η ακρίβεια παραπομπών (citation accuracy), δηλαδή πόσο καλά το LLM στηρίζει τις παραγόμενες απαντήσεις με πραγματικές πηγές— μετριέται με:

- Κάλυψη παραπομπών (citation coverage): ποσοστό απαντήσεων με τουλάχιστον μία σωστή παραπομπή.
- Ακρίβεια/ανάκληση παραπομπών (citation precision/recall): κατά πόσο τα παρατιθέμενα τεκμήρια είναι σχετικά και αληθή.
- Ποσοστό ανακρίβειών (hallucination rate): ποσοστό απαντήσεων που περιέχουν ανύπαρκτες ή εφευρεμένες πληροφορίες.

Σε benchmarking του LegalBench-RAG, τα συστήματα LegalBERT + retriever πέτυχαν ακρίβεια παραπομπών πάνω από 85%, με ποσοστά παισιθήσεων κάτω από 7% — σαφώς χαμηλότερα από καθαρά LLMs χωρίς retrieval [4], [46]. Ωστόσο, όταν το σύνολο δεδομένων είχε πολλαπλές κοντινές πηγές, τα μοντέλα συχνά μπέρδευαν ή συνέθεταν πληροφορίες (composite hallucinations).

6. Σχολιασμός Συνόλων Δεδομένων & Εμπόδια Αξιολόγησης (Dataset Annotation & Evaluation Bottlenecks)

Η ποιότητα της αξιολόγησης εξαρτάται άμεσα από την επισημείωση των συνόλων δεδομένων (annotation of datasets).

- CUAD: επισημασμένο από νομικούς εμπειρογνώμονες, απαιτείται τριπλή συμφωνία (3-way agreement), με δια-επισημαντική συμφωνία F1 93%.
- ContractNLI: συλλογιστική σε επίπεδο ρήτρας (clause-level inference), όπου οι annotators διαφωνούν σε «γκρίζες ζώνες» (π.χ. αν υπάρχει αντίφαση ή όχι).

Το στενό σημείο της επισημείωσης (annotation bottleneck) καθυστερεί τη σύγκριση επιδόσεων, ειδικά για νέες δικαιοδοσίες ή πολυγλωσσικά έργα. Πολλές εταιρείες/οργανισμοί πλέον ενσωματώνουν ενεργή μάθηση (active learning) ή ημιαυτόματη προκαταρκτική επισήμανση (semi-automated pre-tagging) για επιτάχυνση της διαδικασίας.

7. Αξιολόγηση σε Πραγματικό Περιβάλλον & Ανατροφοδότηση από τον Χρήστη (Real-World Deployment & Human-in-the-loop Feedback)

Τα πιο επιτυχημένα συστήματα αξιολογούνται και «στο πεδίο», δηλαδή σε ζωντανές ροές εργασίας (live workflows):

- Έλεγχος εγγράφων (document review): μετρώνται τα σφάλματα που έφτασαν στον τελικό χρήστη και απαιτούσαν διόρθωση (user correction rate).

- Συμμόρφωση (compliance): πόσοι κίνδυνοι συμμόρφωσης εντοπίστηκαν αυτόματα και πόσοι αγνοήθηκαν/παραλείφθηκαν.
- Διάρκεια επεξεργασίας: μέτρηση της εξοικονόμησης χρόνου έναντι της χειροκίνητης αξιολόγησης.

Σύμφωνα με μελέτες, εταιρείες που υιοθέτησαν πλήρως αυτοματοποιημένα συστήματα εξαγωγής ρητρών με συμμετοχή ανθρώπου στη διαδικασία (human-in-the-loop) ανέφεραν μείωση του χρόνου ελέγχου έως 60% και βελτίωση της ανάκλησης για κρίσιμους κινδύνους (critical risks) πάνω από 90% [55], [54].

8. Ερμηνευσιμότητα – Επεξηγησιμότητα και Νομική Βεβαιότητα (Interpretability – Explainability and Legal Confidence)

Τα τελευταία χρόνια, αυξάνεται η απαίτηση για επεξηγήσιμα συστήματα τεχνητής νοημοσύνης (explainable AI):

- Οπτικοποίηση προσοχής (attention visualization): τα μοντέλα εμφανίζουν το «σημείο» του συμβολαίου ή της απόφασης που αποτέλεσε τη βάση για την πρόβλεψη.
- Βαθμολόγηση νομικής βεβαιότητας (legal confidence scoring): όχι μόνο δυαδικό αποτέλεσμα, αλλά και βαθμολογία/εξήγηση για το «πόσο σίγουρο» είναι το μοντέλο για κάθε πρόβλεψη/εξαγωγή.
- Διαδραστική επικύρωση (interactive validation): ο νομικός χρήστης μπορεί να αποδεχθεί ή να απορρίψει κάθε απάντηση με ένα κλικ, ενισχύοντας τον κύκλο ανατροφοδότησης μέσω ενεργής μάθησης (active learning feedback loop).

9. Προκλήσεις & Μελλοντικές Τάσεις

- Ανθεκτικότητα (robustness): τα μοντέλα πρέπει να αντεπεξέρχονται σε «εκτός πεδίου» (out of domain) συμβόλαια/αποφάσεις (π.χ. νέοι νόμοι, καινοφανείς ρήτρες).
- Μεροληψία και Δικαιοσύνη (bias & fairness): αποφυγή εύκολης παρανόησης λόγω προκαταλήψεων του συνόλου εκπαίδευσης (training set biases).
- Τυποποίηση μετρικών (standardization of metrics): απουσία ενιαίων πρωτοκόλλων αξιολόγησης μεταξύ πλατφορμών/συνόλων δεδομένων.
- Ενσωμάτωση στη νομική πρακτική (integration with legal practice): τα εξελιγμένα συστήματα εστιάζουν πλέον στην ενσωμάτωση με υφιστάμενα συστήματα διαχείρισης νομικών εγγράφων (legal document management systems, DMS), όχι μόνο σε εργαστηριακές συνθήκες (lab conditions).

Η πρακτική εφαρμογή μετρικών αποτυπώνεται στο παρακάτω διάγραμμα, το οποίο παρουσιάζει τα αποτελέσματα αξιολόγησης του συνόλου δεδομένων LeCaRD. Στο διάγραμμα αυτό παρουσιάζονται οι επιδόσεις διαφόρων μεθόδων ανάλυσης ομοιότητας νομικών υποθέσεων, όπου διαπιστώνεται ότι τα νομικά ευαισθητοποιημένα μοντέλα (legal-aware models), όπως το Legal-BERT, υπερέρχουν σημαντικά σε μετρικές όπως το MAP και το nDCG

σε σύγκριση με παραδοσιακές προσεγγίσεις (BM25, TF-IDF). Τα αποτελέσματα αυτά επιβεβαιώνουν τη σημασία των εξειδικευμένων embeddings και της βαθύτερης σημασιολογικής κατανόησης.

Όπως φαίνεται στο διάγραμμα 4.5, η εξέλιξη των legal-aware μοντέλων οδηγεί σε βελτίωση όχι μόνο της ακρίβειας, αλλά και της πρακτικής χρησιμότητας των συστημάτων, καθιστώντας τα πιο αποτελεσματικά εργαλεία για τη νομική έρευνα και την υποστήριξη της ανάλυσης ομοιότητας. Ωστόσο, η έννοια της συνάφειας παραμένει σύνθετη και συχνά εξαρτάται από το νομικό και πρακτικό πλαίσιο της αξιολόγησης.

	Model	P@5	P@10	MAP	NDCG@10	NDCG@20	NDCG@30
Common query set	BM25	0.423	0.410	0.490	0.726	0.790	0.883
	TF-IDF	0.348	0.305	0.480	0.789	0.830	0.847
	LMIR	0.460	0.430	0.511	0.766	0.813	0.896
Controversial query set	BM25	0.348	0.283	0.463	0.745	0.821	0.903
	TF-IDF	0.157	0.113	0.379	0.812	0.841	0.852
	LMIR	0.357	0.326	0.443	0.779	0.837	0.911
Overall query set	BM25	0.406	0.381	0.484	0.731	0.797	0.888
	TF-IDF	0.304	0.261	0.457	0.795	0.832	0.848
	LMIR	0.436	0.406	0.495	0.769	0.818	0.900
Test set	BM25	0.380	0.350	0.498	0.739	0.804	0.894
	TF-IDF	0.270	0.215	0.459	0.817	0.836	0.853
	LMIR	0.450	0.435	0.512	0.769	0.807	0.896
	BERT	0.470	0.430	0.568	0.774	0.821	0.899

Σχήμα 4.5: Αποτελέσματα αξιολόγησης διαφορετικών μεθόδων ανάλυσης ομοιότητας νομικών υποθέσεων στο σύνολο δεδομένων LeCaRD, με τις βασικές μετρικές P@5, P@10, MAP και nDCG.[3].

4.3 Πρόβλεψη Δικαστικών Αποφάσεων

Η πρόβλεψη δικαστικών αποφάσεων μέσω τεχνικών μηχανικής μάθησης και Μεγάλων Γλωσσικών Μοντέλων (LLMs) αποτελεί ένα από τα πλέον καινοτόμα της νομικής πληροφορικής. Η δυνατότητα των υπολογιστικών συστημάτων να εκτιμούν την έκβαση μιας υπόθεσης, αξιοποιώντας ιστορικά δεδομένα, νομικά επιχειρήματα και πρότυπα συλλογιστικής, αναδεικνύει τόσο τις τεράστιες δυνατότητες όσο και τους περιορισμούς της τεχνητής νοημοσύνης στον χώρο της απονομής δικαιοσύνης. Παράλληλα, η αξιοποίηση τέτοιων εργαλείων εγείρει ζητήματα που σχετίζονται με τα όρια της αυτοματοποίησης, τη σημασία της ανθρώπινης ερμηνείας και την αξιοπιστία των προβλεπτικών μοντέλων στη νομική πράξη.

Τα τελευταία χρόνια, πλήθος μελετών έχουν προτείνει μοντέλα για την πρόβλεψη αποφάσεων σε δικαστήρια όπως το Ευρωπαϊκό Δικαστήριο Δικαιωμάτων του Ανθρώπου (ECHR), το Ανώτατο Δικαστήριο των ΗΠΑ, το Ανώτατο Λαϊκό Δικαστήριο της Κίνας και εθνικά δικαστήρια στην Τουρκία[57], Βραζιλία[58], Καναδά[59] και Φιλιππίνες[12]. Οι προσεγγίσεις κυμαίνονται από κλασικά υποστηρικτικά διανυσματικά μηχανήματα (SVMs) και recurrent networks (LSTM/GRU), έως μετασχηματιστές όπως το BERT και τα εξειδικευμένα δίκτυα όπως το Lawformer, το QAJudge και το LawLLM.

Η πρόβλεψη μπορεί να αφορά διαφορετικές πτυχές: το αν παραβιάστηκε ένα άρθρο, το είδος της ποινής, την ετυμηγορία ή ακόμη και τη ρητορική δομή της απόφασης. Εξίσου σημαντικό είναι και το ζήτημα της εξήγησης της πρόβλεψης όχι μόνο τι προβλέπει το μοντέλο, αλλά και γιατί. Γι' αυτό η ανάλυση της πρόβλεψης δεν είναι μόνο αριθμητική αλλά και ποιοτική, με έμφαση στην αναγνωσιμότητα, τη διαφάνεια και τη νομική τεκμηρίωση της

παραγόμενης απόφασης.

Στις ενότητες που ακολουθούν, παρουσιάζονται οι μεθοδολογικές βάσεις, οι τεχνικές εφαρμογές και οι τρόποι αξιολόγησης των συστημάτων πρόβλεψης δικαστικών αποφάσεων, καθώς και οι προκλήσεις και περιορισμοί που προκύπτουν από την εφαρμογή τους στην πράξη.

4.3.1 Μεθοδολογία Προγνωστικών Μοντέλων

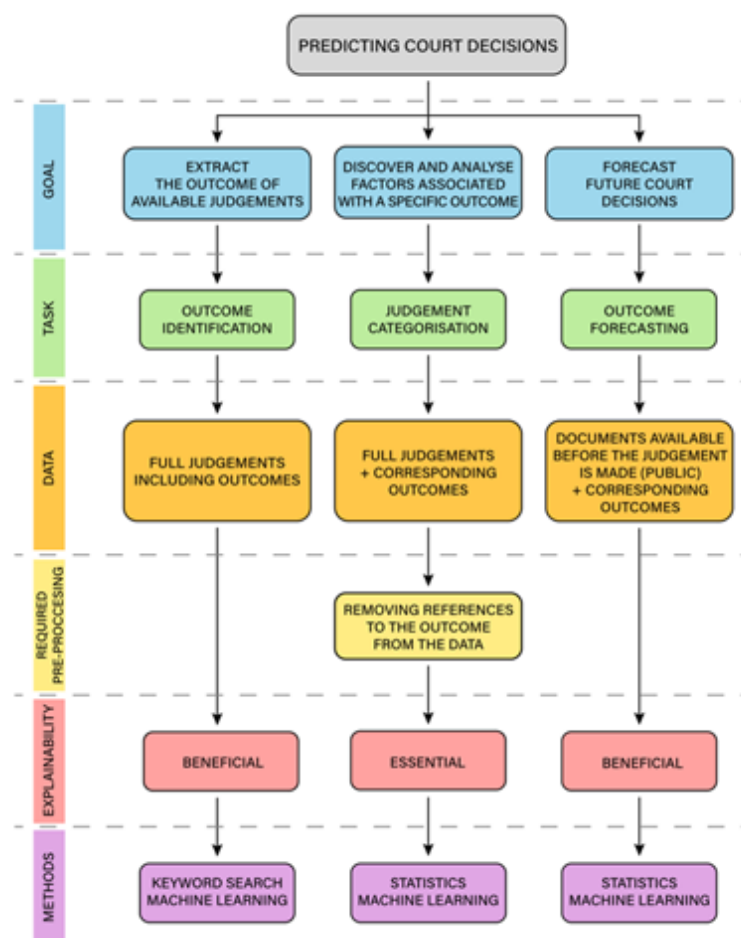
Η ανάπτυξη αξιόπιστων προγνωστικών μοντέλων για δικαστικές αποφάσεις προϋποθέτει τη χρήση εξειδικευμένων και υψηλής ποιότητας νομικών δεδομένων, καθώς και την εφαρμογή προηγμένων μεθόδων προεπεξεργασίας, ώστε να εξαχθεί η μέγιστη δυνατή πληροφορία από τα αδόμητα και συχνά ανομοιογενή νομικά κείμενα.

Η δομή των δεδομένων σε σύγχρονα σύνολα δεδομένων συνήθως περιλαμβάνει το πλήρες κείμενο της απόφασης, την αναλυτική περιγραφή των πραγματικών περιστατικών (facts), το εφαρμοστέο νομικό πλαίσιο (νόμοι, άρθρα), καθώς και μεταδεδομένα όπως ημερομηνία, τύπος δικαστηρίου και σύνθεση δικαστών. Ιδιαίτερη σημασία έχει η επιστημονικά τεκμηριωμένη επιλογή των εκβάσεων, π.χ. παραβίαση, μη παραβίαση ή τύπος ποινής. Η εσωτερική συνοχή, η καθαρότητα και η αναλυτικότητα της κάθε εγγραφής επιδρούν άμεσα στην ποιότητα των προβλέψεων του μοντέλου [54], [56], [11].

Η φάση της προεπεξεργασίας (preprocessing) είναι πολύ κρίσιμη. Περιλαμβάνει βασικές ενέργειες όπως καθαρισμό κειμένου (αφαίρεση θορύβου, νορμαλιζατιον χαρακτήρων, ενιαία μορφοποίηση αριθμών και ημερομηνιών), τοκενιζατιον (κατάτμηση σε λέξεις/φράσεις), και λεμματιζατιον για μείωση της πολυμορφίας. Εξειδικευμένες τεχνικές, όπως ο εντοπισμός νομικών οντοτήτων μέσω Named Entity Recognition, επιτρέπουν την ανάδειξη κρίσιμων στοιχείων, όπως άρθρα νόμων, ημερομηνίες, δικαστές και νομικά προηγούμενα [56], [60]. Σε μεγάλα και πολύπλοκα κείμενα, εφαρμόζονται μέθοδοι κατακερματισμού (chunking) ή ιεραρχικής τμηματοποίησης (hierarchical segmentation), ώστε το μοντέλο να επεξεργάζεται τμηματικά τα κείμενα χωρίς να χάνεται το συνολικό νόημα.

Επιπλέον, η επιλογή της μεθόδου αναπαράστασης (Bag-of-Words, TF-IDF, word embeddings, contextual embeddings όπως BERT/LegalBERT) καθορίζει σε σημαντικό βαθμό την απόδοση του μοντέλου. Τα παραδοσιακά sparse representations (BoW, TF-IDF) παραμένουν χρήσιμα ως baseline, όμως τα σύγχρονα contextual embeddings είναι απαραίτητα για την αποτύπωση των βαθύτερων νοηματικών και συντακτικών συσχετισμών, ειδικά σε νομικά κείμενα υψηλής πολυπλοκότητας [56], [11].

Στο παρακάτω διάγραμμα 4.6[54], παρουσιάζεται ένα χαρακτηριστικό διάγραμμα ροής που απεικονίζει τους στόχους και τις απαιτήσεις των τριών βασικών εργασιών στην αυτόματη ανάλυση και πρόβλεψη δικαστικών αποφάσεων: (α) αναγνώριση εκβάσεων (outcome identification), (β) κατηγοριοποίηση αποφάσεων βάσει εκβάσεων (outcome-based judgment categorisation), (γ) πρόβλεψη εκβάσεων (outcome forecasting).



Σχήμα 4.6: Διάγραμμα ροής που απεικονίζει τους στόχους και τις απαιτήσεις τις 3 εργασίες πρόβλεψης δικαστικών αποφάσεων

Το διάγραμμα ροής απεικονίζει τα διακριτά στάδια επεξεργασίας σε συστήματα πρόβλεψης δικαστικών αποφάσεων, με κάθε τύπο εργασίας να ακολουθεί διαφορετικό πιπελινε όσον αφορά τα δεδομένα εισόδου, τις τεχνικές και τον βαθμό ερμηνευσιμότητας. Συγκεκριμένα, η αναγνώριση εκβάσεων (Outcome Identification) βασίζεται σε πλήρη κείμενα αποφάσεων και αξιοποιεί τεχνικές ανάκτησης πληροφορίας και απλούς ταξινομητές· η κατηγοριοποίηση αποφάσεων (Judgement Categorisation) εστιάζει στην ταξινόμηση υποθέσεων ανάλογα με το αποτέλεσμα, αξιοποιώντας feature extraction και σύγχρονους ταξινομητές (SVM, BERT-based κ.λπ.), ενώ η πρόβλεψη εκβάσεων (Outcome Forecasting) αποτελεί το πιο απαιτητικό στάδιο, καθώς το σύστημα λειτουργεί αποκλειστικά με τα πραγματικά περιστατικά και χωρίς πρόσβαση στην τελική απόφαση, απαιτώντας βαθύτερη κατανόηση του νομικού πλαισίου και προηγμένες μεθόδους επεξεργασίας.

Η δυσκολία αυξάνεται σημαντικά από τα απλά στάδια εξαγωγής προς τη γνήσια πρόβλεψη, επιβάλλοντας τη χρήση πιο σύνθετων αρχιτεκτονικών, ενισχυμένης ακρίβειας στην προεπεξεργασία και εξελιγμένων μηχανισμών ερμηνευσιμότητας (explainability)[54]. Η επιλογή κατάλληλης αρχιτεκτονικής για το εκάστοτε προγνωστικό μοντέλο επηρεάζει άμεσα τόσο την ακρίβεια όσο και τη γενικευσιμότητα του συστήματος, με τη σύγχρονη έρευνα να μετατοπίζεται διαρκώς προς βαθιές νευρωνικές αρχιτεκτονικές (deep neural architectures), που επιτρέπουν αποτελεσματική διαχείριση της πολυπλοκότητας των νομικών δεδομένων.

Αρχικά, οι πρώτες προσπάθειες στον χώρο χρησιμοποίησαν αλγορίθμους όπως οι μηχανές διανυσμάτων υποστήριξης (Support Vector Machines, SVM), τα τυχαία δάση (Random Forests) και τη λογιστική παλινδρόμηση (Logistic Regression) πάνω σε αραιές αναπαραστάσεις κειμένου (sparse text representations), όπως το TF-IDF και το Bag-of-Words, BoW.

Η υιοθέτηση βαθιών νευρωνικών δικτύων (deep neural networks) σηματοδότησε μία νέα εποχή για το νομικό NLP. Μοντέλα όπως τα Convolutional Neural Networks (CNNs), Recurrent Neural Networks (RNNs, π.χ. LSTM, GRU), και ειδικότερα τα BiLSTM με μηχανισμούς προσοχής (attention), αποδείχθηκαν θετικά αποτελεσματικά στην επεξεργασία ακολουθιακών δεδομένων και στην ανάδειξη κρίσιμων τμημάτων του κειμένου [56], [11].

Για πολύ μεγάλες αποφάσεις, εφαρμόζονται ιεραρχικές μετασχηματιστικές αρχιτεκτονικές (hierarchical transformers), όπου το κείμενο διασπάται σε επιμέρους segments ή chunks, κάθε ένα από τα οποία επεξεργάζεται ξεχωριστά και στη συνέχεια τα embeddings ενοποιούνται για την τελική πρόβλεψη [56].

Στο πλαίσιο της ποράλληλης μάθησης (multi-task learning), τα μοντέλα εκπαιδεύονται ταυτόχρονα σε περισσότερα από μια εργασία. Για παράδειγμα, σε ένα δικαστικό σύνολο δεδομένων, ένα δίκτυο μπορεί να προβλέπει παράλληλα τόσο το είδος της παραβίασης όσο και το άρθρο που αφορά, ή ακόμα και το εύρος της προβλεπόμενης ποινής. Αυτή η προσέγγιση υλοποιείται με διαμοιραζόμενη ραχοκοκαλιά (shared backbone), όπου τα πρώτα επίπεδα του μοντέλου μαθαίνουν κοινή αναπαράσταση του νομικού κειμένου, και στη συνέχεια εξειδικευμένα κεφάλια (task-specific heads) παράγουν τις αντίστοιχες προβλέψεις για κάθε υποεργασία [11].

Το βασικό πλεονέκτημα της πολυεργασιακής εκπαίδευσης είναι ότι επιτρέπει στο μοντέλο να μεταφέρει γνώση μεταξύ συναφών tasks, βελτιώνοντας τη σταθερότητα και τη γενικευσιμότητα των αποτελεσμάτων. Επιπλέον, επιτυγχάνεται σημαντική μείωση του κινδύνου υπερεκμάθησης (overfitting), καθώς το δίκτυο αναγκάζεται να μάθει αντιπροσωπευτικά χαρακτηριστικά (features) που είναι χρήσιμα σε περισσότερα από ένα προβλήματα.

Πρακτικά, η εκπαίδευση πραγματοποιείται με έναν ενιαίο βρόχο βελτιστοποίησης (optimization loop), όπου η συνολική συνάρτηση κόστους (loss function) αποτελεί συνδυασμό των επιμέρους απωλειών (losses) κάθε task, π.χ.:

$$\mathcal{L}_{total} = \beta_1 \mathcal{L}_{outcome} + \beta_2 \mathcal{L}_{article} + \beta_3 \mathcal{L}_{penalty}$$

όπου τα β_i είναι βάρη που προσαρμόζονται ανάλογα με τη σημασία κάθε υπο-εργασίας (task) ή την ανισορροπία κλάσεων (class imbalance).

Παράλληλα, σε πιο σύνθετα νομικά περιβάλλοντα, εφαρμόζονται υβριδικές προσεγγίσεις όπου τα βαθιά νευρωνικά δίκτυα (deep neural networks) συνδυάζονται με πιθανοκρατικά γραφικά μοντέλα (probabilistic graphical models), όπως τα Bayesian networks [61]. Σε αυτή τη διάταξη, τα ενσωματώματα εξόδου (output embeddings) ή οι αρχικές προβλέψεις του νευρωνικού μοντέλου τροφοδοτούν τις πιθανότητες των κόμβων σε ένα Bayesian δίκτυο. Έτσι, το τελικό σύστημα δεν παρέχει μόνο μια «μαύρου κουτιού» πρόβλεψη, αλλά ποσοτικοποιεί την αβεβαιότητα και επιτρέπει real-time ενημέρωση των πιθανοτήτων κάθε έκβασης με βάση νέα αποδεικτικά στοιχεία.

Τα συστήματα αυτά προσφέρουν μεγαλύτερη διαφάνεια και ερμηνευσιμότητα (explain-

ability), επιτρέποντας στο χρήστη (δικαστή, νομικό, ερευνητή) να αναλύσει γιατί το σύστημα καταλήγει σε συγκεκριμένη πρόβλεψη, βασιζόμενο στην ποσοτική ανάλυση των στοιχείων και όχι μόνο στη «δαισθητική» κρίση ενός νευρωνικού δικτύου.

Η επιτυχία κάθε μοντέλου πρόβλεψης δικαστικών αποφάσεων εξαρτάται όχι μόνο από την επιλογή της κατάλληλης αρχιτεκτονικής, αλλά και από τη μεθοδική διαδικασία εκπαίδευσης και βελτιστοποίησης των παραμέτρων του. Το πρώτο κρίσιμο βήμα είναι η σωστή διάσπαση του dataset σε σύνολα εκπαίδευσης (training), επικύρωσης (validation) και ελέγχου (test), με στόχο την αξιόπιστη αξιολόγηση της γενικευσιμότητας του μοντέλου. Η διαδικασία αυτή ακολουθεί συχνά στρατηγικές διαστρωμάτωσης (stratified splitting), ώστε να διατηρείται η αναλογία των διαφορετικών κατηγοριών εκβάσεων σε όλα τα υποσύνολα [56], [11].

Ένα από τα βασικά προβλήματα στην εκπαίδευση μοντέλων με νομικά δεδομένα είναι η ανισορροπία κατηγοριών (class imbalance), καθώς σε πολλές δικαιοδοσίες ορισμένες εκβάσεις εμφανίζονται σημαντικά συχνότερα από άλλες – για παράδειγμα, η απόφαση «μη παραβίασης» υπερτερεί αριθμητικά της «παραβίασης» στο ECHR. Για την αντιμετώπιση αυτού του φαινομένου εφαρμόζονται διάφορες τεχνικές, όπως η χρήση σταθμισμένων συναρτήσεων απωλειών (weighted loss functions) π.χ. σταθμισμένη διασταυρούμενη εντροπία (weighted cross-entropy) ή focal loss ώστε το μοντέλο να δίνει μεγαλύτερη βαρύτητα στα λάθη που αφορούν σπάνιες κατηγορίες. Επιπλέον, αξιοποιούνται μέθοδοι ενίσχυσης δεδομένων (data augmentation) ειδικά προσαρμοσμένες στη νομική πληροφορία, όπως η αναδιάταξη προτάσεων, η εναλλαγή νομικών παραπομπών ή η εισαγωγή παραπλήσιων αποφάσεων ως νέα δείγματα [11]. Τέλος, συχνά εφαρμόζονται τεχνικές υπερδειγματοληψίας ή υποδειγματοληψίας των κατηγοριών, ώστε να ενισχυθεί η παρουσία των μειονοτικών labels κατά την εκπαίδευση του μοντέλου.

Κατά την ίδια τη φάση της εκπαίδευσης, χρησιμοποιούνται βέλτιστοι αλγόριθμοι βελτιστοποίησης όπως Adam ή AdamW [62], ενώ για μεγαλύτερη σταθερότητα εφαρμόζονται τεχνικές όπως batch normalization, dropout, early stopping και learning rate scheduling. Στα μοντέλα transformers, η χρήση mixed-precision training συμβάλλει στη μείωση της μνήμης και στην επιτάχυνση της εκπαίδευσης χωρίς απώλεια ακρίβειας [56].

Η τελική επιλογή των υπερπαραμέτρων (hyperparameters) (π.χ. ρυθμός μάθησης (learning rate), μέγεθος παρτίδας (batch size), αριθμός στρωμάτων (number of layers) γίνεται συνήθως μέσω εξαντλητικής αναζήτησης (grid search) ή βελτιστοποίησης Bayes (Bayesian optimization) πάνω στο σύνολο επικύρωσης (validation set), ενώ επιπλέον αξιοποιούνται τεχνικές διασταυρούμενης επικύρωσης (cross-validation) όπου το σύνολο δεδομένων το επιτρέπει.

Η αξιολόγηση της απόδοσης δεν περιορίζεται στην απλή μέτρηση της ακρίβειας (accuracy) ή της ακρίβειας θετικών (precision), αλλά γίνεται με πλήρεις μετρικές όπως macro/micro F1-score, πίνακες σύγχυσης (confusion matrices) ανά κατηγορία, ROC/AUC (ειδικά για δυαδικές εργασίες), καθώς και ανάλυση διαμέτρου (calibration analysis) ώστε να διαπιστώνεται η αξιοπιστία των προβλέψεων [56], [63]. Επιπλέον, σε προχωρημένες μελέτες, η αξιολόγηση περιλαμβάνει ανάλυση ερμηνευσιμότητας (explainability analysis) με οπτικοποίηση προσοχής (attention visualization), ανάλυση κοντινότερων γειτόνων (nearest neighbor analysis) και benchmarking με συμμετοχή ανθρώπινου παράγοντα (human-in-the-loop benchmarking), όπου νομικοί συγκρίνουν τις προβλέψεις με πραγματικά προηγούμενα [63].

Ένα από τα βασικότερα ζητούμενα στα συστήματα αυτόματης πρόβλεψης δικαστικών αποφάσεων είναι η ερμηνευσιμότητα (interpretability) και η διαφάνεια των αποτελεσμάτων. Στο νομικό πεδίο, καμία πρόβλεψη όσο τεχνικά άρτια κι αν είναι δεν μπορεί να υιοθετηθεί χωρίς πειστική και τεκμηριωμένη εξήγηση του γιατί οδηγηθήκαμε σε αυτήν. Οι μηχανισμοί ερμηνείας αποκτούν κομβικό ρόλο τόσο για την αποδοχή του συστήματος από τους χρήστες (δικαστές, νομικούς, ερευνητές), όσο και για την αποφυγή ακούσιων σφαλμάτων ή προκαταλήψεων.

Στα σύγχρονα βαθιά μοντέλα, ο μηχανισμός προσοχής (attention mechanism) λειτουργεί ως βασικό εργαλείο ερμηνείας. Μέσω οπτικοποίησης προσοχής (attention visualization), π.χ. με θερμοχάρτες (heatmaps), ο ερευνητής μπορεί να εντοπίσει ποια σημεία του κειμένου «βαραίνουν» περισσότερο στην απόφαση του συστήματος. Επιπλέον, η εφαρμογή τεχνικών ανάλυσης κοντινότερων γειτόνων (nearest neighbor analysis) επιτρέπει την ανάδειξη παρόμοιων υποθέσεων βάσει αποστάσεων ενσωματώσεων (embedding distances), δίνοντας πρακτικά εξήγηση με βάση προηγούμενα (precedent-based explanation) το μοντέλο δείχνει σε ποιες προηγούμενες περιπτώσεις βασίζει την πρόβλεψή του [63].

Παράλληλα, σημαντικό ρόλο διαδραματίζει η ανάλυση ισοτιμίας (fairness analysis), δηλαδή η μελέτη πιθανών μεροληψιών (biases) στο μοντέλο. Τα νομικά σύνολα δεδομένων (datasets) συχνά παρουσιάζουν έμφυτες μεροληψίες (inherent bias) (π.χ. κοινωνικές/πολιτισμικές διαφορές, υπερεκπροσώπηση συγκεκριμένων αποτελεσμάτων). Για την αντιμετώπιση αυτών των κινδύνων, εφαρμόζονται εργαλεία ελέγχου ισοτιμίας (fairness auditing tools), ομαδοποιημένες μετρικές αξιολόγησης (group-specific evaluation metrics), αλλά και τεχνικές ανωνυμοποίησης δεδομένων (data anonymization) και απομερόληψης (de-biasing) στην προεπεξεργασία [56].

4.3.2 Ανάλυση Εκβάσεων Υποθέσεων

Η ανάλυση των εκβάσεων δικαστικών υποθέσεων αποτελεί βασικό εργαλείο αξιολόγησης συστημάτων τεχνητής νοημοσύνης στο δίκαιο και ταυτόχρονα προσφέρει βαθύτερη κατανόηση της ίδιας της νομικής διαδικασίας. Η έκβαση κάθε υπόθεσης δεν είναι απλώς μια δυαδική επιλογή, αλλά το αποτέλεσμα μιας σύνθετης συλλογιστικής που βασίζεται σε πραγματικά περιστατικά, εφαρμοζόμενα άρθρα, νομολογιακές παραδόσεις και συχνά πολιτισμικούς ή κοινωνικούς παράγοντες[56][60].

Σε τεχνικό επίπεδο, η ανάλυση εκβάσεων ξεκινά με την εξαγωγή και κανονικοποίηση των αποτελεσμάτων από τα ακατέργαστα δεδομένα. Τα labels εμφανίζονται συνήθως σε μη ομοιογενή μορφή και απαιτείται εξατομικευμένη χαρτογράφηση, συνδυάζοντας κανόνες και τεχνικές μηχανικής μάθησης, με τη συμβολή νομικών για τη διόρθωση αμφιλεγόμενων ή ασαφών περιπτώσεων[56]. Κατόπιν, τα σύγχρονα υπολογιστικά μοντέλα (όπως transformers, graph neural networks, υβριδικές αρχιτεκτονικές) παρέχουν τόσο την τελική πρόβλεψη όσο και πιθανοτικές εξόδους για κάθε κατηγορία. Αυτές οι εξόδους υφίστανται διακρίβωση (calibration) με τεχνικές όπως Platt scaling ή isotonic regression, διασφαλίζοντας ότι οι προβλεπόμενες πιθανότητες ανταποκρίνονται στις πραγματικές συχνότητες επιτυχίας[63]. Η ποσοτικοποίηση και η διαχείριση της αβεβαιότητας (uncertainty quantification) παραμένουν κρίσιμες τόσο για την ερμηνευσιμότητα όσο και για την αποδοχή των αποτελεσμάτων στην

πράξη.

Σε πολλά νομικά σύνολα δεδομένων απαιτείται η ταυτόχρονη πρόβλεψη πολλαπλών εκβάσεων (όπως για διαφορετικά άρθρα ή για πολλούς εναγόμενους), γεγονός που καθιστά απαραίτητη τη χρήση πολυκατηγορικών μοντέλων (multi-label transformers, graph-based models) και τεχνικών ταυτόχρονης βελτιστοποίησης (joint optimization) και εξισορρόπησης απωλειών (loss balancing)[11]. Η ορθή αντιμετώπιση πολυκατηγορικών και ασαφών εκβάσεων ενισχύει την αξιοπιστία και την πρακτική αξία των συστημάτων πρόβλεψης.

Η συγκριτική ανάλυση εκβάσεων σε διαφορετικές δικαιοδοσίες και νομικά σύνολα δεδομένων αναδεικνύει τη μεγάλη πρόκληση της γενικευσιμότητας (generalization) και την ανάγκη για τεχνικές προσαρμογής πεδίου (domain adaptation) και μεταφοράς γνώσης. Κάθε δικαστικό σύστημα ενσωματώνει ιδιαιτερότητες που δεν αποτυπώνονται πάντα ως στατιστικά μοτίβα, καθιστώντας αναγκαία την προσαρμογή και τον επανασχεδιασμό των μοντέλων ανάλογα με το εκάστοτε νομικό και πολιτισμικό πλαίσιο[56][60][63].

Η ερμηνευσιμότητα, η διαφάνεια και η διεπισημονική συνεργασία παραμένουν βασικές προϋποθέσεις για τη λειτουργία των συστημάτων τεχνητής νοημοσύνης ως αξιόπιστων εργαλείων στον χώρο του δικαίου.

4.3.3 Μετρικές Αξιολόγησης και Αποτελέσματα

Η αποτίμηση της απόδοσης των συστημάτων πρόβλεψης δικαστικών αποφάσεων αποτελεί κρίσιμο στάδιο για την επιστημονική τεκμηρίωση, τη νομική αξιοπιστία και την κοινωνική αποδοχή των εφαρμογών τεχνητής νοημοσύνης στον χώρο του δικαίου. Η επιλογή των μετρικών αξιολόγησης δεν είναι ουδέτερη: επηρεάζει το πώς ερμηνεύονται τα αποτελέσματα και ποια συμπεράσματα προκύπτουν ως έγκυρα σε κάθε νομικό περιβάλλον. Επιπλέον, οι ίδιες οι μετρικές φέρουν εγγενείς περιορισμούς ως προς την ικανότητά τους να αποτυπώνουν την πολυπλοκότητα και την ηθική διάσταση των νομικών αποφάσεων.

Οι βασικές μετρικές περιλαμβάνουν:

- **Ακρίβεια (accuracy):** Ποσοστό ορθών προβλέψεων στο σύνολο των παραδειγμάτων. Αν και δημοφιλής, η ακρίβεια υποτιμά τα προβλήματα κλάσεων με ανισόρροπη κατανομή, ιδίως σε datasets όπου η πλειοψηφία των αποφάσεων είναι παραβίαση ή ένοχος ([56], [11]).
- **Ανάκληση (recall):** Ποσοστό των πραγματικά θετικών περιπτώσεων που προβλέφθηκαν σωστά. Είναι κρίσιμη όταν έχει σημασία να εντοπιστούν όλες οι παραβιάσεις, ακόμη και αν γίνονται κάποια σφάλματα.
- **Ευστοχία (precision):** Ποσοστό των θετικών προβλέψεων που ήταν σωστές. Ιδιαίτερα σημαντική όταν η λανθασμένη επισήμανση (π.χ. ένοχος σε αθώο) έχει υψηλό κόστος.
- **F1-score:** Ο αρμονικός μέσος όρος ανάκλησης και ευστοχίας, που εξισορροπεί τις δύο μετρικές σε σενάρια με ανισοβαρείς κλάσεις.
- **Macro-F1 / Micro-F1:**
 - Macro-F1 δίνει ίσο βάρος σε κάθε κατηγορία, κατάλληλο για μελέτη μειονοτικών εκβάσεων.

- Micro-F1 υπολογίζεται συνολικά, δίνει έμφαση στις πλειοψηφικές κλάσεις.

Σε πολυετικετικές εργασίες (multi-label classification), όπως η πρόβλεψη ταυτόχρονων παραβιάσεων σε υποθέσεις του Ευρωπαϊκού Δικαστηρίου Δικαιωμάτων του Ανθρώπου[64], η αξιολόγηση της απόδοσης βασίζεται σε μια σειρά εξειδικευμένων μετρικών. Ο Hamming loss εκφράζει το ποσοστό των ετικετών (labels) που αποδίδονται λανθασμένα σε κάθε εγγραφή, με χαμηλότερες τιμές να υποδηλώνουν καλύτερη απόδοση του συστήματος. Η ακριβής ταύτιση (exact match) ή κοινή ακρίβεια (joint accuracy) μετρά το ποσοστό των περιπτώσεων όπου όλα τα λαβελς προβλέφθηκαν σωστά, παρέχοντας μια αυστηρή αλλά αντιπροσωπευτική εκτίμηση της συνολικής επίδοσης. Παράλληλα, οι μετρικές ROC-AUC (Ρεχειερ Οπερατινγ ήαραστεριστις) και PR-AUC (Πρεσισιον-Ρεσαλλ) αποδεικνύονται ιδιαίτερα χρήσιμες σε καταστάσεις με περιορισμένο αριθμό θετικών κατηγοριών, καθώς αξιολογούν την ικανότητα του μοντέλου να διαχωρίζει αποτελεσματικά τις διάφορες κλάσεις. Ο συντελεστής συσχέτισης του Ματthews (Matthews Correlation Coefficient, MCC) αποτελεί μια ισορροπημένη μετρική ακόμη και σε περιβάλλοντα με σημαντική ανισορροπία μεταξύ των κατηγοριών. Τέλος, σημαντικό ρόλο διαδραματίζουν και οι μέθοδοι βαθμονόμησης και αξιοπιστίας (calibration), διασφαλίζοντας ότι οι προβλεπόμενες πιθανότητες αντικατοπτρίζουν με ακρίβεια τα πραγματικά ποσοστά επιτυχίας.

Ένα κρίσιμο, συχνά υποτιμημένο στοιχείο, είναι η βαθμονόμηση των πιθανοτήτων (calibration) που αποδίδει το μοντέλο. Συχνά τα μοντέλα δείχνουν υψηλή ακρίβεια αλλά η βεβαιότητά τους είναι ψευδής — ένα φαινόμενο που έχει τεράστια σημασία σε νομικά περιβάλλοντα, όπου η αξιοπιστία της απόφασης πρέπει να τεκμηριώνεται. Χαρακτηριστικές τεχνικές όπως reliability diagrams και expected calibration error (ECE) εφαρμόζονται σε [63], αναδεικνύοντας το πόσο υπερσιγουρα ή αβέβαια είναι τα μοντέλα στην πράξη.

Η επιστημονική λογοδοσία (scientific accountability) απαιτεί τη συνεχή παρακολούθηση και αναφορά των failure cases, δηλαδή εκείνων των προβλέψεων που οδηγούν σε εντελώς εσφαλμένα αποτελέσματα. Όπως φαίνεται σε [54], τέτοιες περιπτώσεις αφορούν συχνά υποθέσεις με πολύπλοκο νομικό πλαίσιο, ασαφή πραγματικά περιστατικά ή ασυνήθιστες νομικές καταλήξεις.

Χαρακτηριστικά παραδείγματα εφαρμογής των παραπάνω μετρικών και τεχνικών εντοπίζονται σε διεθνή σύνολα δεδομένων και νομικές δικαιοδοσίες, καταδεικνύοντας τόσο την πρόοδο όσο και τις ανοιχτές προκλήσεις στον χώρο. Στο Chinese Supreme Court Dataset, η χρήση βαθιών νευρωνικών δικτύων με ενσωματώσεις τύπου BERT οδήγησε σε υψηλές τιμές macro-F1, ωστόσο η κοινή ακρίβεια (joint accuracy) παρέμεινε χαμηλή, αναδεικνύοντας τα όρια των σημερινών μεθόδων σε πλήρη πολυκατηγορική πρόβλεψη [11]. Στο Indian Legal Dataset, διαπιστώθηκε ότι η καθαρότητα των ετικετών επηρεάζει έντονα την αξιοπιστία και τη γενικευσιμότητα των μοντέλων, υπογραμμίζοντας τη σημασία της επιμελούς επιμέλειας δεδομένων (data curation) [60]. Αντίστοιχα, στο σύνολο δεδομένων του Συμβουλίου της Ευρώπης (Council of Europe corpus), η εφαρμογή μοντέλων όπως το Legal-BERT βελτίωσε την απόδοση σε δείκτες όπως το F1, αλλά μείωσε τη βαθμονόμηση (calibration), γεγονός που εντείνει τις ανησυχίες για την ερμηνευσιμότητα (interpretability) των αποτελεσμάτων. Η έρευνα σε νέες δικαιοδοσίες, όπως το νομικό σύστημα της Τουρκίας, επιβεβαιώνει την προσαρμοστικότητα των σύγχρονων μεθόδων, αλλά και τις ιδιαίτερες προκλήσεις που προ-

κύπτουν από τη διαφορετικότητα των δικαστηρίων και τη φύση των υποθέσεων. Σε αυτή τη μελέτη, deep learning μοντέλα όπως τα LSTM / BiLSTM με μηχανισμό attention επέδειξαν υψηλή αποτελεσματικότητα, με την απόδοση να εξαρτάται αφενός από τη δομή του δικαστηρίου και αφετέρου από την πολυπλοκότητα των υποθέσεων [57].

Πέραν της επιλογής μετρικών και μεθόδων αξιολόγησης, η πρόσφατη βιβλιογραφία αναδεικνύει τη σημασία της γλωσσικής και θεματικής εξειδίκευσης των μοντέλων. Μελέτες όπως αυτή του BERTimbau αποδεικνύουν ότι μονογλωσσικά μοντέλα, προεκπαιδευμένα σε τοπικά δεδομένα, υπερέχουν των γενικών πολυγλωσσικών προσεγγίσεων τόσο στη διαχείριση της νομικής ορολογίας όσο και στη συνολική απόδοση [51]. Επιπρόσθετα, η εμφάνιση πολυδιάστατων benchmarks όπως το [5] μετατοπίζει το ενδιαφέρον από τις κλασικές μετρικές προς πιο σύνθετα σενάρια αξιολόγησης, τα οποία περιλαμβάνουν όχι μόνο την ανάκτηση πληροφοριών και άρθρων, αλλά και τη σύνθεση, τεκμηρίωση και ανάλυση πλήρων νομικών απαντήσεων [5]. Παρά την πρόοδο, οι μελέτες καταδεικνύουν ότι τα LLMs παρουσιάζουν αξιοσημείωτα κενά στη νομική ανάλυση συγκριτικά με τους καλύτερους ανθρώπινους αξιολογητές, γεγονός που αναδεικνύει τη διαρκή ανάγκη για ερμηνευσιμότητα (explainability) και ανθρώπινη επίβλεψη.

Οι μετρικές από μόνες τους δεν επαρκούν για την αξιολόγηση των συστημάτων πρόβλεψης δικαστικών αποφάσεων. Το ζητούμενο δεν είναι μόνο η τεχνική απόδοση, αλλά και ο βαθμός αποδοχής τους από το νομικό σώμα και την κοινωνία. Η ακρίβεια χρειάζεται να συνδυάζεται με ερμηνευσιμότητα, ενώ η ανάκληση δεν αποκτά αξία αν δεν υπάρχει εμπιστοσύνη στους μηχανισμούς λογικής εξαγωγής της απόφασης. Η σύγχρονη ερευνητική προσέγγιση στρέφεται πλέον σε υβριδικά σχήματα αξιολόγησης που συνδυάζουν τυπικές μετρικές, βαθμολογία, αναλύσεις αποτυχημένων περιπτώσεων (failure case inspection) και ανθρώπινη επιβεβαίωση (human-in-the-loop).

Η διεθνής βιβλιογραφία καταδεικνύει ότι, παρά τις βελτιώσεις στις ποσοτικές μετρήσεις, τα πιο εξελιγμένα μοντέλα συχνά υστερούν σε διαφάνεια και κοινωνική αποδοχή ιδίως σε περιβάλλοντα όπου οι αποφάσεις φέρουν σημαντικό νομικό και κοινωνικό βάρος[65][39]. Η ενσωμάτωση εναλλακτικών δεικτών, όπως η συστημική ερμηνευσιμότητα και η calibration των πιθανοτήτων, καθίσταται αναγκαία για την ουσιαστική επιχειρησιακή εφαρμογή τέτοιων συστημάτων.

Συνολικά, η πρόοδος στον τομέα της πρόβλεψης δικαστικών αποφάσεων δεν εξαρτάται μόνο από την τεχνολογική υπεροχή του μοντέλου, αλλά και από την ποιότητα και την πληρότητα των αξιολογητικών μηχανισμών που το πλαισιώνουν. Η διαμόρφωση ενός πολυδιάστατου, ενιαίου πλαισίου αξιολόγησης παραμένει ζητούμενο για την έρευνα και θεμέλιο για τη μελλοντική ενσωμάτωση αυτών των εργαλείων στην απονομή της δικαιοσύνης.

4.4 Έξυπνα Δικαστήρια και Συστήματα Ψηφιακής Δικαιοσύνης

Τα έξυπνα δικαστήρια συνιστούν μια σύγχρονη αρχιτεκτονική παρέμβαση στη λειτουργία της δικαιοσύνης, όπου τεχνολογίες όπως η τεχνητή νοημοσύνη, τα μεγάλα γλωσσικά μοντέλα, το blockchain, οι ψηφιακές ταυτότητες και τα συνεργατικά συστήματα ενσωματώνονται στη διαδικασία απονομής του δικαίου[66]. Στόχος δεν είναι μόνο η επιτάχυνση των διαδικασιών ή η αποφόρτιση του διοικητικού έργου των δικαστηρίων, αλλά η ριζική αναδιαμόρφωση της

σχέσης πολίτη-δικαιοσύνης μέσω διαφάνειας, πρόσβασης και λογοδοσίας.

Η πιο συστηματική υλοποίηση αυτής της φιλοσοφίας συναντάται στην Κίνα, όπου το Ανώτατο Λαϊκό Δικαστήριο έχει θεσπίσει το εθνικό Smart Court System, συνδυάζοντας ηλεκτρονική κατάθεση, αυτοματοποιημένη εξαγωγή δεδομένων, πρόβλεψη αποφάσεων και παραγωγή προσχεδίων με χρήση εργαλείων NLP και κανόνων βασισμένων σε δεδομένα (data-driven rule-based systems) [67]. Παράλληλα, άλλες χώρες υιοθετούν διαφορετικές εκδοχές ψηφιακής δικαιοσύνης, όπως η χρήση blockchain για τη διασφάλιση της ακεραιότητας αποδεικτικών στοιχείων στην Ουκρανία [68], ή τα υβριδικά e-courts στην Ινδονησία, που επικεντρώνονται στην προσβασιμότητα και την εξ αποστάσεως συμμετοχή των διαδίκων [69].

Η πανδημία COVID-19 επιτάχυνε δραματικά τη μετάβαση, με πλήθος δικαστηρίων διεθνώς να υιοθετούν πλατφόρμες τηλεδιασκέψεων, ψηφιακές καταθέσεις και απομακρυσμένες ακροάσεις, μεταβάλλοντας ταυτόχρονα τόσο τις διαδικαστικές πρακτικές όσο και την ίδια τη φυσιογνωμία της δικαστικής διαδικασίας [70]. Ωστόσο, οι τεχνολογικές αυτές λύσεις αναδεικνύουν και νέες προκλήσεις: ο κίνδυνος αδιαφανούς αυτοματισμού, η ανάγκη για ερμηνευσιμότητα και διαφάνεια, οι κοινωνικές ανισότητες στην πρόσβαση, καθώς και η προστασία των θεμελιωδών δικαιωμάτων.

Κεντρικό διακύβευμα παραμένει η ισορροπία ανάμεσα στην τεχνολογική καινοτομία και τη θεσμική νομιμοποίηση. Η εισαγωγή τεχνητής νοημοσύνης, αυτοματοποιημένων συστημάτων ανάλυσης και ψηφιακών ταυτοτήτων μπορεί να βελτιώσει σημαντικά την αποδοτικότητα και τη συνέπεια, όμως μόνο υπό την προϋπόθεση της ύπαρξης σαφών μηχανισμών ελέγχου, λογοδοσίας και επανεξέτασης. Η διεθνής εμπειρία δείχνει ότι η επιτυχία των έξυπνων δικαστηρίων κρίνεται τελικά από τη διασταύρωση τεχνολογικής αποτελεσματικότητας και θεσμικής εγκυρότητας [12][66].

4.4.1 Retrieval-Augmented Generation σε Νομικά Περιβάλλοντα

Όπως παρουσιάστηκε αναλυτικά στην ενότητα 4.1, η αρχιτεκτονική Retrieval-Augmented Generation (RAG) αποτελεί μία συνδυαστική προσέγγιση που αξιοποιεί ταυτόχρονα τεχνικές ανάκτησης πληροφορίας (retrieval) και παραγωγής απαντήσεων (generation) μέσω Μεγάλων Γλωσσικών Μοντέλων. Πιο συγκεκριμένα, τα συστήματα RAG αναζητούν σχετικά νομικά τεκμήρια (π.χ. αποφάσεις, συμβάσεις, άρθρα) από βάσεις δεδομένων και στη συνέχεια τα ενσωματώνουν ως περιβάλλον συμφραζομένων (context) στο ερώτημα που τίθεται, προκειμένου το γλωσσικό μοντέλο να συνθέσει απαντήσεις που είναι τεκμηριωμένες, επίκαιρες και ερμηνεύσιμες.

Η προσέγγιση αυτή γεφυρώνει το χάσμα μεταξύ παραδοσιακών μηχανών αναζήτησης (όπως η αναζήτηση με λέξεις-κλειδιά ή η σημασιολογική αναζήτηση) και της σύγχρονης γλωσσικής κατανόησης, προσφέροντας ένα ισχυρό υπόβαθρο για νομικές εφαρμογές που απαιτούν ακρίβεια και διαφάνεια.

Παράλληλα, προηγμένες τεχνικές έχουν ενσωματωθεί στα συστήματα RAG ώστε να ενισχυθεί η αξιοπιστία και η πληρότητα των απαντήσεων. Για παράδειγμα, η τεχνική του πολυσταδιακού συλλογισμού (multi-hop reasoning) επιτρέπει στο μοντέλο να αντλεί πληροφορίες από πολλαπλές πηγές ή αποσπάσματα και να τις συνδυάζει σε ένα ενιαίο συμπέρασμα, κάτι ιδιαίτερα χρήσιμο σε περιπτώσεις όπου η νομική πληροφορία είναι κατακεκομμένη. Αντίστοι-

χα, η δυναμική κατασκευή προτροπών (dynamic prompt assembly) επιτρέπει την επιλογή και ανασύνθεση των πλέον κατάλληλων αποσπασμάτων από τις πηγές, με βάση το περιεχόμενο του ερωτήματος, ενισχύοντας έτσι τη συνάφεια και την ακρίβεια των απαντήσεων.

Τέλος, η τεχνική ανίχνευσης παραπομπών (citation tracing) προσθέτει ένα επίπεδο διαφάνειας, καθώς επιτρέπει στο σύστημα να αναφέρει με σαφήνεια τις πηγές από τις οποίες αντλήθηκαν τα δεδομένα που χρησιμοποιήθηκαν στην τελική απάντηση, στοιχείο κρίσιμο σε νομικά πλαίσια όπου απαιτείται δυνατότητα ελέγχου και τεκμηρίωσης.

Η τυπική αρχιτεκτονική RAG περιλαμβάνει τα εξής βασικά στάδια:

1. Προεπεξεργασία και Εμπλουτισμός Νομικών Δεδομένων

Συλλογή μεγάλου όγκου νομικών εγγράφων (αποφάσεις, συμβόλαια, άρθρα, νομολογία). Προκαταρκτική επεξεργασία (καθαρισμός, normalization, chunking σε αποσπάσματα 256–512 tokens, εμπλουτισμός με μεταδεδομένα όπως ημερομηνία, κατηγορία υπόθεσης, δικαστήριο, σχετικές διατάξεις). Η διαδικασία αυτή εξασφαλίζει τη δομημένη και διαλειτουργική διαχείριση της πληροφορίας.

2. Δημιουργία Διανυσματικής Αναπαράστασης (Embedding):

Κάθε απόσπασμα νομικού κειμένου μετατρέπεται σε διάνυσμα υψηλών διαστάσεων μέσω εξειδικευμένων γλωσσικών μοντέλων όπως LegalBERT, CaseLawBERT, ή custom fine-tuned μετασχηματιστών για το εκάστοτε νομικό σύστημα. Αυτά τα embeddings αποτυπώνουν τόσο τη θεματική όσο και τη σημασιολογική συνάφεια του περιεχομένου.

3. Δημιουργία Δείκτη Ανάκτησης (Vector Indexing):

Τα embeddings αποθηκεύονται σε εξειδικευμένες βάσεις δεδομένων διανυσμάτων (π.χ. FAISS, Weaviate, Qdrant, Pinecone), επιτρέποντας αποδοτική αναζήτηση σε μεγάλης κλίμακας σε σώματα κειμένων μεγάλης κλίμακας. Οι βάσεις αυτές προσφέρουν τόσο πυκνή (dense) όσο και υβριδική (hybrid) αναζήτηση, εξασφαλίζοντας ακρίβεια και ταχύτητα στην ανάκτηση.

4. Ανάκτηση Συναφών Εγγράφων (Retrieval):

Όταν ο χρήστης υποβάλλει ένα ερώτημα (query), αυτό μετατρέπεται επίσης σε embedding και αναζητούνται τα πιο σχετικά αποσπάσματα με βάση μέτρα ομοιότητας (cosine similarity, dot product). Εξειδικευμένοι αλγόριθμοι ρερανκινγκ (όπως MMR ή BM25+dense hybrid) χρησιμοποιούνται για τη βελτίωση της σχετικότητας των αποτελεσμάτων.

5. Παραγωγή Απάντησης (Generation):

Τα επιλεγμένα αποσπάσματα τροφοδοτούν το prompt ενός Μεγάλου Γλωσσικού Μοντέλου (LLM, π.χ. GPT-4, Llama 3), το οποίο συνθέτει μια απάντηση που βασίζεται στα ανακτημένα τεκμήρια. Προχωρημένες τεχνικές περιλαμβάνουν συμπίεση ερωτήματος (prompt compression), αναδιάταξη αποσπασμάτων (chunk reordering), συλλογισμό πολλαπλών βημάτων (multi-hop reasoning) και ανίχνευση παραπομπών (citation tracing), ώστε η απάντηση να παραμένει τεκμηριωμένη και επαληθεύσιμη.

Η τεχνολογία RAG έχει πλέον καθιερωθεί σε ένα ευρύ φάσμα νομικών εφαρμογών που απαιτούν άμεση, ακριβή και τεκμηριωμένη επεξεργασία πληροφορίας. Πλατφόρμες νομικής έρευνας όπως το Lexis+ AI και το Westlaw Edge αξιοποιούν RAG pipelines για προηγμένη αναζήτηση νομολογίας, αυτόματη περίληψη αποφάσεων και εντοπισμό σχετικών προηγούμενων, ενώ ενσωματώνουν σημασιολογική ανάκτηση (semantic retrieval), δυναμική υποβολή ερωτημάτων (dynamic prompting), ανίχνευση παραπομπών (citation tracing) και επανακατάταξη (reranking) για τη μείωση των παραισθήσεων (hallucinations) και την αύξηση της αξιοπιστίας των απαντήσεων [71], [4].

Παράλληλα, στην αυτοματοποιημένη επεξεργασία συμβολαίων, οι αρχιτεκτονικές RAG επιτρέπουν στοχευμένη αναζήτηση ρητρών, έλεγχο συμμόρφωσης (compliance) και αυτόματη επισήμανση κρίσιμων όρων, αξιοποιώντας εξειδικευμένες ενσωματώσεις (embeddings), κατακερματισμένη ευρετηρίαση (chunked indexing) και μετα-επεξεργασία αποτελεσμάτων με νομικά ευαισθητοποιημένα LLMs [72].

Επιπλέον, τα RAG συστήματα χρησιμοποιούνται ως εργαλεία υποστήριξης δικαστών, προσφέροντας δυναμική αναζήτηση και συσχέτιση νομολογιακών πηγών, αυτόματη σύνθεση συνοψισμένων προηγούμενων και επισήμανση σχετικών άρθρων του νόμου, με την απόδοση αυτών των εργαλείων να αξιολογείται μέσω εξειδικευμένων μετρικών επικάλυψης (overlap metrics) και δεικτών ερμηνευσιμότητας (explainability indices) [43].

Τέλος, ιδιαίτερη έμφαση δίνεται στη διαχείριση πολυγλωσσικών και πολυδιάστατων δεδομένων, με τις σύγχρονες πλατφόρμες να υποστηρίζουν αποτελεσματική αναζήτηση και ανάλυση ακόμη και σε περιβάλλοντα με πολύπλοκη ορολογία και ποικιλία νομικών συστημάτων.

Ειδικά σε ευρωπαϊκά και διεθνή νομικά περιβάλλοντα, η RAG αρχιτεκτονική αξιοποιεί πολυγλωσσικές ενσωματώσεις για επεξεργασία νομολογίας σε διαφορετικές γλώσσες και έννομες τάξεις. Αυτό είναι κρίσιμο για διασυννοριακές υποθέσεις, συγκριτική ανάλυση και ενιαία πρόσβαση σε πολυγλωσσικά σώματα κειμένων [53].

Η ταχεία εξέλιξη της RAG συνοδεύεται από σημαντικές τεχνολογικές καινοτομίες, όπως η υβριδική ανάκτηση (hybrid retrieval), που συνδυάζει αραιές (sparse, π.χ. TF-IDF, λέξεις-κλειδιά) και πυκνές (dense, βασισμένες σε ενσωματώσεις) μεθόδους, ενισχύοντας την ανάκληση (recall) και την ακρίβεια (precision), ιδίως σε σενάρια με περιορισμένα δεδομένα ανά γλώσσα [4]. Επιπλέον, η ενσωμάτωση γράφων γνώσης (knowledge graphs) στα RAG pipelines επιτρέπει την αξιοποίηση δομημένων σχέσεων μεταξύ νομικών οντοτήτων για βαθύτερη κατανόηση και επεξήγηση των συσχετισμών [43]. Οι σύγχρονες αρχιτεκτονικές υιοθετούν επίσης τεχνικές πολυσταδιακής συλλογιστικής (multi-hop reasoning), όπου πληροφορίες ανακτώνται αλυσιδωτά από πολλαπλές πηγές πριν τη σύνθεση της τελικής απάντησης, καθώς και δυναμική ανασύνθεση prompts (dynamic prompt assembly) με επιλογή αποσπασμάτων βάσει συμφραζομένων (context-aware chunk selection), με στόχο την καλύτερη τεκμηρίωση και συμπύκνωση της πληροφορίας [71], [43].

4.4.2 Μελέτη περίπτωσης: Εφαρμογή σε διαφορετικές δικαιοδοσίες

Η δυναμική των αρχιτεκτονικών RAG αναδεικνύεται ακόμα περισσότερο μέσα από τις διαφορετικές στρατηγικές εφαρμογής τους σε εθνικά και υπερεθνικά νομικά συστήματα.

Η παρακάτω ανάλυση μελετά συγκριτικά παραδείγματα υλοποίησης σε Κίνα, Ευρωπαϊκή Ένωση, ΗΠΑ, Τουρκία και Φιλιππίνες, αναδεικνύοντας τα τεχνικά, θεσμικά και κοινωνικά χαρακτηριστικά κάθε προσέγγισης.

Η πλέον εντυπωσιακή εφαρμογή συστημάτων ψηφιακής δικαιοσύνης εντοπίζεται στην Κίνα, όπου το Ανώτατο Λαϊκό Δικαστήριο (Supreme People's Court, SPC) έχει υλοποιήσει ένα πλήρως ψηφιοποιημένο Smart Court System, συνδυάζοντας αρχές τεχνητής νοημοσύνης, ηλεκτρονικής διακυβέρνησης και ανάλυσης μεγάλων δεδομένων [67]. Η πλατφόρμα αξιοποιεί τεχνικές NLP-RAG ώστε να επιτρέπει την αυτοματοποιημένη επεξεργασία υποθέσεων, την αναζήτηση σχετικής νομολογίας και τη σύνταξη προσχεδίων αποφάσεων. Συγκεκριμένες εφαρμογές, όπως το Xiao Zhi Fa (AI Θυδγε), μπορούν να εκδικάζουν αυτόματα απλές αστικές και εμπορικές διαφορές μέσω σημασιολογικής αναζήτησης και παραγωγής φυσικής γλώσσας (semantic search και natural language generation), ενώ τα διαδικτυακά δικαστήρια (e-courts) λειτουργούν εξ ολοκλήρου εξ αποστάσεως, παρέχοντας πλήρη ψηφιακή εμπειρία στους διαδίκους. Μέχρι το 2022, περισσότερες από τρία εκατομμύρια υποθέσεις είχαν επιλυθεί μέσω αυτών των συστημάτων. Ωστόσο, σοβαρές ανησυχίες παραμένουν ως προς την ερμηνευσιμότητα (explainability), την ιχνηλασιμότητα (traceability) και τη θεσμική λογοδοσία, με αποτέλεσμα η διαφάνεια και η ανεξαρτησία της δικαιοσύνης να παραμένουν αμφιλεγόμενες [73].

Ενδεικτική είναι η περίπτωση της Ελβετίας, όπου το ζήτημα της αυτόματης νομικής μετάφρασης αποκτά ιδιαίτερη βαρύτητα λόγω της συνύπαρξης τεσσάρων επίσημων γλωσσών και της ανάγκης για άμεση πρόσβαση σε νομικά κείμενα σε όλες τις γλώσσες. Το έργο SwiLTra-Bench εισήγαγε μια μεγάλη πολυγλωσσική βάση δεδομένων με πάνω από 180.000 αντιστοιχισμένα νομικά κείμενα, δικαστικές αποφάσεις και ανακοινώσεις τύπου, καλύπτοντας τις επίσημες γλώσσες της Ελβετίας και τα αγγλικά. Σκοπός του ήταν να αξιολογήσει και να βελτιώσει τις επιδόσεις των μεγάλων γλωσσικών μοντέλων στη μετάφραση νομικών κειμένων[74].

Στις Ηνωμένες Πολιτείες[10], η ενσωμάτωση LLMs και αρχιτεκτονικών RAG χαρακτηρίζεται από έντονη εμπορική αξιοποίηση, κυρίως μέσω του ιδιωτικού τομέα και των νομικών τεχνολογικών εταιρειών (legal tech companies). Πλατφόρμες όπως Lexis+ AI, Casetext Co-Counsel και Westlaw Edge αξιοποιούν fine-tuned language models και προηγμένες τεχνικές αναζήτησης και παραγωγής (semantic retrieval, re-ranking, dynamic generation) για τη νομική έρευνα, την ανάλυση εγγράφων και την υποστήριξη των δικηγόρων σε καθημερινές εργασίες. Η χρήση eDiscovery με active learning και ομαδοποίηση (clustering) έχει ήδη υπερβεί τις ανθρώπινες δυνατότητες ως προς τον όγκο και την ταχύτητα επεξεργασίας δεδομένων. Ωστόσο, το θεσμικό πλαίσιο παραμένει ασαφές, με ελάχιστη κρατική παρέμβαση και απουσία ομοσπονδιακής ρύθμισης για τη χρήση AI σε δικαστικό επίπεδο, με αποτέλεσμα περιστατικά εσφαλμένης χρήσης να αναδεικνύουν την ανάγκη για αυξημένο oversight.

Στη Βραζιλία[75], η χρήση τεχνητής νοημοσύνης έχει εξαπλωθεί σε περισσότερα από τα μισά ανώτατα και ανώτερα δικαστήρια, με πρωτοβουλίες όπως το σύστημα VICTOR για αυτοματοποιημένη προκαταρκτική ανάλυση υποθέσεων και το ATHOS για την υποστήριξη των δικαστών στη διαχείριση μεγάλων όγκων εγγράφων. Το ρυθμιστικό πλαίσιο που θέσπισε το Εθνικό Συμβούλιο Δικαιοσύνης επιβάλλει αυστηρούς όρους διαφάνειας, προστασίας προσωπικών δεδομένων και ασφάλειας, ενώ η χρήση κρατικών πηγών και η δυνατότητα ελέγχου

των συστημάτων θεωρούνται απαραίτητες για τη διασφάλιση της ποιότητας των αποφάσεων.

Οι Φιλιππίνες[12] παρουσιάζουν ένα ενδιαφέρον παράδειγμα συνεργασίας του Ανώτατου Δικαστηρίου με το Ινστιτούτο Προηγμένης Επιστήμης και Τεχνολογίας για την υιοθέτηση τεχνικών NLP και ML στην πρόβλεψη και ανάλυση δικαστικών αποφάσεων. Το αναπτυχθέν σύστημα ανέλυσε τη δομή και το περιεχόμενο χιλιάδων αποφάσεων, με χρήση μοντέλων BERT και LEGAL-BERT, και πέτυχε ποσοστά ακρίβειας 78–82% στη δυαδική ταξινόμηση υποθέσεων (binary classification). Η χρήση ελαφριών υποσυστημάτων RAG επέτρεψε τη διαγλωσσική αναφορά σε σχετικές αποφάσεις, ενώ το Ανώτατο Δικαστήριο διατήρησε επικουρικό αλλά όχι αποφασιστικό ρόλο για το σύστημα. Η εμπειρία των Φιλιππίνων δείχνει ότι ακόμη και δικαστικά συστήματα με περιορισμένους πόρους μπορούν να αξιοποιήσουν τεχνικές αιχμής, αρκεί να υπάρχει συνεργασία με την ερευνητική κοινότητα και θεσμική υποστήριξη.

Τέλος η Τουρκία [57] έχει επιδείξει σημαντική ερευνητική δραστηριότητα στον τομέα της πρόβλεψης δικαστικών αποφάσεων μέσω τεχνικών μηχανικής μάθησης και γλωσσικής επεξεργασίας. Ερευνητικά πιπελινες πρόβλεψης στο Yargitay (Ανώτατο Δικαστήριο) βασίστηκαν σε BERT-based embeddings, vector retrieval και binary classification για την εκτίμηση της πιθανότητας παραβίασης, με ποσοστά ακρίβειας άνω του 75% σε συγκεκριμένες κατηγορίες υποθέσεων. Παρόλο που η θεσμική ενσωμάτωση των συστημάτων αυτών βρίσκεται ακόμη σε πρώιμο στάδιο, η Τουρκία λειτουργεί ως χαρακτηριστικό παράδειγμα εφαρμογής ΑΙ σε δικαστικά συστήματα με περιορισμένα, ετερογενή ή ημιδομημένα δεδομένα [57]. Νεότερες έρευνες δοκίμασαν επίσης διαγλωσσικά ενσωματώματα (cross-lingual embeddings) και RAG pipelines για την πειραματική πρόβλεψη εκβάσεων μεταξύ τουρκικών και ευρωπαϊκών υποθέσεων.

Η επιτυχία τέτοιων εφαρμογών προϋποθέτει θεσμική προσαρμογή, διαφανείς κανόνες λογοδοσίας και ενσωμάτωση στους υφιστάμενους μηχανισμούς απονομής δικαίου. Συνολικά, η τεχνολογική καινοτομία μπορεί να ενισχύσει σημαντικά τη δικαιοσύνη, μόνο εφόσον λαμβάνει υπόψη τις νομικές, κοινωνικές και πολιτισμικές ιδιαιτερότητες κάθε χώρας, λειτουργώντας συμπληρωματικά και όχι ανταγωνιστικά προς το παραδοσιακό δικαιοσύνη πλαίσιο.

4.5 Εργαλεία Νομικής Υποβοήθησης

Η ανάπτυξη της τεχνητής νοημοσύνης και ειδικά των Μεγάλων Γλωσσικών Μοντέλων έχει διαμορφώσει ένα νέο οικοσύστημα εργαλείων για την παροχή νομικής υποστήριξης. Τα εργαλεία αυτά, που περιλαμβάνουν νομικούς ψηφιακούς βοηθούς (virtual legal assistants), συστήματα συνομιλίας (chatbots), εφαρμογές απευθείας εξυπηρέτησης πελατών και ολοκληρωμένες πλατφόρμες διαχείρισης υποθέσεων, αποτελούν πια κρίσιμη υποδομή για τον εκσυγχρονισμό της νομικής πρακτικής και την ενίσχυση της πρόσβασης στη δικαιοσύνη.

Τα σύγχρονα εργαλεία νομικής υποβοήθησης βασίζονται σε τεχνολογίες επεξεργασίας φυσικής γλώσσας (Natural Language Processing), διανυσματικών αναπαραστάσεων (embeddings), και συχνά ενσωματώνουν μηχανισμούς ανάκτησης και ενισχυμένης παραγωγής (RAG) για την παραγωγή τεκμηριωμένων και προσωποποιημένων απαντήσεων. Η διασύνδεση αυτών των συστημάτων με βάσεις δεδομένων νομοθεσίας, νομικής αρθρογραφίας και προηγούμενων αποφάσεων επιτρέπει την αυτοματοποίηση κρίσιμων διαδικασιών, τη μείωση

του χρόνου εξυπηρέτησης και την παροχή νομικής πληροφόρησης υψηλής ποιότητας ακόμα και σε μη ειδικούς χρήστες.

Η ενότητα αυτή εξετάζει τις βασικές κατηγορίες εργαλείων νομικής υποβοήθησης, με έμφαση στη χρήση συστημάτων συνομιλίας (chatbots) και εικονικών νομικών βοηθών (virtual legal assistants) για τη διαμεσολάβηση μεταξύ πολιτών και νομικών υπηρεσιών, στις εφαρμογές με άμεση πρόσβαση στον πελάτη (client-facing applications) που διευκολύνουν τη νομική επικοινωνία και συμβουλευτική, καθώς και στην ολοκληρωμένη ενσωμάτωση των μεγάλων γλωσσικών μοντέλων (LLMs) στα εργαλεία διαχείρισης υποθέσεων (case management tools) και τις συνεργατικές νομικές εργασίες (collaborative legal workflows). Ιδιαίτερη έμφαση δίνεται στην τεχνική ανάλυση, την πρακτική υλοποίηση και την κριτική αξιολόγηση των λύσεων αυτών, αναδεικνύοντας τα οφέλη, τις προκλήσεις και τις προοπτικές που προσφέρει η ενσωμάτωσή τους στη σύγχρονη νομική πρακτική.

4.5.1 Νομικά Chatbots και Εικονικοί Βοηθοί

Τα νομικά συστήματα συνομιλίας (legal chatbots) και οι εικονικοί βοηθοί (virtual assistants) έχουν εξελιχθεί σε προηγμένες τεχνολογικές λύσεις που μετασχηματίζουν την αλληλεπίδραση πελατών, νομικών επαγγελματιών και δικαστικών συστημάτων με τη νομική πληροφορία. Η υλοποίηση αυτών των εργαλείων συνδυάζει τεχνικές επεξεργασίας φυσικής γλώσσας, αρχιτεκτονικές μεγάλων γλωσσικών μοντέλων και pipelines αναζήτησης τεκμηρίων (RAG) για την παροχή ακριβών, τεκμηριωμένων και ερμηνεύσιμων απαντήσεων [72].

Η τυπική τεχνική υποδομή τέτοιων εργαλείων αποτελείται από διαδοχικά στάδια που ξεκινούν με την επεξεργασία φυσικής γλώσσας (NLP pipeline), η οποία περιλαμβάνει τον τεμαχισμό (tokenization), την αναγνώριση οντοτήτων (entity recognition) και τη σημασιολογική ανάλυση (semantic parsing), επιτρέποντας στο σύστημα να κατανοεί σύνθετα νομικά ερωτήματα [72]. Ακολουθεί η κατηγοριοποίηση προθέσεων (intent classification), όπου εξειδικευμένα μοντέλα μετασχηματιστών (transformers), όπως τα LegalBERT, CaseLawBERT ή GreekLegalBERT, χρησιμοποιούνται για την ακριβή ανίχνευση νομικών εννοιών και την αναγνώριση της πρόθεσης του χρήστη. Το στάδιο της ανάκτησης και παραγωγής (retrieval and generation) αξιοποιεί τεχνικές διανυσματικής αναζήτησης (vector retrieval) με διανυσματικές αναπαραστάσεις (embeddings), π.χ. μέσω text-embedding-ada-002 ή multilingual-e5-base, καθώς και συστήματα RAG για τη βελτίωση της τεκμηρίωσης των απαντήσεων. Τέλος, η διαχείριση διαλόγου (dialogue management) υλοποιείται με τη βοήθεια μηχανών κατάστασης (state machines) ή μηχανισμών βασισμένων σε προσοχή (attention-based), ώστε το σύστημα να διαχειρίζεται πολυσταδιακές αλληλεπιδράσεις και να παρέχει διευκρινίσεις όπου απαιτείται [4].

Η χρήση νομικών chatbots προϋποθέτει αυξημένα μέτρα ασφάλειας δεδομένων, ειδικά σε εφαρμογές που περιλαμβάνουν ευαίσθητες πληροφορίες (π.χ. προσωπικά δεδομένα ή στοιχεία υποθέσεων). Εφαρμόζονται συχνά on-premise εγκαταστάσεις (on-premise deployments) με κρυπτογράφηση, ανωνυμοποίηση αρχείων καταγραφής (logs), και περιορισμένη πρόσβαση στα δεδομένα (access control) [72], [71].

Τα νομικά συστήματα συνομιλίας (legal chatbots) και οι εικονικοί βοηθοί (virtual assistants) αποτελούν κομβικό σημείο σύγκλισης μεταξύ τεχνητής νοημοσύνης (artificial intelli-

gence, AI) και νομικής πρακτικής (legal practice), ενσωματώνοντας τεχνικές αιχμής (state-of-the-art) στην επεξεργασία φυσικής γλώσσας (natural language processing, NLP), ανάκτηση πληροφοριών (retrieval), επεξηγησιμότητα (explainability) και ασφάλεια δεδομένων (data security). Για έναν μηχανικό, η πρόκληση δεν είναι μόνο η δημιουργία ενός λειτουργικού βοτ, αλλά και η διασφάλιση αξιοπιστίας (reliability), επεκτασιμότητας (scalability) και θεσμικής συμμόρφωσης (regulatory compliance) σε ένα από τα πιο απαιτητικά πεδία εφαρμογής της τεχνητής νοημοσύνης [72], [4], [71].

Οι εφαρμογές με διεπαφή προς τον πελάτη (client-facing applications) αποτελούν δυναμικό πεδίο της νομικής τεχνολογίας, ενσωματώνοντας προηγμένες τεχνικές AI, LLMs και RAG για αυτοματοποίηση της επικοινωνίας και της εξυπηρέτησης στον νομικό τομέα. Σε αντίθεση με τα παραδοσιακά συστήματα, οι σύγχρονες λύσεις αξιοποιούν vector search και τεχνικές εννοιολογικής συνάφειας για παροχή εξατομικευμένων απαντήσεων και αυτόματη παραγωγή νομικών εγγράφων. Κομβικό ρόλο διαδραματίζει ο συνδυασμός διαδραστικών ροών εργασίας (interactive workflows) και αυτοματοποιημένης επεξεργασίας δεδομένων, που επιτρέπουν τη δυναμική προσαρμογή στα χαρακτηριστικά κάθε χρήστη—όπως συμβαίνει σε case intake portals όπου το αίτημα του πελάτη αξιολογείται και δρομολογείται κατάλληλα.

Η γλωσσική παραγωγή (language generation) μέσω LLMs επιτρέπει τη σύνθεση απαντήσεων σε φυσική γλώσσα και την αυτόματη δημιουργία εγγράφων, όπως συμβάσεων και υπομνημάτων, είτε με templates είτε με εξατομικευμένες ροές. Πλατφόρμες όπως DocuSign, LegalZoom και νομικά portals υλοποιούν RAG pipelines για να διασφαλίζουν τεκμηριωμένες και αξιόπιστες παραδόσεις [71]. Η διανυσματική ανάκτηση (vector retrieval) με semantic embeddings ενισχύει την αναζήτηση σχετικών νομικών τεκμηρίων, ενώ η ενσωμάτωση μηχανισμών αυτόματης ανίχνευσης ελλείπων ή ανακριβών απαντήσεων αυξάνει την αξιοπιστία [43]. Κρίσιμες προκλήσεις αποτελούν η επαλήθευση (validation) και η ασφάλεια δεδομένων (data privacy, GDPR compliance), καθώς και η προστασία προσωπικών δεδομένων μέσω κρυπτογράφησης και ανωνυμοποίησης. Η διαλειτουργικότητα (interoperability) με DMS, συστήματα διαχείρισης υποθέσεων και πλατφόρμες e-signature διασφαλίζει ολιστική εμπειρία χρήσης, ενώ η διαρκής ενημέρωση των βάσεων γνώσης είναι αναγκαία για την προσαρμογή σε νομοθετικές αλλαγές.

Η αποτελεσματικότητα αυτών των εφαρμογών μετράται με δείκτες όπως ο βαθμός ικανοποίησης χρηστών, το ποσοστό αυτοματοποιημένης διεκπεραίωσης και η ακρίβεια νομικών απαντήσεων, υπό διαρκείς ελέγχους και audits [43], [72]. Τελικά, οι εφαρμογές με διεπαφή προς τον πελάτη μετασχηματίζουν το πεδίο των νομικών υπηρεσιών, προσφέροντας ταχύτητα, αυτοματοποίηση και διαφάνεια, υπό την προϋπόθεση αυστηρού τεχνικού σχεδιασμού, συνεχούς ποιοτικής παρακολούθησης και συμμόρφωσης με το νομοθετικό και δεοντολογικό πλαίσιο [71], [43], [72].

4.5.2 Ενσωμάτωση σε Συστήματα Διαχείρισης Νομικής Πρακτικής

Η σύγχρονη ενσωμάτωση τεχνολογιών τεχνητής νοημοσύνης και ειδικά μεγάλων γλωσσικών μοντέλων (LLMs) σε συστήματα διαχείρισης νομικής πρακτικής (Practice Management Systems – PMS) δεν περιορίζεται στην απλή αυτοματοποίηση εργασιών, αλλά εισάγει νέες αρχιτεκτονικές, διαλειτουργικότητα και βάθος ανάλυσης δεδομένων που μετασχηματίζουν

ριζικά τη λειτουργία των νομικών γραφείων. Οι πλατφόρμες αυτές υιοθετούν modular αρχιτεκτονική, επιτρέποντας την ευέλικτη διασύνδεση υποσυστημάτων, όπως η διαχείριση υποθέσεων, εγγράφων, πελατών και οικονομικών συναλλαγών, με προηγμένα υποσυστήματα επεξεργασίας φυσικής γλώσσας και μηχανικής μάθησης.

Σε τεχνικό επίπεδο, η διαχείριση υποθέσεων (case management) βασίζεται στη δυναμική σύνδεση μεταδεδομένων (τύπος υπόθεσης, προθεσμίες, πελάτες) με τεκμηριωμένα αποσπάσματα από τη βάση γνώσης του γραφείου. Η διαχείριση εγγράφων υλοποιείται με χρήση semantic search, indexing και version control, όπου κάθε έγγραφο μετατρέπεται σε διανυσματική αναπαράσταση (embedding) και συνδέεται με την τρέχουσα υπόθεση ή το πελατολόγιο, διευκολύνοντας τον εντοπισμό και την αξιολόγηση κρίσιμων πληροφοριών [72].

Η αυτοματοποίηση διαδικασιών επιτυγχάνεται μέσω ροών εργασίας (workflows) που ενεργοποιούνται είτε με βάση συγκεκριμένες ενέργειες (π.χ. εισαγωγή νέας υπόθεσης, ενημέρωση εγγράφου) είτε με συμβάντα που διακινούνται σε έναν δίαυλο συμβάντων (event bus), όπως τα Kafka, RabbitMQ. Τα μηνύματα που παράγονται, δρομολογούνται στα κατάλληλα υποσυστήματα (modules), όπως modules επεξεργασίας εγγράφων, καταγραφής ελέγχου (audit logging), ειδοποιήσεων και ανάλυσης δεδομένων (analytics), ενώ τα μεγάλα γλωσσικά μοντέλα (LLMs) καλούνται κατά απαίτηση (on-demand) να αναλάβουν tasks όπως αυτόματη σύνταξη εγγράφων, δημιουργία απαντήσεων σε πελάτες ή αξιολόγηση συμβολαίων [4], [71]. Σε αυτό το πλαίσιο, το pipeline Ανάκτησης-Ενισχυμένης Παραγωγής (RAG pipeline) συνδυάζει σημασιολογική ανάκτηση (semantic retrieval) με προσαρμοσμένο ζοντεξτ (contextual prompting), έτσι ώστε το LLM να απαντά σε ερωτήματα χρησιμοποιώντας τεκμηριωμένα αποσπάσματα από το νομικό αρχείο του γραφείου.

Η διαλειτουργικότητα επιτυγχάνεται με υλοποίηση RESTful ή GraphQL APIs, επιτρέποντας στα υποσυστήματα PMS να ανταλλάσσουν διασυνδέουν δεδομένα με εξωτερικά εργαλεία (όπως πίνακες ανάλυσης δεδομένων (analytics dashboards), πλατφόρμες κανονιστικής συμμόρφωσης (compliance platforms) ή υπηρεσίες cloud τρίτων (third-party cloud services)) και να ενσωματώνουν μεγάλα γλωσσικά μοντέλα (LLMs) που αναλαμβάνουν επιμέρους διεργασίες, όπως σύνοψη εγγράφων (document summarization), σύνταξη (drafting) και νομική αναζήτηση (legal search). Ένα σύγχρονο πρότυπο είναι η χρήση αρθρωτών ευφυών πρακτόρων (modular AI agents), όπου κάθε agent εξειδικεύεται σε διαφορετική θεματική ενότητα (π.χ. εργατικό δίκαιο, ανάλυση συμβολαίων), δέχεται ερωτήματα (queries) και επιστρέφει αποτελέσματα μέσω έξυπνης δρομολόγησης (intelligent routing), μεγιστοποιώντας έτσι την εξειδίκευση και την επεκτασιμότητα του συστήματος [71], [43].

Η ενσωμάτωση βάσεων διανυσμάτων (vector databases) — όπως τα FAISS, Pinecone, Qdrant — επιτρέπει τη διαρκή ενημέρωση των διανυσματικών αναπαραστάσεων (embeddings) όλων των εγγράφων και την πραγματοποίηση σημασιολογικής αναζήτησης (semantic search) όχι μόνο με βάση το πλήρες κείμενο, αλλά και με φίλτρα που λαμβάνουν υπόψη το συμφραζόμενο (context-aware filters), όπως ο τύπος υπόθεσης, ο πελάτης, η ημερομηνία ή το δικαστήριο. Αυτό επιτρέπει τη γρήγορη ανάκτηση, ομαδοποίηση (clustering) και προτεραιοποίηση εγγράφων για κάθε χρήση — από νομική έρευνα (legal research) μέχρι προετοιμασία για κατάθεση σε δικαστήριο (preparation for court filings) [4].

Σε επίπεδο διασφάλισης ποιότητας και αξιοπιστίας, η υιοθέτηση τεχνικών όπως έλεγχος πρόσβασης βάσει ρόλων (role-based access control), πλήρης κρυπτογράφηση δεδομένων (σε

κατάσταση αδράνειας και κατά τη μεταφορά — at rest & in transit), ιχνηλάτηση συμβάντων με δυνατότητα ελέγχου (auditable event trails), καθώς και αυτοματοποιημένη ανωνυμοποίηση των δεδομένων εκπαίδευσης για τα LLMs διασφαλίζει τη συμμόρφωση με κανονισμούς (compliance) όπως οι GDPR, CCPA και θεσμικά στάνταρδς [72]. Επιπλέον, κάθε αυτοματοποιημένη διαδικασία (π.χ. παραγωγή εγγράφου) συνοδεύεται από μηχανισμούς ιχνηλασιμότητας παραπομπών (citation tracing) και επεξηγησιμότητας (explainability modules), ώστε ο τελικός χρήστης να μπορεί να εντοπίσει το μονοπάτι συλλογισμού (reasoning path) και τις πηγές που χρησιμοποιήθηκαν για κάθε πρόταση ή απόφαση.

Η εμπειρική αξιολόγηση τέτοιων ολοκληρωμένων Practice Management Systems – PMS έχει αναδείξει σημαντικά οφέλη. Σύμφωνα με μελέτες, η υλοποίηση αυτοματοποιημένων ροών εργασίας (workflow automation) και σημασιολογικής ανάκτησης εγγράφων μειώνει κατά 40–60% τον χρόνο επεξεργασίας υποθέσεων και εγγράφων, ενώ αυξάνει το ποσοστό ορθότητας των αναφορών τιμολόγησης (billing reports) και συμμόρφωσης (compliance reports) σε επίπεδα άνω του 95% [71], [43]. Η χρήση αρθρωτών πρακτόρων LLM (modular LLM agents) έχει αποδειχθεί αποτελεσματική στη μείωση λαθών, στην ενίσχυση της διαφάνειας και στη βελτίωση της ικανοποίησης των χρηστών — καθώς οι δικηγόροι επικεντρώνονται πλέον σε ουσιαστικές νομικές εργασίες και όχι σε γραφειοκρατικές διαδικασίες.

Επιλεγμένες μελέτες περίπτωσης (case studies) δείχνουν ότι δικηγορικά γραφεία που εφάρμοσαν ολοκληρωμένη ενσωμάτωση LLMs και modules αυτοματοποίησης ροής εργασιών (workflow automation modules), πέτυχαν δραστική μείωση του χρόνου ελέγχου συμβολαίων και παραγωγής νομικών γνωμοδοτήσεων, αύξηση της συνέπειας (consistency) στα παραγόμενα αποτελέσματα και μείωση των σφαλμάτων συμμόρφωσης (compliance errors) σε πραγματικές ροές εργασίας [71], [43], [72]. Στον δημόσιο τομέα, η ενσωμάτωση τέτοιων πρακτικών είχε ως αποτέλεσμα τη μείωση του χρόνου ελέγχου κρατικών συμβάσεων κατά 52%, ενώ σε εκπαιδευτικά προγράμματα (law tech clinics) καταγράφηκε εντυπωσιακή βελτίωση στην ταχύτητα εκμάθησης των νέων εργαλείων από δικηγόρους και φοιτητές νομικής [71], [72].

4.6 Αξιολόγηση και Διασφάλιση Ποιότητας

Η συστηματική αξιολόγηση των συστημάτων τεχνητής νοημοσύνης και των LLMs στον νομικό τομέα αποτελεί βασική προϋπόθεση για την υιοθέτηση και την ασφαλή αξιοποίησή τους στην πράξη. Η μεθοδολογία της αξιολόγησης έχει εξελιχθεί σημαντικά τις τελευταίες δεκαετίες. Αρχικά, οι μετρικές επικεντρώνονταν σε παραδοσιακά σχήματα ακρίβειας (accuracy), ανάκλησης (recall), και F1-score κυρίως σε συστήματα ανάκτησης πληροφορίας (Information Retrieval). Η ανάπτυξη των word embeddings έφερε νέες μεθοδολογίες, όπως η μέτρηση ομοιότητας (cosine similarity) και τα tasks αναλογιών λέξεων, ενώ οι μετρικές διακρίθηκαν πλέον σε ενδογενείς (intrinsic) και εξωγενείς (extrinsic) [76].

Ωστόσο, με την έλευση των LLMs το πεδίο της αξιολόγησης μετατοπίστηκε δραστικά. Όπως επισημαίνεται και στο άρθρο Revisiting Word Embeddings in the LLM Era, οι παραδοσιακές μετρικές αποδείχθηκαν ανεπαρκείς για την αποτίμηση των πιο σύνθετων, σημασιολογικών και συμπραζόμενων ιδιοτήτων των LLMs [76]. Η αξιολόγηση σήμερα βασίζεται σε ένα πολυδιάστατο σύνολο μετρικών, όπως οι BERTScore, BLEU, ROUGE, Mean Recip-

rocal Rank (MRR), normalized Discounted Cumulative Gain (nDCG), αλλά και μετρικές όπως η τεκμηριωσιμότητα (faithfulness), ο δείκτης παραισθήσεων (hallucination rate), η ακρίβεια παραπομπών (citation accuracy), με έμφαση στην αποτίμηση της τεκμηρίωσης, της αξιοπιστίας και της συμμόρφωσης με το εκάστοτε πεδίο εφαρμογής (domain)..

Ειδικά στο νομικό πεδίο, η αξιολόγηση των LLMs αποκτά ακόμη μεγαλύτερη πολυπλοκότητα, καθώς δεν αρκεί η γλωσσική ακρίβεια — απαιτείται μέτρηση της νομικής ακρίβειας (legal accuracy), της εγκυρότητας και της συμμόρφωσης (compliance), της συνέπειας και της διαφάνειας (explainability). Η τεκμηρίωση των απαντήσεων, η ανίχνευση παρασιώσεων (hallucinations) και η επαληθευσιμότητα (verifiability) αποτελούν κρίσιμα στοιχεία για την πραγματική αποδοχή των legal AI tools στην πράξη. Για το λόγο αυτό, τα σύγχρονα συστήματα αξιολόγησης στηρίζονται πλέον σε benchmarking frameworks ειδικά προσαρμοσμένα για το δίκαιο (π.χ. LegalBench, LeCoDe, LegalRAG), που συνδυάζουν αυτοματοποιημένες μετρικές, ανθρώπινη επιβεβαίωση (human-in-the-loop) και πολυεπίπεδη ανάλυση ποιότητας [77], [43].

4.6.1 Μεθοδολογίες σύγκρισης επιδόσεων

Η αξιολόγηση και η συγκριτική αποτίμηση (benchmarking) συστημάτων τεχνητής νοημοσύνης στον νομικό τομέα έχει εξελιχθεί σε έναν ιδιαίτερα απαιτητικό, πολυπαραγοντικό κλάδο της επιστήμης. Η επιστημονική συζήτηση πλέον υπερβαίνει τα παραδοσιακά metrics, εστιάζοντας σε frameworks και ενδοτασείς που προσομοιώνουν την πολυπλοκότητα της πραγματικής νομικής πράξης, ενσωματώνοντας βαθιές τεχνικές μετρικές, μεθοδολογίες ανθρωποκεντρικής αξιολόγησης και συνεχή σύγκριση μεταξύ ετερογενών συστημάτων. Η εξέλιξη αυτή οφείλεται τόσο στην άνοδο των μεγάλων γλωσσικών μοντέλων (LLMs) όσο και στην εμφάνιση ειδικών αρχιτεκτονικών RAG, που συνδυάζουν αναζήτηση, τεκμηρίωση και γλωσσική παραγωγή σε πραγματικό χρόνο.

Η αξιολόγηση με το LegalBench-RAG είναι πολυεπίπεδη. Για τα tasks ανάκτησης (retrieval tasks) χρησιμοποιούνται μετρικές όπως η ανάκληση στην κορυφή ($\text{recall}@k$), ο μέσος αμοιβαίος βαθμός κατάταξης (Mean Reciprocal Rank – MRR), η κανονικοποιημένη αθροιστική απομείωση κέρδους (normalized Discounted Cumulative Gain – nDCG) και το ποσοστό επιτυχίας (hit rate), που αποτιμούν τόσο την πληρότητα όσο και την ιεράρχηση των ανακτημένων τεκμηρίων.

Στα tasks παραγωγής (generation tasks), η αξιολόγηση της τεκμηριωσιμότητας (faithfulness) γίνεται τόσο αυτόματα (μέσω επικάλυψης με τα ανακτημένα αποσπάσματα – overlap with retrieved passages), όσο και με πανελς νομικών που βαθμολογούν την ορθότητα και την επάρκεια των παραπομπών (citations). Η ακρίβεια παραπομπών (citation accuracy) αποτιμάται με αυτόματους ελεγκτές (automated checkers), ενώ η ανίχνευση παραισθήσεων (hallucination rate) γίνεται με αυτόματη διασταύρωση των απαντήσεων με τα gold standards.

Επιπλέον, εφαρμόζεται ιχνηλασιμότητα προέλευσης (provenance tracking) και εξαγωγή της διαδρομής συλλογισμού (extraction reasoning path), ώστε να ελέγχεται αν το μοντέλο μπορεί να τεκμηριώσει πλήρως το συλλογισμό του [78].

Παρά τα πλεονεκτήματα του LegalBench-RAG, παραμένουν συγκεκριμένα προβλήματα, κυρίως η εξάρτηση από επισημειωμένα (annotated) νομικά σώματα κειμένων (legal corpora)

που περιορίζονται συχνά σε αγγλοσαξονικές έννομες τάξεις. Η μετρική πληρότητας (faithfulness metric), όσο χρήσιμη κι αν είναι, δεν εξασφαλίζει απόλυτα τη νομική πληρότητα, ιδίως όταν τα ανακτημένα τεκμήρια είναι ελλιπή ή αμφιλεγόμενα. Τέλος, η ανάγκη για αξιολόγηση βάσει χρυσού προτύπου (gold standard evaluation) με χειροκίνητη επιμέλεια καθιστά την επέκταση του benchmarking σε άλλες γλώσσες (όπως τα ελληνικά) κοστοβόρα και χρονοβόρα.

Σύμφωνα με τα πειραματικά δεδομένα της βιβλιογραφίας[77], τα LLMs που εκπαιδεύονται με επισημειωμένους διαλόγους πολλαπλών γύρων (annotated multi-turn dialogues) αυξάνουν τη σχετικότητα και την πληρότητα των απαντήσεων κατά 24% σε σχέση με μοντέλα που έχουν εκπαιδευτεί μόνο σε δεδομένα ενός γύρου (single-turn data). Τα ποσοστά παραισθήσεων (hallucination rates) παραμένουν χαμηλά (7–9%) όταν επιβάλλεται απαίτηση για παραπομπές (citation requirement), αλλά αυξάνονται αισθητά σε ερωτήματα με περίπλοκα πραγματικά περιστατικά ή αμφισημία. Ένα κρίσιμο εύρημα του LeCoDe είναι ότι οι νομικοί αξιολογούν ως πιο σημαντική τη διαφάνεια (transparency), π.χ. την ιχνηλασιμότητα παραπομπών (citation provenance), σε σχέση με το ύφος ή τη ρητορική φλυαρία των LLMs. Ωστόσο, η αξιοπιστία της ανθρώπινης αξιολόγησης (human evaluation) εξακολουθεί να αμφισβητείται, καθώς παρουσιάζονται σημαντικές διαφορές στη βαθμολόγηση ανάμεσα σε έμπειρους νομικούς.

Για πιο εξειδικευμένες εργασίες, όπως η ανάλυση συμβολαίων και η ανίχνευση αντιφάσεων ή ασυνέπειας (contradictions or inconsistencies), τα benchmarks όπως το ContractNLI και το CUAD[79] αξιολογούν συνεπαγωγή (entailment), συνοχή (consistency) και εξαγωγή ρητρών (clause extraction) με μετρικές όπως ακρίβεια (accuracy), F1-score ανά ρήτρα, πληρότητα (faithfulness) και επαλήθευση υποστηρικτικών τεκμηρίων (supporting evidence trace). Τα αποτελέσματα δείχνουν ότι τα μοντέλα που έχουν εκπαιδευτεί αποκλειστικά σε νομικά δεδομένα παρουσιάζουν έως και 20% ανώτερη απόδοση σε σχέση με γενικά LLMs, ιδιαίτερα σε απαιτητικές εργασίες όπως η αναγνώριση ασαφών ή νομικά αμφιλεγόμενων ρητρών.

Επιπλέον, σύγχρονες μελέτες για την πρόβλεψη παραπομπών[80][78] και την πραγματικότητα (citation prediction and factuality, π.χ. SaulLM και LegalBench-RAG) δείχνουν ότι η προσαρμογή στο πεδίο (domain-specific fine-tuning) και η προσαρμοσμένη εκπαίδευση με ρητές οδηγίες (explicit instruction tuning) οδηγούν σε υψηλότερη ακρίβεια παραπομπών (citation accuracy) και πληρότητα (faithfulness), μειώνοντας τον κίνδυνο παραισθήσεων (hallucinations), αλλά ότι τα πλήρως αυτοματοποιημένα pipelines εξακολουθούν να απαιτούν παρέμβαση ειδικών (expert intervention) και επισημειωμένα σύνολα δεδομένων αναφοράς (gold-standard datasets) για πραγματικά αξιόπιστη αξιολόγηση.

Συνολικά, η εξέλιξη των μεθοδολογιών αξιολόγησης (benchmarking methodologies) στη νομική πληροφορική καθορίζεται από την ανάγκη πολυεπίπεδης, επιστημονικά εδραιωμένης και διαλειτουργικής αξιολόγησης. Ο συνδυασμός τεχνικών μετρικών (technical metrics), ανθρωποκεντρικής επιμέλειας (human-centric curation) και ανάπτυξης ανοικτών, πολυγλωσσικών συνόλων δεδομένων (open, multilingual datasets) διασφαλίζει τη βιωσιμότητα και τη θεσμική αποδοχή των εργαλείων τεχνητής νοημοσύνης για το δίκαιο (legal AI tools) σε κάθε δικαιοδοσία.

4.6.2 Ανίχνευση και Πρόληψη Παραισθήσεων

Η ανίχνευση και η πρόληψη παραισθήσεων (hallucinations) στα συστήματα τεχνητής νοημοσύνης, ιδίως στα μεγάλα γλωσσικά μοντέλα που εφαρμόζονται στο δίκαιο, αποτελεί θεμελιώδη επιστημονικό και τεχνολογικό στόχο. Οι παραισθήσεις συνίστανται στην παραγωγή εσφαλμένων, ανυπόστατων ή ανακριβώς τεκμηριωμένων απαντήσεων, οι οποίες δύνανται να εμφανιστούν είτε σε επίπεδο πραγματικών περιστατικών (factual hallucinations), είτε σε λανθασμένες νομικές παραπομπές (citation hallucinations), είτε σε εσφαλμένο συλλογισμό ή νομική ερμηνεία (reasoning hallucinations), ακόμα και σε αντιφατικές τοποθετήσεις εντός του ίδιου κειμένου (self-contradiction). Ειδικά στο νομικό πλαίσιο, η ύπαρξη τέτοιων λαθών ενέχει σοβαρούς κινδύνους, καθώς μπορεί να οδηγήσει σε λανθασμένες αποφάσεις, εσφαλμένες νομικές συμβουλές ή αποδόμηση της θεσμικής αξιοπιστίας του συστήματος.

Η σύγχρονη βιβλιογραφία [81][72] προχωρά σε σαφή κατηγοριοποίηση των παραισθήσεων (hallucinations) που εμφανίζουν τα συστήματα τεχνητής νοημοσύνης σε νομικές εφαρμογές. Πραγματολογικές παραισθήσεις (factual hallucinations) αφορούν την παραγωγή ανυπόστατων ή κατασκευασμένων γεγονότων που δεν προκύπτουν από το δοθέν εισαγωγικό κείμενο (input). Αντίστοιχα, παρατηρούνται παραισθήσεις παραπομπής (citation hallucinations), δηλαδή παραπομπές σε ανύπαρκτα, λανθασμένα ή ανακριδή άρθρα, αποφάσεις ή νομικά προηγούμενα (precedents) χωρίς πραγματική ύπαρξη. Επίσης, ιδιαίτερα σημαντικές είναι οι παραισθήσεις συλλογισμού (reasoning hallucinations), όπου το μοντέλο ακολουθεί λανθασμένη ή μη τεκμηριωμένη πορεία επιχειρηματολογίας, οδηγώντας σε εσφαλμένα συμπεράσματα. Τέλος, συναντώνται και απαντήσεις που εμπεριέχουν αντιφάσεις ή αυτοαντιφάσεις (contradiction/self-contradiction), με το μοντέλο να παράγει ασύμβατες ή αλληλοαναιρούμενες θέσεις στο ίδιο ή σε διαδοχικά αποσπάσματα.

Τα σύγχρονα συστήματα συνδυάζουν πολλαπλές τεχνικές για τη διαχείριση των παραισθήσεων[81]. Πρωτίστως, εφαρμόζουν αυτοματοποιημένη ανίχνευση αλληλοεπικάλυψης (automated overlap detection), ελέγχοντας αν οι παραγόμενες παραπομπές ή ισχυρισμοί υπάρχουν πράγματι στο εισαγωγικό πλαίσιο (input context) ή στο σύνολο ανάκτησης (retrieval set). Η ανίχνευση βασίζεται σε μετρικές επικάλυψης (overlap scores), συσχέτισης n-gram, καθώς και συσχέτισης εννοιών μέσω embeddings similarity ή μετρικών όπως η Attributable to Retrieved Content (ARC). Επιπλέον, αξιοποιούνται τεχνικές προσεγγιστικής ταύτισης και θεματικής αντιστοίχισης (fuzzy matching και domain-specific mapping), με χρήση αλγορίθμων εντοπισμού μικρών παραλλαγών ή παραφράσεων, ώστε να διαπιστώνεται η ύπαρξη πραγματικής νομικής βάσης (legal support) πίσω από κάθε παραπομπή. Παράλληλα, τα υποσυστήματα ελέγχου γεγονότων (fact-checking modules) εφαρμόζουν αυτόματη εξαγωγή νομικών ισχυρισμών και διασταύρωση με εξωτερικές βάσεις δεδομένων (external legal databases), ώστε να επαληθεύεται η ύπαρξη σχετικής νομολογίας ή νομοθεσίας μέσω τεχνικών εννοιολογικής συνεπαγωγής (entailment-based verification). Τέλος, ιδιαίτερα κρίσιμη παραμένει η ανάλυση σφαλμάτων με συμμετοχή ειδικών (human-in-the-loop error analysis), όπου εξειδικευμένοι νομικοί ελέγχουν λεπτές ή αμφιλεγόμενες περιπτώσεις που δεν ανιχνεύονται πάντα από τα αυτόματα συστήματα, όπως οι υπαινιγμοί εφαρμογής διατάξεων ή οι αμφιλεγόμενες νομολογιακές θέσεις.

Σε πειραματικά δεδομένα, η εφαρμογή υβριδικών πλαισίων (hybrid frameworks, δηλ.

αυτοματοποιημένη και ανθρώπινη ανίχνευση) αυξάνει σημαντικά την ανάκληση ανίχνευσης (recall) παραισθήσεων, μειώνει τα ψευδώς αρνητικά (false negatives) και προσφέρει διαφάνεια και αξιοπιστία, στοιχείο απαραίτητο για εφαρμογή σε δικαστικό ή συμβουλευτικό επίπεδο [78].

Η αποτελεσματική πρόληψη των παραισθήσεων (hallucination prevention) αποτελεί σύνθετη τεχνική και οργανωτική πρόκληση. Οι πιο αξιόπιστες στρατηγικές σήμερα περιλαμβάνουν την παραγωγή με ενισχυμένη ανάκτηση (retrieval-augmented generation, RAG), όπου οι απαντήσεις περιορίζονται σε τεκμηριωμένα τμήματα (segments) που έχουν ανακτηθεί από αξιόπιστες βάσεις, μειώνοντας σημαντικά τις παραισθήσεις όσο αυξάνεται η κάλυψη (coverage) του συστήματος. Παράλληλα, εφαρμόζεται περιορισμένη παραγωγή και υποχρεωτική παραπομπή (constrained generation, explicit citation requirement), μέσω κατάλληλων εντολών (prompts) που επιβάλλουν τεκμηρίωση και εμποδίζουν την παραγωγή απάντησης χωρίς έγκυρη πηγή. Επίσης, αξιοποιείται η τεχνική των εντολών αποχής και της εκτίμησης αβεβαιότητας (abstain prompting, uncertainty estimation), επιτρέποντας στο μοντέλο να «αρνηθεί» να απαντήσει σε περίπτωση ανεπαρκούς τεκμηρίωσης, αυξάνοντας έτσι την ακρίβεια (precision), ειδικά σε ασαφή νομικά ερωτήματα (ambiguous queries). Επιπλέον, οι κύκλοι ενεργής μάθησης και ανατροφοδότησης (active learning, feedback loops) με συμμετοχή χρηστών και επισημάνσεις παραισθήσεων ενισχύουν τη συνεχή βελτίωση μέσω επανεκπαίδευσης (retraining, fine-tuning) του μοντέλου.

Παρά την πρόοδο, εξακολουθούν να υφίστανται σημαντικές προκλήσεις. Τα benchmarks ανίχνευσης παραισθήσεων παραμένουν συχνά περιορισμένα ως προς τη γλώσσα και τη δικαιοδοσία, ενώ η ανθρώπινη υποκειμενικότητα δυσχεραίνει την τυποποίηση της αξιολόγησης. Πολλά LLMs προσαρμόζονται ώστε να αποφεύγουν εύκολα ανιχνεύσιμες παραισθήσεις, αλλά υστερούν στην αναγνώριση σύνθετων ή σπάνιων νομικών ερωτημάτων. Η δυσκολία ανίχνευσης επιτείνεται όταν τα νομικά δεδομένα παραμένουν κλειστά ή μη ανοιχτά (closed data). Έτσι, η επιστημονική συζήτηση μετατοπίζεται προς τη δημιουργία ανοιχτών, πολυγλωσσικών benchmarks παραισθήσεων, όπως το HalluLens και τα σύνολα δεδομένων του LegalBench-RAG, καθώς και προς τη διαρκή συνεργασία τεχνικών και νομικών για τη διασφάλιση της ασφάλειας και αξιοπιστίας στα νομικά συστήματα τεχνητής νοημοσύνης.

4.6.3 Συγκριτική Ανάλυση Απόδοσης

Η συγκριτική ανάλυση απόδοσης (comparative performance analysis) αποτελεί θεμέλιο της επιστημονικής αξιολόγησης των νομικών συστημάτων τεχνητής νοημοσύνης (legal AI systems). Δεν επαρκεί η απλή παράθεση τεχνικών αποτελεσμάτων· απαιτείται πολυεπίπεδη διερεύνηση που λαμβάνει υπόψη τα χαρακτηριστικά των μοντέλων, τα σύνολα δεδομένων (datasets), τις γλωσσικές και θεσμικές ιδιαιτερότητες, τις μετρικές (metrics), την αξιοπιστία της αξιολόγησης και, κυρίως, τη συμβατότητα των ευρημάτων με τις απαιτήσεις της νομικής πράξης.

Η απόδοση ενός νομικού συστήματος τεχνητής νοημοσύνης επηρεάζεται από πολλαπλούς παράγοντες: το είδος και την πολυπλοκότητα του έργου (task) —π.χ. ανάκτηση υποθέσεων (case retrieval), απάντηση νομικών ερωτημάτων (legal question answering), εξαγωγή όρων συμβολαίων (contract clause extraction)—, την προσαρμογή του LLM στο εκάστοτε

πεδίο (domain adaptation), τη γλωσσική επικάλυψη (language coverage), το είδος του συστήματος ανάκτησης (retriever: αραιό sparse, πυκνό dense, υβριδικό hybrid), την ύπαρξη δευτερεύοντος ταξινομητή (reranker), την απαίτηση παραπομπών (citation requirement) και την ύπαρξη επισημασμένου χρυσού προτύπου (annotated gold standard). Η διεθνής βιβλιογραφία, ειδικά τα αποτελέσματα του LegalBench-RAG, έχει τεκμηριώσει ότι δεν υπάρχει ενιαία βέλτιστη λύση, κάθε σύστημα επεξεργασίας (pipeline) παρουσιάζει διαφορετικά πλεονεκτήματα και μειονεκτήματα (strengths & weaknesses) ανάλογα με τη ρύθμιση και το εκάστοτε σενάριο χρήσης (use case).

Τα μοντέλα που έχουν υποστεί περαιτέρω εκπαίδευση (fine-tuning) σε νομικά σώματα κειμένων (legal corpora), όπως τα LegalBERT και Lawformer, αποδίδουν σταθερά καλύτερα ως προς την αληθοφάνεια (faithfulness) και την ακρίβεια παραπομπών (citation accuracy), ενώ τα γενικού σκοπού LLMs (general-purpose LLMs) εμφανίζουν αυξημένη ευελιξία σε ερωτήματα ελεύθερου κειμένου (free-text queries), αλλά και υψηλότερη πιθανότητα εμφάνισης παραισθήσεων (hallucinations) και λαθών στις παραπομπές (citation errors). Τα πυκνά συστήματα ανάκτησης (dense retrievers, π.χ. SPLADE, ColBERT) επιτυγχάνουν σημαντικά υψηλότερη ανάκληση στην (recall@10) σε tasks ανάκτησης, ωστόσο υστερούν σε ερμηνευσιμότητα (interpretability) σε σχέση με υβριδικές λύσεις (hybrid sparse+dense, π.χ. BM25+dense reranking), οι οποίες προσφέρουν καλύτερη ισορροπία ανάμεσα σε ανάκληση και επεξηγησιμότητα (recall και explainability) [43], [78].

Η σύγκριση μεταξύ εμπορικών διεπαφών προγραμματισμού (commercial APIs) και ανοιχτού κώδικα συστημάτων επεξεργασίας (open-source pipelines) αναδεικνύει τα πλεονεκτήματα και τα μειονεκτήματα κάθε προσέγγισης. Τα εμπορικά LLMs (commercial LLMs, π.χ. GPT-4, Claude) εμφανίζουν υψηλή απόδοση σε γενικά έργα, εξαιρετική γλωσσική φυσικότητα και πολυγλωσσική στήριξη, αλλά λειτουργούν ως μαύρα κουτιά (black boxes) με περιορισμένη επεξηγησιμότητα και διαφάνεια (transparency), καθώς και αβεβαιότητα ως προς την πηγή των παραπομπών (citation provenance). Αντίθετα, τα ακαδημαϊκά LLMs ανοιχτού κώδικα (open-source legal LLMs) παρέχουν πλήρη έλεγχο, υψηλή ερμηνευσιμότητα (interpretability) και δυνατότητα προσαρμογής (fine-tuning) σε εθνικά ή θεματικά σύνολα δεδομένων, αλλά συχνά υστερούν σε κάλυψη (coverage) —ιδίως για γλώσσες με λιγότερους πόρους (low-resource languages)— και σε εύρος νομικών θεμάτων.

Ιδιαίτερη προσοχή απαιτεί η συγκριτική ανάλυση μεταξύ διαφορετικών δικαιοδοσιών (jurisdictions). Η απόδοση των μοντέλων σε σύνολα δεδομένων όπως το GreekBarBench ή το SwiLTra-Bench αναδεικνύει τις εγγενείς προκλήσεις των LLMs σε γλώσσες με περιορισμένα δεδομένα, όπου η ακρίβεια (accuracy) και η αληθοφάνεια (faithfulness) μπορούν να μειωθούν έως και 20% σε σύγκριση με τα αγγλικά benchmarks, υπογραμμίζοντας την ανάγκη για τοπική τεχνολογική ανάπτυξη και εμπλουτισμό των χρυσών προτύπων (gold standard datasets) [5], [74].

Η αξιολόγηση με πραγματικά ερωτήματα (real-world queries), όπως η ανάκτηση υποθέσεων από ελληνικά δικαστήρια ή η πολυγλωσσική νομική μετάφραση στην Ελβετία, δείχνει ότι τα εμπορικά μοντέλα απαιτούν επιπλέον ενίσχυση (domain adaptation, feedback, human-in-the-loop), ενώ τα εξειδικευμένα νομικά μοντέλα (legal-tuned models) πρέπει να συμπληρώνονται από εξειδικευμένο σύστημα ανάκτησης, ενεργή μάθηση (active learning) και χειροκίνητη επιμέλεια (manual curation).

Η ουσιαστική συγκριτική ανάλυση απόδοσης αποτελεί το επιστημονικό εργαλείο που αποκαλύπτει τα πλεονεκτήματα και τα όρια κάθε τεχνολογικής προσέγγισης. Η αλυσίδα της απόδοσης δεν καθορίζεται μόνο από την ακρίβεια (accuracy) και την ανάκληση (recall), αλλά και από τη διαφάνεια (transparency), την επεξηγησιμότητα (explainability), τη νομική ανθεκτικότητα (legal robustness) και τη θεσμική αποδοχή (institutional acceptance). Η ορθολογική επιλογή και ενσωμάτωση των κατάλληλων συστημάτων επεξεργασίας (pipelines) περνά αναγκαστικά από πολυεπίπεδη, τεκμηριωμένη (evidence-based) και επιστημονικά εδραιωμένη συγκριτική αποτίμηση.

4.6.4 Μετρήσεις Νομικής Ακρίβειας

Η αξιολόγηση της νομικής ακρίβειας (legal accuracy) στα σύγχρονα συστήματα τεχνητής νοημοσύνης αποτελεί ίσως τη μεγαλύτερη επιστημονική πρόκληση στον χώρο της νομικής πληροφορικής. Δεν αρκεί η επιτυχία σε συμβατικές μετρικές κατηγοριοποίησης, καθώς κάθε απάντηση οφείλει να είναι πραγματολογικά ακριβής, τεκμηριωμένη ως προς τις παραπομπές, διαφανής στη λογική αλυσίδα που ακολουθεί και επαληθεύσιμη από ανεξάρτητους αξιολογητές. Τα πλέον προηγμένα benchmarks, όπως τα LegalBench-RAG, LeCoDe, GreekBarBench, έχουν εισαγάγει σύνθετες μετρήσεις, συνδυάζοντας την ακριβή ταύτιση (exact match), την ακρίβεια παραπομπών (citation accuracy), τη διαφάνεια συλλογιστικής (explainability) και την αληθοφάνεια (faithfulness) των απαντήσεων. Παρά την πρόοδο, η πραγματική νομική ορθότητα εξακολουθεί να είναι δύσκολο να τυποποιηθεί, καθώς επηρεάζεται από το εκάστοτε θεσμικό και πολιτισμικό πλαίσιο, αλλά και από τη δυνατότητα των συστημάτων να αντεπεξέρχονται σε ερωτήματα με πολλαπλές ή αμφισβητούμενες ερμηνείες. Η ερευνητική προσπάθεια επικεντρώνεται πλέον στη δημιουργία ανοικτών, διαλειτουργικών και θεσμικά έγκυρων μετρικών και συνόλων δεδομένων, ώστε να διασφαλίζεται όχι μόνο η τεχνική επάρκεια, αλλά και η ουσιαστική αποδοχή των συστημάτων τεχνητής νοημοσύνης στη νομική πράξη.

Η ουσιαστική αποτίμηση της νομικής ακρίβειας (legal accuracy) και της πρακτικής χρησιμότητας των συστημάτων τεχνητής νοημοσύνης στη νομική πράξη απαιτεί benchmarking όχι μόνο σε αυτόματες μετρικές, αλλά και σε ρεαλιστικά, πολυεπίπεδα σενάρια με σύγκριση προς τις πραγματικές επιδόσεις ανθρώπινων ειδικών. Ένα από τα πλέον σύγχρονα και απαιτητικά benchmarks στον χώρο αποτελεί το GreekBarBench[5], το οποίο αξιολογεί τις δυνατότητες των LLMs στην κατανόηση, ανάλυση και τεκμηρίωση ελληνικών δικαστικών αποφάσεων.

Το [5] στηρίζεται σε ένα ευρύ corpus αποφάσεων από πέντε διαφορετικούς τομείς του ελληνικού δικαίου (αστικό, ποινικό, εμπορικό, δημόσιο, κώδικας δικηγόρων) και υλοποιεί πολυδιάστατη αξιολόγηση σε τρία επίπεδα:

- α) ορθή απόδοση των πραγματικών γεγονότων (Facts)
- β) ακριβής ταυτοποίηση των νομοθετικών διατάξεων στις οποίες στηρίζεται η απόφαση (Cited Articles)
- γ) τεκμηριωμένη νομική ανάλυση (Analysis).

Η αξιολόγηση των αποτελεσμάτων γίνεται τόσο με αυτόματες μετρικές όσο και με τη συμμετοχή ομάδων νομικών ειδικών, επιτρέποντας άμεση σύγκριση με την ανθρώπινη απόδοση.

Στον σχήμα 4.7 παρουσιάζονται ο πίνακας με τα αναλυτικά αποτελέσματα για 13 διαφορετικά LLMs και ομάδες ανθρώπινων αξιολογητών. Τα δεδομένα αναδεικνύουν τη δυναμική πρόοδο των σύγχρονων LLMs: σε επίπεδο Facts, τα μοντέλα υψηλής απόδοσης (π.χ. Gemini-2.5-Flash, GPT-4.1, Claude-3.7-Sonnet) αγγίζουν ή ξεπερνούν τη μέση ανθρώπινη απόδοση, ενώ η απόδοση των Top-5% δικηγόρων παραμένει το ανώτατο όριο. Ωστόσο, στη διάσταση Cited Articles, ακόμη και τα πιο εξελιγμένα μοντέλα δυσκολεύονται σημαντικά, εμφανίζοντας μειωμένη ακρίβεια σε σχέση με τους ανθρώπινους ειδικούς – κάτι που επιβεβαιώνει τη σημασία της τεκμηριωμένης παραπομπής και της θεσμικής αληθοφάνειας.

Models	Civil Law				Criminal Law				Commercial Law				Public Law				Lawyers' Code				Overall			
	f	c	a	avg	f	c	a	avg	f	c	a	avg	f	c	a	avg	f	c	a	avg	f	c	a	avg
Top-5%	-	-	-	9.00	-	-	-	10.00	-	-	-	10.00	-	-	-	9.20	-	-	-	9.18	-	-	-	8.87
Experts	-	-	-	6.80	-	-	-	8.29	-	-	-	8.70	-	-	-	7.70	-	-	-	7.39	-	-	-	7.78
Gemini-2.5-F	8.8	8.4	8.4	8.53	8.4	7.9	8.5	8.28	8.6	8.0	8.3	8.27	8.7	8.0	8.1	8.28	8.9	8.5	8.5	8.62	8.7	8.2	8.4	8.40
GPT-4.1	8.8	8.1	8.4	8.44	8.5	8.2	8.2	8.28	8.4	8.1	8.3	8.27	8.6	7.7	8.1	8.14	8.8	8.2	8.4	8.48	8.6	8.0	8.3	8.32
Claude-3-7	8.5	7.2	7.4	7.72	8.2	6.9	7.0	7.37	7.6	6.9	7.4	7.31	8.6	7.2	7.6	7.79	8.5	8.2	8.2	8.29	8.3	7.3	7.5	7.71
GPT-4.1-mini	8.3	7.1	7.3	7.57	7.7	6.4	6.9	7.01	8.4	7.4	7.4	7.76	8.4	7.3	7.6	7.75	8.5	7.6	7.9	7.98	8.3	7.2	7.4	7.63
Gemma-3-27B	7.8	5.5	5.9	6.39	6.7	4.9	4.9	5.51	6.8	5.2	5.6	5.88	8.1	5.9	6.1	6.68	7.8	6.6	6.6	7.01	7.5	5.7	5.8	6.33
L-Krikri-8B	7.0	5.3	5.6	5.95	7.1	5.3	5.1	5.84	6.7	5.3	5.4	5.79	7.9	6.3	6.2	6.78	7.4	6.4	6.4	6.74	7.2	5.7	5.8	6.24

Σχήμα 4.7: Αναλυτικά αποτελέσματα αξιολόγησης νομικής ακρίβειας σε πραγματικές δικαστικές αποφάσεις, όπως αποτυπώνονται στο GreekBarBench για 13 LLMs και ομάδες ανθρώπινων αξιολογητών, ανά τομέα δικαίου και τύπο ερωτήματος. [5].

Τα παραπάνω ευρήματα τεκμηριώνουν τη σύνθετη φύση της νομικής ακρίβειας και δείχνουν ότι, παρά την ταχεία πρόοδο της τεχνητής νοημοσύνης, η πλήρης εξίσωση της μηχανικής με την ανθρώπινη κατανόηση και τεκμηρίωση αποτελεί ακόμη ανοιχτό επιστημονικό ζητούμενο. Η διαρκής επικαιροποίηση των συνόλων δεδομένων, η πολυεπίπεδη αξιολόγηση και η ενεργός εμπλοκή νομικών panels στην αξιολόγηση αποτελούν αναγκαίες προϋποθέσεις για την ασφαλή και υπεύθυνη υιοθέτηση των legal AI συστημάτων στη σύγχρονη νομική πράξη.

Η περαιτέρω ενίσχυση των εργαλείων τεχνητής νοημοσύνης στα νομικά συστήματα απαιτεί διαρκή επικαιροποίηση των δεδομένων, πολυεπίπεδη αξιολόγηση και ενεργή συνεργασία με νομικούς, ώστε τα μοντέλα να ανταποκρίνονται όχι μόνο σε τεχνικά, αλλά και σε ουσιαστικά θεσμικά και δεοντολογικά κριτήρια. Μόνο υπό αυτές τις προϋποθέσεις μπορεί να διασφαλιστεί ότι η τεχνητή νοημοσύνη θα λειτουργεί συμπληρωματικά – και όχι ανταγωνιστικά – προς την ανθρώπινη κρίση, διαμορφώνοντας ένα αξιόπιστο και θεσμικά αποδεκτό μέλλον για την εφαρμογή της στη νομική πράξη.

4.7 Συμπεράσματα

Τα σύγχρονα συστήματα νομικής τεχνητής νοημοσύνης αξιοποιούν μεγάλα γλωσσικά μοντέλα (LLMs), αρχιτεκτονικές ανάκτησης-ενισχυμένης παραγωγής (RAG) και προηγμένες τεχνικές ανάλυσης δεδομένων για να μετασχηματίσουν τον τρόπο λειτουργίας της δικαιοσύνης. Η ραγδαία πρόοδος στην αυτόματη νομική έρευνα, την πρόβλεψη αποφάσεων και την υποβοήθηση δικαστών και πολιτών συνοδεύεται από νέες προκλήσεις ως προς τη δια-

φάνεια, την αξιοπιστία, τη βαθμονόμηση των μοντέλων και τη διαχείριση παραισθήσεων (hallucinations).

Η διεθνής εμπειρία δείχνει ότι η επιτυχής ενσωμάτωση των τεχνολογιών αυτών δεν εξαρτάται μόνο από τους ίδιους τους αλγορίθμους, αλλά και από το θεσμικό πλαίσιο, την ανθρώπινη επίβλεψη και τη διαρκή αξιολόγηση της νομικής ακρίβειας και της ηθικής διάστασης κάθε λύσης. Πολυγλωσσικά και ανοιχτά benchmarks, διαλειτουργικότητα συστημάτων και συστηματική συνεργασία τεχνικών και νομικών ειδικών αναδεικνύονται ως θεμελιώδεις προϋποθέσεις για τη μετάβαση σε ένα μέλλον έξυπνης, διαφανούς και κοινωνικά αποδεκτής ψηφιακής δικαιοσύνης.

Η τεχνολογική υπεροχή δεν αρκεί από μόνη της. Η ουσιαστική πρόκληση είναι να εξασφαλιστεί η ισορροπία μεταξύ καινοτομίας, θεσμικής λογοδοσίας και διασφάλισης του κράτους δικαίου, ώστε τα εργαλεία της τεχνητής νοημοσύνης να λειτουργούν πραγματικά προς όφελος της δίκαιης και ορθής απονομής δικαιοσύνης στην ψηφιακή εποχή.

Ρυθμιστικό και ηθικό πλαίσιο

Η ραγδαία ενσωμάτωση της τεχνητής νοημοσύνης (AI) και των μεγάλων γλωσσικών μοντέλων (LLMs) στη νομική πράξη μεταβάλλει ριζικά τόσο τη λειτουργία της δικαιοσύνης όσο και τον ρόλο του νομικού επαγγελματία. Η μετάβαση από παραδοσιακές διαδικασίες σε ψηφιακές πλατφόρμες φέρνει στο προσκήνιο νέες προκλήσεις σχετικά με τη διαφάνεια, την εξηγησιμότητα, τη λογοδοσία και τον σεβασμό των θεμελιωδών δικαιωμάτων.

Σε αυτό το πλαίσιο, τίθενται κρίσιμα ερωτήματα για τη νομική ευθύνη σε περιπτώσεις σφαλμάτων, την αντιμετώπιση της αλγοριθμικής μεροληψίας, τη διασφάλιση της ιδιωτικότητας και τον ρόλο του ανθρώπινου παράγοντα στη λήψη αποφάσεων. Οι διεθνείς ρυθμιστικές πρωτοβουλίες, όπως ο Ευρωπαϊκός Κανονισμός για την Τεχνητή Νοημοσύνη[6], θέτουν πλέον σαφείς όρους συμμόρφωσης, διαφάνειας και αξιολόγησης κινδύνου.

Το παρόν κεφάλαιο εξετάζει το ρυθμιστικό και ηθικό υπόβαθρο της χρήσης τεχνητής νοημοσύνης στο δίκαιο, με έμφαση στη νομική ευθύνη, τη διαφάνεια, τα διεθνή ρυθμιστικά πρότυπα, την αλγοριθμική δεοντολογία και τις προϋποθέσεις για ουσιαστική συνεργασία ανθρώπου και αλγορίθμου στο δικαστικό οικοσύστημα.

5.1 Νομική Ευθύνη και Διαφάνεια

Η έννοια της νομικής διαφάνειας συνιστά θεμέλιο της λειτουργίας του δικαίου και του κράτους δικαίου, καθώς διασφαλίζει το δικαίωμα των πολιτών και των διαδίκων να γνωρίζουν όχι μόνο το αποτέλεσμα μιας δικαστικής ή διοικητικής απόφασης, αλλά κυρίως τη διαδικασία λήψης αυτής της απόφασης, τα θεμελιώδη της στοιχεία, τους εμπλεκόμενους παράγοντες και τα μέσα ελέγχου ή αμφισβήτησής της. Στο πλαίσιο της εισαγωγής τεχνητής νοημοσύνης και ιδιαίτερα των μεγάλων γλωσσικών μοντέλων στο δικαστικό και διοικητικό οικοσύστημα, η νομική διαφάνεια αποκτά νέα διάσταση και σημασία, καθώς η αλγοριθμική λήψη αποφάσεων ενέχει κινδύνους αδιαφάνειας, 'μαύρου κουτιού' (black box effect) και διασποράς ευθύνης [82].

Η νομική διαφάνεια διαφέρει από την τεχνική εξηγησιμότητα (explainability) ή τη διαφάνεια του αλγορίθμου ως προς τον σκοπό και το εύρος της: ενώ η τεχνική διαφάνεια αφορά την 'κατανόηση' των εσωτερικών λειτουργιών του συστήματος από τους προγραμματιστές ή τους ειδικούς, η νομική διαφάνεια εστιάζει στην υποχρέωση του θεσμικού συστήματος να παρέχει στους διαδίκους, τους δικαστές και το κοινωνικό σύνολο επαρκή πληροφόρηση για τα βασικά στοιχεία και τα κριτήρια που διαμόρφωσαν την απόφαση, έτσι ώστε να μπορε-

ί να ασκηθεί ουσιαστικός έλεγχος, να τεκμηριώνεται η λογοδοσία, και να διασφαλίζεται η δικαιοσύνη [82], [13], [83].

Τα βασικά θεσμικά εργαλεία που κατοχυρώνουν τη νομική διαφάνεια στα συστήματα τεχνητής νοημοσύνης είναι:

- Η υποχρέωση τεκμηρίωσης και αιτιολόγησης κάθε αλγοριθμικής απόφασης (όπως απαιτείται από τον AI Act και τη δικονομική νομοθεσία).
- Η υποχρέωση ενημέρωσης των διαδίκων ότι χρησιμοποιείται αλγοριθμικό σύστημα στη διαδικασία λήψης απόφασης.
- Η δυνατότητα challenge και επανεξέτασης της απόφασης, με παροχή επαρκών στοιχείων για το πώς οδηγήθηκε το σύστημα στο τελικό αποτέλεσμα [82], [84], [85].
- Η υιοθέτηση μέτρων ελεγκτικότητας (auditability) και η διατήρηση διαδρομών ελέγχου (audit trails), ώστε να καθίσταται δυνατή η αναδρομική εξέταση της διαδικασίας λήψης αποφάσεων [13].

Η προστασία θεμελιωδών δικαιωμάτων, όπως η δικαστική προστασία, η ισότητα των όπλων (equality of arms) και το δικαίωμα σε δίκαιη δίκη (due process), προϋποθέτει την ύπαρξη ουσιαστικής διαφάνειας σε κάθε στάδιο της διαδικασίας. Σε περιβάλλοντα όπου η απόφαση λαμβάνεται ή επηρεάζεται από τα μεγάλα γλωσσικά μοντέλα, η αδυναμία κατανόησης ή αμφισβήτησης της αλγοριθμικής λογικής μπορεί να οδηγήσει σε ουσιαστική αδυναμία προσβολής της απόφασης και να υπονομεύσει τη νομιμότητα και την εμπιστοσύνη στη δικαιοσύνη [13], [83].

Η διεθνής εμπειρία, όπως στην υπόθεση COMPAS[86] στις Η.Π.Α, καταδεικνύει ότι ο περιορισμένος βαθμός διαφάνειας των αλγορίθμων ή η λειτουργία συστημάτων ως μαύρα κουτιά (black boxes) επιφέρει σοβαρά προβλήματα λογοδοσίας, περιορίζει το δικαίωμα προσβολής της απόφασης και διακυβεύει τη δίκαιη μεταχείριση των διαδίκων [83]. Αντίθετα, συστήματα που ενσωματώνουν δυνατότητες εξήγησης, διατήρηση audit trails και θεσμικά κατοχυρωμένο δικαίωμα αμφισβήτησης (challenge) συμβάλλουν ουσιαστικά στην ενίσχυση της νομιμότητας και της κοινωνικής αποδοχής [82], [85], [83].

Η διασφάλιση της διαφάνειας στη νομική χρήση των LLMs δεν μπορεί να στηριχθεί αποκλειστικά σε τεχνικές λύσεις, αλλά απαιτεί συνδυασμό θεσμικών εγγυήσεων, κανονιστικών απαιτήσεων και πρακτικών αυδιτινγ. Ο AI Act [6] και άλλες σύγχρονες ρυθμίσεις εισάγουν την υποχρέωση για διαφάνεια στην ενημέρωση, στην αιτιολόγηση και στη δυνατότητα ελέγχου των αυτοματοποιημένων αποφάσεων (άρθρα 13–15 AI Act, GDPR art. 22). Παράλληλα, η τήρηση πλήρους και προσβάσιμου ιστορικού ελέγχου (audit trail) για κάθε βήμα της αλγοριθμικής διαδικασίας καθίσταται κρίσιμη για την αποκατάσταση της αλυσίδας απόφασης και την απόδοση ευθύνης (accountability) όπου απαιτείται.

Εν τέλει, η κατοχύρωση ενός διαφανούς (transparency) και επαληθεύσιμου πλαισίου λειτουργίας των μεγάλων γλωσσικών μοντέλων (LLMs) στη δικαιοσύνη διασφαλίζει τη δυνατότητα ελέγχου (auditability), αμφισβήτησης (challenge) και επανεξέτασης (review) των αποφάσεων, προστατεύει τα δικαιώματα των διαδίκων και ενισχύει την εμπιστοσύνη της κοινωνίας στην τεχνητή νοημοσύνη ως εργαλείο απονομής δικαιοσύνης. Η διαρκής θεσμική

εργήγορη, η ενσωμάτωση τεχνικών εξηγησιμότητας (explainability) και ελέγξιμότητας (auditability), καθώς και η ρητή αναγνώριση του δικαιώματος αμφισβήτησης (challenge), αποτελούν τα εχέγγυα για ουσιαστική και αποτελεσματική νομική διαφάνεια (transparency) στο ψηφιακό νομικό οικοσύστημα

5.1.1 Απαιτήσεις Επεξηγησιμότητας

Η εξηγησιμότητα (explainability) στα LLMs προσεγγίζεται ως πολυεπίπεδος στόχος που συνδυάζει τεχνικές, νομικές και θεσμικές απαιτήσεις. Σε τεχνικό επίπεδο, η κύρια δυσκολία πηγάζει από τη φύση των μοντέλων μαύρου κουτιού, τα οποία εκπαιδεύονται σε τεράστια σύνολα δεδομένων (datasets), με δισεκατομμύρια παραμέτρους, σε διαδικασίες που είναι από τη βάση τους μη γραμμικές και δύσκολα αναγώγιμες σε ανθρώπινη λογική.

Οι βασικές μέθοδοι που χρησιμοποιούνται για την εξηγησιμότητα (explainability) στα μεγάλα γλωσσικά μοντέλα (LLMs) αποσκοπούν στη βελτίωση της διαφάνειας και της ερμηνευσιμότητας της λειτουργίας τους, ιδίως σε κρίσιμα νομικά περιβάλλοντα. Μια από τις πιο διαδεδομένες τεχνικές είναι η οπτικοποίηση της προσοχής (attention visualization), όπου τα βάρη προσοχής (attention weights) που χρησιμοποιούν τα transformer-based LLMs παρέχουν ενδείξεις για το ποιες λέξεις ή φράσεις προσέχει το μοντέλο κατά την παραγωγή απαντήσεων. Παρότι προσφέρουν μια πρώτη εικόνα για τη διαδικασία, τα attention patterns δεν ταυτίζονται απαραίτητα με τη λογική του ανθρώπου και η ερμηνεία τους παραμένει αμφιλεγόμενη [82]. Παράλληλα, αξιοποιούνται προσεγγιστικά μοντέλα (surrogate models) όπως τα LIME και SHAP, τα οποία λειτουργούν ως τοπικά ερμηνεύσιμα υποδείγματα και προσεγγίζουν τη συμπεριφορά του LLM για συγκεκριμένες εισόδους (inputs), προσφέροντας εξηγήσεις που όμως δεν καλύπτουν το σύνολο της λειτουργίας του μοντέλου.

Η σημαντικότητα χαρακτηριστικών και οι χάρτες επιρροής (feature importance και saliency maps) αποτελούν επίσης βασικές προσεγγίσεις, καθώς αξιολογούν πώς συγκεκριμένες λέξεις ή χαρακτηριστικά συμβάλλουν στην τελική πρόβλεψη, μέσω μεθόδων οπισθοδιάδοσης (backpropagation) ή διαβάθμισης (gradient-based approaches). Επιπλέον, η ανάλυση αντιπαραδειγμάτων (counterfactual analysis) εξετάζει πώς θα άλλαζε η έξοδος του μοντέλου αν τροποποιούνταν ορισμένα inputs, συμβάλλοντας στη διάγνωση προκαταλήψεων (bias) και μη προφανών εξαρτήσεων.

Τέλος, ιδιαίτερα σημαντικό ρόλο διαδραματίζει η καταγραφή του ιστορικού ελέγχου και της προέλευσης δεδομένων (audit trails και data provenance), δηλαδή η διατήρηση πλήρους ιστορικού των δεδομένων εισόδου, των αποφάσεων και των εσωτερικών καταστάσεων του μοντέλου. Παρά το υψηλό υπολογιστικό κόστος, αυτή η τεχνική είναι καθοριστική για τη διασφάλιση της συμμόρφωσης (compliance) σε απαιτητικές νομικές εφαρμογές [85].

Η εξηγησιμότητα (explainability) στα Μεγάλα Γλωσσικά Μοντέλα παραμένει σε μεγάλο βαθμό προσχηματική. Τα παραγόμενα επεξηγηματικά κείμενα είναι συχνά μεταγενέστερα (post hoc) και δεν αποκαλύπτουν τον πραγματικό τρόπο που το μοντέλο καταλήγει στο τελικό παράδειγμα. Για παράδειγμα, νομικές αποφάσεις που βασίζονται σε LLMs (π.χ. προτάσεις ποινής ή εκτίμησης κινδύνου (risk assessments) ενδέχεται να ενσωματώνουν προκατάληψη (bias) ή αυθαίρετες αποφάσεις (arbitrary decisions), χωρίς να παρέχεται δυνατότητα εντοπισμού της αιτίας.

Υπάρχουν τεκμηριωμένες αποτυχίες:

- Η υπόθεση COMPAS (Η.Π.Α) ανέδειξε τη συστημική προκατάληψη (systemic bias) που μπορεί να εμφανιστεί σε εργαλεία εκτίμησης κινδύνου (risk assessment tools), τα οποία λειτουργούν χωρίς δυνατότητα ουσιαστικής εξήγησης των αποφάσεων τους [85].
- Σε δικαστικά συστήματα της Κίνας, η εξηγησιμότητα περιορίζεται στην επίσημη παρουσίαση αποτελεσμάτων χωρίς πλήρη αποτύπωση της αλγοριθμικής λογικής.

Ο Κανονισμός για την Τεχνητή Νοημοσύνη της Ευρωπαϊκής Ένωσης [6] επιβάλλει ρητές απαιτήσεις (explicit requirements) για την εξηγησιμότητα (explainability) των συστημάτων υψηλού κινδύνου (high-risk systems). Συγκεκριμένα, προβλέπεται η τεκμηρίωση των λογικών βημάτων λήψης απόφασης, η παραγωγή εξόδου κατανοητής από τον άνθρωπο (human-understandable output), η τήρηση πλήρους ιστορικού ενεργειών (audit trails) καθώς και η δυνατότητα αμφισβήτησης των αποφάσεων.

Η απαίτηση αυτή προϋποθέτει την τεχνική ενσωμάτωση μηχανισμών παρακολούθησης και εξηγήσιμης τεκμηρίωσης σε κάθε βήμα της επεξεργασίας[6]. Για παράδειγμα, τα audit trails δεν αποτελούν απλώς αρχεία καταγραφής, αλλά πρέπει να τεκμηριώνουν με ακρίβεια: (α) τα δεδομένα εισόδου, (β) τα ενδιάμεσα στάδια ανάκτησης ή πρόβλεψης, (γ) τις επιλογές του μοντέλου και (δ) το τελικό αποτέλεσμα. Επιπλέον, η εξηγησιμότητα σε νομικά πλαίσια απαιτεί τεκμηρίωση με παραπομπές σε νομικές πηγές, αλλά και λογική ανάλυση που να μπορεί να ελεγχθεί από τον διάδικο ή το δικαστήριο.

Η εκπλήρωση αυτών των απαιτήσεων δεν είναι απλώς τεχνική, αλλά συνιστά νομική υποχρέωση βάσει του AI Act. Συνεπώς, η εφαρμογή τεχνικών εργαλείων όπως οι μηχανισμοί retrieval logging, η διατήρηση πλήρους αποτύπωσης κάθε ενέργειας (traceability), και η δυνατότητα ερμηνεύσιμης εξόδου μέσω μεταμοντελικών τεχνικών (post-hoc explainability methods) αποτελούν κρίσιμους άξονες συμμόρφωσης με το θεσμικό πλαίσιο. Σε αντίθετη περίπτωση, η χρήση LLMs σε δικαστικά περιβάλλοντα ενδέχεται να παραβιάζει βασικές αρχές δικαιοσύνης και διαφάνειας.

Η χρήση συστημάτων που λειτουργούν ως μαύρα κουτιά, χωρίς επαρκή δυνατότητα αιτιολόγησης ή επαλήθευσης της λογικής ακολουθίας πίσω από την εκάστοτε απόφαση, έχει προκαλέσει έντονες επιφυλάξεις στη σύγχρονη νομική θεωρία και νομολογία. Ιδίως στο ευρωπαϊκό πλαίσιο προστασίας θεμελιωδών δικαιωμάτων, διατυπώνεται η άποψη ότι η απουσία διαφάνειας ενδέχεται να συνιστά παραβίαση του δικαιώματος σε δίκαιη δίκη, όπως κατοχυρώνεται στο άρθρο 6 της ΕΣΔΑ, όταν η απόφαση επηρεάζει έννομα συμφέροντα και δεν παρέχεται η δυνατότητα κατανόησης ή αμφισβήτησης της αιτιολόγησής της. Η μεταφορά της έννοιας της αιτιολόγησης από το ανθρώπινο στο αλγοριθμικό περιβάλλον αποτελεί ανοιχτό ερμηνευτικό και θεσμικό ζήτημα, το οποίο αναμένεται να τεθεί ενώπιον των ευρωπαϊκών δικαστηρίων σε επόμενες υποθέσεις.

5.1.2 Μοντέλα Λογοδοσίας

Η διαμόρφωση αποτελεσματικών πλαισίων λογοδοσίας (accountability frameworks) για τα συστήματα τεχνητής νοημοσύνης, και ειδικότερα για τα μεγάλα γλωσσικά μοντέλα (LLMs),

αποτελεί μία από τις μεγαλύτερες προκλήσεις στη σύγχρονη τεχνολογική και νομική πραγματικότητα. Η πολυπλοκότητα, η αδιαφάνεια (opacity) των αλγοριθμικών αποφάσεων και η πληθώρα φορέων που εμπλέκονται σε κάθε στάδιο ανάπτυξης, υλοποίησης και χρήσης των συστημάτων αυτών καθιστούν τη σαφή κατανομή ευθύνης ιδιαίτερα δύσκολη. Ως αποτέλεσμα, κάθε δικαιοδοσία ακολουθεί διαφορετική προσέγγιση, ανάλογα με το θεσμικό της πλαίσιο, τη σχέση κράτους-αγοράς, και το επίπεδο προστασίας που επιδιώκει για τους πολίτες και τους οργανισμούς της.

Η σύγκριση των ρυθμιστικών προσεγγίσεων της Ευρωπαϊκής Ένωσης, των Ηνωμένων Πολιτειών Αμερικής και της Κίνας αναδεικνύει τις ποικίλες εκφάνσεις και προκλήσεις της λογοδοσίας στην εποχή της τεχνητής νοημοσύνης. Κάθε σύστημα προβάλλει διαφορετικές προτεραιότητες και ενσωματώνει διακριτές δικλίδες ασφαλείας, με αποτέλεσμα να προκύπτουν σημαντικές αποκλίσεις τόσο ως προς την προστασία των δικαιωμάτων των χρηστών όσο και ως προς την αποτελεσματικότητα των μηχανισμών ελέγχου και απονομής ευθύνης.

Ο νέος ευρωπαϊκός κανονισμός για την τεχνητή νοημοσύνη εισάγει ίσως το πιο αυστηρό και λεπτομερές πλαίσιο λογοδοσίας παγκοσμίως[6]. Προβλέπει υποχρεωτική τεκμηρίωση (documentation) και δυνατότητα ελέγχου (auditability) σε όλα τα στάδια του κύκλου ζωής του συστήματος, καθώς και συνεχή παρακολούθηση και αξιολόγηση κινδύνων και επιδόσεων. Η ανθρώπινη εποπτεία (human oversight) αποτελεί βασική αρχή, με τον άνθρωπο να διατηρεί τον τελικό λόγο στη λήψη αποφάσεων (the “last word”). Ο πάροχος του συστήματος τεχνητής νοημοσύνης είναι υπεύθυνος για τη διαδικασία εκπαίδευσης (training), την παρακολούθηση (monitoring) και την τεκμηρίωση (documentation), ενώ στους οργανισμούς-φορείς υλοποίησης ανατίθεται η ευθύνη για την επιλογή κατάλληλου συστήματος, την προσαρμογή στις ειδικές ανάγκες του πεδίου εφαρμογής (domain adaptation) και την τακτική αξιολόγηση των σχετικών κινδύνων. Η ευθύνη (liability) συχνά κατανέμεται μεταξύ των εμπλεκόμενων μερών, ιδίως σε σενάρια όπου ο άνθρωπος συμμετέχει ενεργά στη λήψη αποφάσεων (human-in-the-loop). Σε αυτές τις περιπτώσεις, η ευθύνη επιμερίζεται μεταξύ παρόχου, φορέα υλοποίησης (deployer) και τελικού λήπτη της απόφασης (decision-maker), απαιτώντας σαφή κατανομή αρμοδιοτήτων μέσω συμβάσεων, εσωτερικών πολιτικών και διαδικασιών αναφοράς (reporting).

Το ρυθμιστικό πλαίσιο για την τεχνητή νοημοσύνη στις ΗΠΑ χαρακτηρίζεται από κατακερματισμό, με έμφαση στην ευθύνη προϊόντος (product liability) και τη νομολογία (case law). Οι πάροχοι μεγάλων γλωσσικών μοντέλων ή εργαλείων τεχνητής νοημοσύνης (LLMs/AI tools) μπορούν να θεωρηθούν υπεύθυνοι για ζημιές (damages) εάν αποδειχθεί ότι το σύστημα παρουσίαζε ελαττώματα ή προκαταλήψεις (bias) που θα μπορούσαν να είχαν εντοπιστεί ή προληφθεί μέσω επαρκούς δοκιμής. Παράλληλα, οι χρήστες – όπως δικαστές – τεκμαίρεται ότι δρουν ορθά, ωστόσο, σε περιπτώσεις δυσμενούς έκβασης, συχνά ακολουθούνται συλλογικές αγωγές (class actions) κατά του παρόχου, κυρίως για ζητήματα προκατάληψης ή ανεπαρκούς διαφάνειας (failure in transparency). Η απουσία ενιαίου ομοσπονδιακού πλαισίου ευθύνης επιτρέπει το φαινόμενο του forum shopping, δηλαδή την επιλογή δικαστικής περιφέρειας με πιο ευνοϊκούς όρους για τον ενάγοντα.

Το ρυθμιστικό πλαίσιο της Κίνας δίνει έμφαση στον κρατικό έλεγχο (state auditing) και τη διασφάλιση λογοδοσίας κυρίως προς το κράτος, παρά προς τον πολίτη. Το Basic Law of Artificial Intelligence επιβάλλει υποχρεωτική λογοδοσία και ιχνηλασιμότητα (accountability

and traceability), ωστόσο η διαφάνεια προς τον πολίτη παραμένει συχνά περιορισμένη[73].

Σε περιπτώσεις αποτυχίας (π.χ. προκατειλημμένο αποτέλεσμα, λανθασμένη συμβουλή), η αλυσίδα ευθύνης διασπάται μεταξύ δημιουργού του συστήματος (developer), φορέα υλοποίησης (deployer) και τελικού χρήστη (user), ενώ συχνά ανακύπτουν ζητήματα αποκλίνουσας ερμηνείας συμβολαίων, αποζημιώσεων (insurance claims) και συγκρούσεων δικαίου (conflict of laws). Στα υβριδικά σχήματα όπου ο άνθρωπος συνεργάζεται με το μοντέλο (LLM+human), είναι κρίσιμη η καταγραφή του ποιος είχε πρόσβαση σε ποια δεδομένα, ποιοι κίνδυνοι/περιορισμοί διατυπώθηκαν, και αν υπήρξε η κατάλληλη εκπαίδευση των χρηστών.

Για να αποτυπωθούν με συγκριτικό και συστηματικό τρόπο οι διαφορές ανάμεσα στα βασικά ρυθμιστικά πλαίσια, ο παρακάτω πίνακας συνοψίζει τις βασικές παραμέτρους τεκμηρίωσης, εξηγησιμότητας, ελέγχου, ανθρώπινης εποπτείας, κατανομής ευθύνης και επιπέδου συμμόρφωσης στα μεγάλα γλωσσικά μοντέλα (LLMs). Μέσα από αυτή τη σύγκριση γίνεται φανερό ότι η Ευρωπαϊκή Ένωση ακολουθεί προσέγγιση υψηλής θεσμικής προστασίας και ρητής λογοδοσίας, οι Ηνωμένες Πολιτείες επικεντρώνονται στην ευθύνη προϊόντος και τη νομολογιακή αντιμετώπιση, ενώ το κινεζικό πλαίσιο εστιάζει κυρίως στον κρατικό έλεγχο και στη λογοδοσία προς τις δημόσιες αρχές.

5.2 Ρυθμιστικό Τοπίο

Η ανάπτυξη και εφαρμογή συστημάτων τεχνητής νοημοσύνης στον νομικό τομέα εξελίσσεται σε ένα ολοένα και πιο σύνθετο ρυθμιστικό περιβάλλον, όπου τα εθνικά και υπερεθνικά πλαίσια συχνά αλληλεπικαλύπτονται ή και συγκρούονται. Η διαμόρφωση κανόνων για τη λογοδοσία, τη διαφάνεια και την προστασία θεμελιωδών δικαιωμάτων αποτελεί πλέον κρίσιμο διακύβευμα για τους εμπλεκόμενους φορείς—από τους προγραμματιστές και τους παρόχους μέχρι τους τελικούς χρήστες και τους θεσμικούς ελεγκτές.

Στη διεθνή σκηνή, συναντάμε αφενός πρωτοποριακά θεσμικά κείμενα όπως ο Ευρωπαϊκός Κανονισμός Τεχνητής Νοημοσύνης και αφετέρου εθνικές προσεγγίσεις που διαφέρουν σημαντικά ως προς τη φιλοσοφία, την εστίαση και το επίπεδο συμμόρφωσης. Ο ρυθμιστικός ανταγωνισμός μεταξύ Η.Π.Α, Ε.Ε και Κίνας καθορίζει την παγκόσμια αγορά νομικής τεχνητής νοημοσύνης, διαμορφώνοντας πρότυπα διαφάνειας, εξηγησιμότητας και λογοδοσίας με άμεσο αντίκτυπο στη νομική ασφάλεια, τα δικαιώματα των πολιτών και την τεχνική καινοτομία [84], [85].

Η υιοθέτηση αυστηρότερων πλαισίων (frameworks), όπως στην Ευρωπαϊκή Ένωση, ή περισσότερο ευέλικτων και αποκεντρωμένων προσεγγίσεων (decentralized approaches), όπως στις Ηνωμένες Πολιτείες, δημιουργεί όχι μόνο τεχνολογικές προκλήσεις για τους προγραμματιστές, αλλά και κρίσιμα νομικά διλήμματα σχετικά με τη διασυνοριακή λειτουργία (cross-border operation) των μεγάλων γλωσσικών μοντέλων (LLMs), τη συμμόρφωση με αντικρουόμενα πρότυπα (standards) και την προστασία των θεμελιωδών δικαιωμάτων (fundamental rights). Ταυτόχρονα, αναδύονται νέα ρυθμιστικά πεδία, όπως η πρόβλεψη για ρυθμιστικούς «χώρους δοκιμών» (regulatory sandboxes), οι διαδικασίες αξιολόγησης επιπτώσεων (impact assessment) και οι δυναμικές αρχιτεκτονικές ελέγχου (dynamic control architectures) που ανταποκρίνονται στη φύση των προσαρμοστικών συστημάτων τεχνητής νοημοσύνης (adaptive AI systems)

5.2.1 Ο Κανονισμός Τεχνητής Νοημοσύνης της ΕΕ και η Νομική Τεχνητή Νοημοσύνη

Η υιοθέτηση του Κανονισμού Τεχνητής Νοημοσύνης της Ευρωπαϊκής Ένωσης [6], που οριστικοποιήθηκε το 2024, σηματοδοτεί μια θεμελιώδη τομή στη ρύθμιση των συστημάτων τεχνητής νοημοσύνης με άμεση και μακροπρόθεσμη επίδραση στο νομικό οικοσύστημα, ειδικά ως προς τη χρήση μεγάλων γλωσσικών μοντέλων (LLMs) για σκοπούς δικαστικής ή διοικητικής λήψης αποφάσεων. Η ευρωπαϊκή προσέγγιση, όπως αποτυπώνεται στον AI Act, είναι χαρακτηριστικά ανθρωποκεντρική και βασισμένη στην αρχή της προφύλαξης των θεμελιωδών δικαιωμάτων (fundamental rights). Εστιάζει συστηματικά στη διασφάλιση της διαφάνειας (transparency), της λογοδοσίας (accountability) και της δυνατότητας ελέγχου (auditability) σε κάθε στάδιο ανάπτυξης και χρήσης συστημάτων ΤΝ, ενώ υιοθετεί μία διαβαθμισμένη λογική αξιολόγησης κινδύνου (risk-based approach), η οποία αντικατοπτρίζει τη δυναμική και την πολυπλοκότητα των αλγοριθμικών τεχνολογιών στο πεδίο του δικαίου[84].

Σύμφωνα με αυτή τη νέα θεσμική αρχιτεκτονική, τα συστήματα τεχνητής νοημοσύνης κατηγοριοποιούνται βάσει του ρυθμιστικού και κοινωνικού τους κινδύνου (regulatory and social risk), διακρίνοντας ρητά ανάμεσα σε εφαρμογές απαράδεκτου κινδύνου (π.χ. κοινωνική βαθμολόγηση ή μη διαφανής επιτήρηση, unacceptable risk applications) και σε εκείνες που χαρακτηρίζονται ως «υψηλού ρίσκου» (high-risk applications). Η δεύτερη κατηγορία αφορά κυρίως συστήματα ΤΝ που ενσωματώνονται σε θεσμικές λειτουργίες με καθοριστικό αντίκτυπο στα ανθρώπινα δικαιώματα, όπως η απονομή δικαιοσύνης, η υποστήριξη λήψης διοικητικών αποφάσεων, η αυτοματοποιημένη επεξεργασία νομικών δεδομένων και οι προτάσεις επί δικαστικών φακέλων. Τα LLMs που αναπτύσσονται ή υλοποιούνται για τέτοιους σκοπούς, ανεξαρτήτως αν λειτουργούν ως εργαλεία υποστήριξης (decision support) ή ως αυτόνομα συστήματα λήψης απόφασης (automated decision-making), υπάγονται υποχρεωτικά στο αυστηρότερο ρυθμιστικό πλαίσιο, αναγνωρίζοντας τον κίνδυνο συστημικών σφαλμάτων ή μεροληψίας (systemic errors, bias) που μπορούν να διακυβεύσουν τη νομιμότητα της δικαιοσύνης[6].

Η ένταξη ενός LLM στο καθεστώς υψηλού κινδύνου δημιουργεί μία σειρά από αυστηρές υποχρεώσεις διαφάνειας (transparency), τεκμηρίωσης (documentation), ερμηνευσιμότητας (explainability) και συνεχούς ελέγχου (continuous monitoring), που εκτείνονται σε όλο τον κύκλο ζωής του συστήματος. Ο Κανονισμός προβλέπει την απαίτηση για πλήρη τεκμηρίωση κάθε φάσης σχεδιασμού, εκπαίδευσης, παραμετροποίησης και λειτουργίας του μοντέλου. Οποιαδήποτε απόφαση, σύσταση ή αποτέλεσμα που παράγεται από το LLM πρέπει να συνοδεύεται από λεπτομερές ιστορικό ελέγχου (audit trail), το οποίο περιλαμβάνει αναλυτική καταγραφή των δεδομένων εισόδου (inputs), των επιλεγμένων αρχιτεκτονικών (architectures), των παραμετρικών ρυθμίσεων (parameter settings), καθώς και των διαδικασιών ελέγχου και επιβεβαίωσης της ορθότητας του αποτελέσματος (validation procedures). Η απαίτηση αυτή επεκτείνεται και στη φάση της εκπαίδευσης (training), επιβάλλοντας στον προγραμματιστή (developer) την υποχρέωση να αποδεικνύει τη νομιμότητα, τη συνάφεια και τη μη μεροληψία (legality, relevance, non-bias) των χρησιμοποιούμενων δεδομένων. Παράλληλα, η ανάγκη για ανιχνευσιμότητα (traceability) καλύπτει και τις αλλαγές που πραγματοποιούνται εκ των υστέρων στο μοντέλο (π.χ. επαναεκπαίδευση, retraining, τροποποίηση παραμέτρων), ώστε

να είναι πάντοτε δυνατή η εξιχνίαση της ακριβούς αλληλουχίας ενεργειών που οδήγησε σε ένα συγκεκριμένο αποτέλεσμα.

Επιπλέον, το ζήτημα της εξηγησιμότητας (explainability) αντιμετωπίζεται όχι απλώς ως τεχνικό ζητούμενο, αλλά ως θεσμική και νομική υποχρέωση. Τα LLMs που εντάσσονται στο νομικό περιβάλλον υποχρεούνται να παρέχουν ανθρώπινα κατανοητές εξηγήσεις (human-understandable explanations) για κάθε κρίσιμη απόφαση ή σύσταση, τόσο προς τον τελικό χρήστη (δικαστή, διοικητικό όργανο, πολίτη) όσο και προς τις αρμόδιες εποπτικές αρχές. Η εξηγησιμότητα πρέπει να βασίζεται σε συγκεκριμένες τεχνικές, όπως επίπεδα ερμηνευσιμότητας (interpretability modules), αναλυτικές περιγραφές των συνόλων δεδομένων (datasets) και των μεταβλητών που επηρέασαν το αποτέλεσμα, χρήση προσεγγιστικών μοντέλων (surrogate models) και μεθόδους οπτικοποίησης (visualization) που αναδεικνύουν τη λογική διαδρομή από την είσοδο (input) στην έξοδο (output). Ωστόσο, ο Κανονισμός τονίζει ότι η διαφάνεια (transparency) δεν μπορεί να είναι προσχηματική ή να εξαντλείται σε τυπικές αιτιολογήσεις απαιτείται ουσιαστική επεξηγησιμότητα ενσωματωμένη εκ σχεδιασμού, ώστε ο χρήστης να μπορεί να αμφισβητήσει αποτελεσματικά μία απόφαση και να ζητήσει αναθεώρηση ή περαιτέρω διευκρινίσεις [82][6].

Ιδιαίτερη έμφαση δίνεται επίσης στη θεσμοθέτηση της ανθρώπινης εποπτείας (human oversight) σε όλα τα στάδια της αλγοριθμικής διαδικασίας[84]. Ο EU AI Act καθιστά σαφές ότι, ακόμη και σε πλήρως αυτοματοποιημένα ή αυτόνομα συστήματα, ο άνθρωπος διατηρεί την αρμοδιότητα παρέμβασης, επανεξέτασης και, όπου χρειάζεται, ανάρτησης του αλγοριθμικού αποτελέσματος. Το δικαίωμα αμφισβήτησης (challenge) και η δυνατότητα προσφυγής σε ανεξάρτητο μηχανισμό ελέγχου ή ένδικο μέσα ενσωματώνονται ως αναγκαία στοιχεία της θεσμικής νομιμότητας, ιδιαίτερα σε εφαρμογές που αγγίζουν το πεδίο του δικαίου και της δημόσιας διοίκησης.

Η διαχείριση του κινδύνου (risk management) και η συνεχής παρακολούθηση των LLMs αποτελούν επίσης κρίσιμους άξονες του θεσμικού πλαισίου. Ο EU AI Act επιβάλλει συστηματική αξιολόγηση των κινδύνων (risk assessment) και εκπόνηση αναλύσεων επιπτώσεων (impact assessments) τόσο πριν από τη διάθεση του συστήματος στην αγορά όσο και κατά τη διάρκεια της λειτουργίας του[13]. Οι απαιτήσεις αυτές μεταφράζονται σε διαρκή έλεγχο για μεροληψία (bias), ανθεκτικότητα (robustness), ισοτιμία (fairness) και μετατόπιση (drift), με ταυτόχρονη υποχρέωση καταγραφής και αναφοράς (logging and reporting) τυχόν περιστατικών παραβίασης δικαιωμάτων, σφαλμάτων ή δυσλειτουργιών προς τις αρμόδιες αρχές. Ο Κανονισμός δημιουργεί ένα ενιαίο πλαίσιο εποπτείας σε ευρωπαϊκό και εθνικό επίπεδο (oversight framework), ενισχύοντας τη διαλειτουργικότητα (interoperability), την αξιοπιστία (reliability) των συστημάτων και την προστασία των πολιτών.

Αξίζει να σημειωθεί πως το ευρωπαϊκό πλαίσιο διαφέρει σημαντικά ως προς την αυστηρότητα και τις απαιτήσεις σε σύγκριση με άλλα διεθνή καθεστώτα. Η Ευρωπαϊκή Ένωση υιοθετεί τον πλέον προωθημένο και δεσμευτικό μηχανισμό ελέγχου (compliance mechanism), με υψηλές απαιτήσεις συμμόρφωσης (compliance), ενισχυμένη λογοδοσία (accountability) και αυστηρές κυρώσεις (sanctions) σε περιπτώσεις μη τήρησης. Το κόστος συμμόρφωσης (compliance cost) αυξάνεται, καθώς οι προγραμματιστές (developers) και οι δημόσιοι φορείς καλούνται να επενδύσουν σε τεχνολογίες διαφάνειας (transparency technologies), εσωτερικές πολιτικές ελέγχου (auditing policies) και διεπιστημονικές ομάδες διακυβέρνησης

(governance teams), ώστε να διασφαλιστεί η πλήρης κάλυψη του ρυθμιστικού βάρους (regulatory burden). Την ίδια στιγμή, το όφελος για το κράτος δικαίου και τους πολίτες είναι ουσιώδες: η απονομή δικαιοσύνης ενισχύεται με πραγματικά ελέγξιμα, διαφανή και ερμηνεύσιμα εργαλεία, ενώ ο κίνδυνος αυθαίρετων ή μεροληπτικών αλγοριθμικών αποτελεσμάτων περιορίζεται δραστικά [82].

Εν κατακλείδι, ο Κανονισμός Τεχνητής Νοημοσύνης της ΕΕ (EU AI Act) σηματοδοτεί την πιο φιλόδοξη μέχρι σήμερα θεσμική προσπάθεια ενσωμάτωσης της τεχνολογικής καινοτομίας στο νομικό πεδίο υπό συνθήκες απόλυτου σεβασμού των θεμελιωδών αρχών του δικαίου και της κοινωνικής δικαιοσύνης (social justice). Η εφαρμογή του αποτελεί πρότυπο για τη διεθνή κοινότητα, δημιουργώντας μια ισχυρή παρακαταθήκη για την υπεύθυνη και κοινωνικά αποδεκτή ανάπτυξη και χρήση των LLMs στον χώρο του δικαίου.

Συνεπώς, οι τεχνικές λύσεις που αναλύονται και στις προηγούμενες ενότητες, όπως οι πλήρεις διαδρομές συλλογισμού, οι μηχανισμοί διασταυρούμενης επαλήθευσης, η δυναμική καταγραφή ενεργειών και οι τεχνικές ανίχνευσης παραισθήσεων, δεν αποτελούν απλώς τεχνολογικές καινοτομίες. Αντιθέτως, συνιστούν ουσιώδεις απαντήσεις στις συγκεκριμένες θεσμικές απαιτήσεις του Κανονισμού, όπως η τεκμηρίωση, η ανιχνευσιμότητα, η διαφάνεια και η εξηγησιμότητα. Η τεχνική συμμόρφωση δεν είναι αποσπασμένη από τις νομικές επιταγές, αλλά αποτελεί αναγκαία προϋπόθεση για την ουσιαστική ενσωμάτωση των LLMs στο ευρωπαϊκό νομικό πλαίσιο, σύμφωνα με το άρθρο 9 του Κανονισμού.

5.2.2 Διεθνείς Ρυθμιστικές Προσεγγίσεις

Η ρύθμιση της τεχνητής νοημοσύνης στον νομικό τομέα δεν ακολουθεί ενιαία παγκόσμια γραμμή, αλλά διαμορφώνεται δυναμικά μέσα από διαφορετικές παραδόσεις, πολιτικές επιλογές και κοινωνικά πλαίσια, με βασικούς πόλους την Ευρωπαϊκή Ένωση, τις Ηνωμένες Πολιτείες Αμερικής και την Κίνα. Κάθε μία από αυτές τις δικαιοδοσίες προσεγγίζει την ανάπτυξη και ενσωμάτωση της Τεχνητής Νοημοσύνης στο δίκαιο με διαφορετική φιλοσοφία, τεχνικές προδιαγραφές και ρυθμιστικά εργαλεία, γεγονός που δημιουργεί τόσο ευκαιρίες όσο και ρυθμιστικές προκλήσεις στην παγκόσμια αγορά νομικής τεχνητής νοημοσύνης.

Στις Ηνωμένες Πολιτείες, το ρυθμιστικό πλαίσιο χαρακτηρίζεται από έντονη αποκέντρωση (decentralization) και πολυμορφία, καθώς η απουσία ενιαίας ομοσπονδιακής νομοθεσίας για τα συστήματα τεχνητής νοημοσύνης (AI) έχει ως αποτέλεσμα η ρύθμιση να προκύπτει κυρίως από μεμονωμένες πολιτειακές πρωτοβουλίες, επαγγελματικά πρότυπα (professional standards), δικαστικά προηγούμενα (case law) και κλαδικές κατευθυντήριες γραμμές (sectoral guidelines)[83]. Η βασική προτεραιότητα των αμερικανικών προσεγγίσεων είναι η προώθηση της καινοτομίας (innovation) και της ανταγωνιστικότητας (competitiveness), με ταυτόχρονη προστασία του καταναλωτή μέσω μηχανισμών εκ των υστέρων (ex post) ευθύνης (liability) και προσφυγής στη δικαιοσύνη, παρά μέσω εκτεταμένων εκ των προτέρων (ex ante) ρυθμιστικών ελέγχων όπως στην Ευρώπη. Έτσι, η λογοδοσία (accountability) στα LLMs που χρησιμοποιούνται σε νομικές εφαρμογές διασφαλίζεται κυρίως μέσω μηχανισμών ευθύνης προϊόντος (product liability), επαγγελματικής ευθύνης (professional liability) και επιμέρους νομολογιακών λύσεων (ad hoc case law) σε περιπτώσεις όπου προκύπτουν προβλήματα μεροληψίας (bias), αδιαφάνειας (opacity) ή ακούσιας βλάβης (harm) προς πολίτες και διαδίκους.

Επιπλέον, οι προγραμματιστές και οι πάροχοι νομικών LLMs στις ΗΠΑ οφείλουν να λαμβάνουν υπόψη τους την αυξανόμενη απαίτηση για ηθική τεχνητή νοημοσύνη (ethical AI), κυρίως μέσα από πρωτοβουλίες αυτορρύθμισης (self-regulation), υιοθέτηση βέλτιστων πρακτικών (best practices) και ανεπίσημα πλαίσια ελέγχου (auditing frameworks), χωρίς ωστόσο να υπόκεινται σε ενιαίο ρυθμιστικό καθεστώς (regulatory regime). Ως αποτέλεσμα, η ασυνεχής αυτή ρυθμιστική πραγματικότητα δημιουργεί περιβάλλον υψηλής νομικής αβεβαιότητας (legal uncertainty) και κίνδυνο δικαστικού φόρουμ (forum shopping), όπου οι εμπλεκόμενοι μπορεί να επιδιώκουν εκδίκαση σε ευνοϊκότερες δικαστικές περιφέρειες, ιδίως σε περιπτώσεις αλγοριθμικών λαθών, συστημικής μεροληψίας (systemic bias) ή αδιαφανών αποφάσεων.

Αντίθετα, το ρυθμιστικό πλαίσιο της Κίνας υιοθετεί μια συγκεντρωτική και κρατικοκεντρική προσέγγιση (centralized, state-centric approach), όπου το κράτος αποτελεί τον βασικό ελεγκτή και εγγυητή της λειτουργίας των συστημάτων τεχνητής νοημοσύνης. Ο «Βασικός Νόμος Τεχνητής Νοημοσύνης» (Basic Law of AI) της Κίνας επιβάλλει απαιτήσεις λογοδοσίας (accountability) και ιχνηλασιμότητας (traceability) στους προγραμματιστές και τους παρόχους συστημάτων ΤΝ, με αυστηρό έλεγχο (auditing), υποχρεωτική καταγραφή και τεκμηρίωση όλων των κρίσιμων παραμέτρων και λειτουργιών. Η εξηγησιμότητα (explainability) και η διαφάνεια (transparency), ωστόσο, εστιάζονται κυρίως στην προστασία των κρατικών συμφερόντων και λιγότερο στη διασφάλιση ατομικών δικαιωμάτων ή στη δυνατότητα αμφισβήτησης (challenge) από πολίτες και επαγγελματίες του δικαίου [87]. Στο πλαίσιο των «έξυπνων δικαστηρίων» (smart courts) που λειτουργούν ήδη σε αρκετές επαρχίες, οι αλγοριθμικές αποφάσεις LLMs τυποποιούνται και τεκμηριώνονται ως προς τη συμμόρφωση με κρατικά πρότυπα (state standards), αλλά η διαφάνεια προς τον τελικό χρήστη ή τον διάδικο παραμένει περιορισμένη. Το βάρος της λογοδοσίας μετατίθεται από την ατομική ή εταιρική ευθύνη προς το κράτος ως τελικό διαχειριστή και εγγυητή, γεγονός που ενισχύει τη ρυθμιστική συνοχή (regulatory coherence) αλλά ταυτόχρονα υπονομεύει τη δυνατότητα ανεξάρτητου ελέγχου ή δημόσιας λογοδοσίας.

Πέραν των τριών βασικών αυτών πυλώνων, καταγράφονται διεθνώς και άλλες σημαντικές ρυθμιστικές πρωτοβουλίες, όπως τα (regulatory sandboxes) σε χώρες του Οργανισμού Οικονομικής Συνεργασίας και ανάπτυξης (ΟΟΣΑ), η προσπάθεια εναρμόνισης μέσω διεθνών προτύπων (ISO/IEC, IEEE) και η σταδιακή ενσωμάτωση κατευθυντήριων γραμμών δεοντολογίας (ethical guidelines) και ισοτιμίας (fairness) στους εθνικούς κώδικες δικαστικής δεοντολογίας (judicial codes of conduct) και διακυβέρνησης δεδομένων (data governance). Η παγκόσμια εικόνα του ρυθμιστικού τοπίου για τα LLMs παραμένει κατακερματισμένη (fragmented) και δυναμική, καθώς κάθε χώρα προσπαθεί να ισορροπήσει μεταξύ προώθησης της καινοτομίας, διασφάλισης των δικαιωμάτων και ενίσχυσης της θεσμικής νομιμότητας (institutional legitimacy).

Αυτό το ετερόκλητο περιβάλλον δημιουργεί, αναπόφευκτα, προκλήσεις για τους προγραμματιστές, τους οργανισμούς και τους τελικούς χρήστες των LLMs σε νομικά περιβάλλοντα [83]. Ενώ η Ευρωπαϊκή Ένωση θέτει τον υψηλότερο θεσμικό πήχη με την εκ των προτέρων ρύθμιση (ex ante regulation), την έμφαση στη διαφάνεια (transparency) και τον σαφή επιμερισμό ευθυνών (explicit liability allocation), οι ΗΠΑ υιοθετούν μια περισσότερο πραγματιστική, εκ των υστέρων (ex post) φιλοσοφία, προάγοντας την ταχεία υιοθέτηση καινοτομιών, αλλά με με-

γαλύτερους κινδύνους για νομική αβεβαιότητα (legal uncertainty) και forum shopping[87]. Από την άλλη, το κινεζικό μοντέλο, παρά τη φαινομενικά υψηλή συμμόρφωση (compliance) και το αυστηρό αυδιτινγ, θυσιάζει συχνά την ατομική προστασία και τον ανεξάρτητο έλεγχο χάριν της κρατικής συνοχής και αποτελεσματικότητας.

Εν κατακλείδι, η σύγκριση των διεθνών ρυθμιστικών προσεγγίσεων αναδεικνύει ότι, παρότι η ανάγκη για λογοδοσία (accountability), διαφάνεια (transparency) και τεχνική συμμόρφωση (compliance) στα LLMs είναι παγκοσμίως αναγνωρισμένη, οι πρακτικές υλοποιήσης, οι προτεραιότητες και τα μέσα επίτευξης της ποικίλλουν δραματικά. Η παγκόσμια τάση οδηγεί σταδιακά σε μεγαλύτερη εναρμόνιση μέσω θεσμικών προτύπων και τεχνολογικών εργαλείων ελέγχου (auditing tools), ωστόσο ο δρόμος προς μια πραγματικά ενοποιημένη, ασφαλή και δίκαιη ρύθμιση παραμένει ανοικτός, απαιτώντας διεπιστημονική συνεργασία (interdisciplinary collaboration) και συνεχή θεσμική καινοτομία.

5.2.3 Υλοποίηση Συμμόρφωσης για Προγραμματιστές

Η διαρκώς αυξανόμενη πολυπλοκότητα και ο κατακερματισμός του διεθνούς ρυθμιστικού πλαισίου επιβάλλουν στους προγραμματιστές (developers) που δραστηριοποιούνται στον χώρο της νομικής τεχνητής νοημοσύνης να υιοθετήσουν δομημένες και πολυεπίπεδες στρατηγικές συμμόρφωσης (compliance strategies). Η υλοποίηση της συμμόρφωσης (compliance implementation) δεν αποτελεί πλέον μία απλή διαδικαστική απαίτηση, αλλά θεμελιώδη προϋπόθεση για την εμπορική διάθεση, την αποδοχή και τη θεσμική ενσωμάτωση των μεγάλων γλωσσικών μοντέλων (LLMs) και των συναφών συστημάτων τεχνητής νοημοσύνης (AI) στα νομικά οικοσυστήματα της Ευρώπης και διεθνώς.

Στο πλαίσιο της Ευρωπαϊκής Ένωσης, η συμμόρφωση με τον EU AI Act συνιστά πολυσύνθετη πρόκληση που διατρέχει όλα τα στάδια ανάπτυξης και υλοποίησης των νομικών LLMs. Κατά πρώτο λόγο, η επιλογή, συλλογή και επεξεργασία των δεδομένων εκπαίδευσης (training data selection, collection and processing) υπόκειται σε αυστηρούς κανόνες τεκμηρίωσης (documentation) και ελέγξιμότητας (auditability), καθώς οι προγραμματιστές (developers) υποχρεούνται να αποδεικνύουν τη νομιμότητα της πηγής των δεδομένων, τη συνάφειά τους με το πεδίο εφαρμογής, καθώς και τη λήψη ειδικών μέτρων για την αποφυγή μεροληψίας (bias mitigation) και την προστασία της ιδιωτικότητας (privacy protection). Το στάδιο της επιμέλειας και προεπεξεργασίας των δεδομένων (data curation and preprocessing) αποκτά ιδιαίτερη βαρύτητα, με την ανάγκη ενσωμάτωσης τεχνολογιών ανίχνευσης και διόρθωσης μεροληψίας (bias detection and correction), ενώ παράλληλα εφαρμόζονται αυτοματοποιημένα εργαλεία ελέγχου (auditing tools) που διασφαλίζουν τον εντοπισμό ασυνεπειών και την τήρηση πλήρους ιστορικού επεξεργασίας (processing logs)[6][84].

Κατά τη φάση του σχεδιασμού και της εκπαίδευσης (design and training) [82] των μεγάλων γλωσσικών μοντέλων, οι προγραμματιστές οφείλουν να ενσωματώνουν μηχανισμούς εξηγησιμότητας (explainability) και διαφάνειας εκ κατασκευής (transparency by design), ώστε η λειτουργία του συστήματος να παραμένει ερμηνεύσιμη (interpretable) και ελέγξιμη τόσο από τους τελικούς χρήστες όσο και από τις αρμόδιες εποπτικές αρχές. Οι πρακτικές αυτές περιλαμβάνουν την υλοποίηση επιπέδων ερμηνευσιμότητας (interpretability layers), τη διατήρηση αναλυτικών αρχείων καταγραφής (detailed operation logs) και τη διασφάλι-

ση δυνατότητας αναδρομής (traceability) όλων των κρίσιμων αποφάσεων και αλλαγών στο pipeline του συστήματος.

Επιπρόσθετα, προβλέπεται η συστηματική διενέργεια αξιολογήσεων κινδύνου και επιπτώσεων (risk and impact assessments), με έμφαση σε δοκιμές για ισοτιμία (fairness), ανθεκτικότητα (robustness) και αντίσταση σε επιθέσεις (adversarial resilience), ώστε να περιορίζεται η πιθανότητα εμφάνισης συστημικών σφαλμάτων ή ακούσιων παραβιάσεων θεμελιωδών δικαιωμάτων.

Καθ' όλη τη διάρκεια της εμπορικής διάθεσης και λειτουργίας του συστήματος, η συμμόρφωση με τις ρυθμιστικές απαιτήσεις (regulatory requirements) προϋποθέτει συνεχή παρακολούθηση της απόδοσης του LLM (performance monitoring), συλλογή ανατροφοδότησης (user feedback) από τους τελικούς χρήστες, καθώς και άμεση αναφορά κάθε περιστατικού προκατάληψης (bias), δυσλειτουργίας (malfunction) ή μη συμμόρφωσης (non-compliance) προς τις αρμόδιες αρχές. Οι υπεύθυνοι ανάπτυξης (developers) οφείλουν να διατηρούν αυτοματοποιημένους μηχανισμούς ελέγχου (automated auditing mechanisms), δυνατότητες επαναφοράς (rollback) και ανάλυσης περιστατικών (forensic analysis), καθώς και οργανωτικές δομές που επιτρέπουν τη συνεχή αναβάθμιση των διαδικασιών συμμόρφωσης, βάσει των επικαιροποιήσεων του ρυθμιστικού πλαισίου (regulatory updates) και των ευρημάτων εσωτερικών ή εξωτερικών ελέγχων. Η λειτουργία πολυεπιστημονικών ομάδων (multidisciplinary teams) που συνδυάζουν νομική, τεχνική και δεοντολογική τεχνογνωσία (legal, technical and ethical expertise) είναι αναγκαία για την ορθή ερμηνεία των κανονιστικών απαιτήσεων και την αποτελεσματική ενσωμάτωσή τους στις διαδικασίες ανάπτυξης και διάθεσης συστημάτων τεχνητής νοημοσύνης.

Στις ΗΠΑ, αν και απουσιάζει ένα ενιαίο, δεσμευτικό πλαίσιο αντίστοιχο του AI Act, η ανάγκη για συμμόρφωση δεν είναι λιγότερο επιτακτική, λόγω της αυξανόμενης πίεσης από τα δικαστήρια, τις επαγγελματικές ενώσεις και τους ίδιους τους πελάτες. Οι προγραμματιστές (developers) υποχρεώνονται να διασφαλίζουν τη διαφάνεια (transparency) των συστημάτων τους, να υιοθετούν βέλτιστες πρακτικές αυτορρύθμισης (self-regulation) και να προετοιμάζονται για ενδεχόμενη λογοδοσία (accountability) στο πλαίσιο αστικών αγωγών, συλλογικών προσφυγών (class actions) ή ρυθμιστικών ελέγχων (regulatory audits) μετά από καταγγελίες για μεροληψία (bias) ή αδιαφάνεια (opacity)[85][83]. Στην πράξη, αυτό μεταφράζεται στην υιοθέτηση διαδικασιών ελέγχου (auditing procedures), τη χρήση ανεξάρτητων αξιολογήσεων ισοτιμίας και ανθεκτικότητας (independent fairness and robustness assessments), και την ανάπτυξη εργαλείων για την παραγωγή αρχείων ελέγχου (audit logs) που μπορούν να αξιοποιηθούν σε διαδικασίες ανάλυσης περιστατικών (forensic analysis).

Το κινεζικό μοντέλο, τέλος, απαιτεί από τους προγραμματιστές συμμόρφωση πρωτίστως με τα κρατικά πρότυπα (state standards), την υλοποίηση διαδικασιών κρατικού ελέγχου (government auditing) και την παροχή διαρκούς τεκμηρίωσης (continuous documentation) που να αποδεικνύει τη συμμόρφωση με τους κεντρικούς κανόνες λογοδοσίας (accountability) και ανιχνευσιμότητας (traceability). Η έμφαση δίνεται κυρίως στη διαφάνεια (transparency) προς το κράτος και όχι απαραίτητα προς τον τελικό χρήστη ή τον πολίτη, γεγονός που διαμορφώνει διαφορετικές τεχνικές και οργανωτικές πρακτικές υλοποίησης συμμόρφωσης.

Συνολικά, η επιτυχής υλοποίηση της συμμόρφωσης από την πλευρά των προγραμματιστών (developers) στον χώρο των νομικών LLMs εξαρτάται από τη βαθιά κατανόηση τόσο

των τεχνικών προκλήσεων της τεχνητής νοημοσύνης όσο και των σύνθετων απαιτήσεων που θέτουν τα εθνικά και υπερεθνικά ρυθμιστικά πλαίσια (national and supranational regulatory frameworks). Η προληπτική ενσωμάτωση της συμμόρφωσης στο σχεδιασμό (compliance by design), η συνεχής παρακολούθηση των διεθνών εξελίξεων και η διατήρηση ευέλικτων εσωτερικών πολιτικών ελέγχου (internal auditing policies) και διαχείρισης κινδύνων (risk management) αποτελούν το βασικό υπόβαθρο για την επιτυχημένη, ασφαλή και κοινωνικά αποδεκτή διάθεση νομικών συστημάτων τεχνητής νοημοσύνης στην αγορά.

5.3 Ηθικές Παραμέτρους

Η ραγδαία διείσδυση των μεγάλων γλωσσικών μοντέλων και των προηγμένων συστημάτων τεχνητής νοημοσύνης στο νομικό οικοσύστημα θέτει στο προσκήνιο κρίσιμα ηθικά ερωτήματα που υπερβαίνουν το στενό πεδίο της νομικής συμμόρφωσης. Η δυναμική λειτουργία των LLMs, η αυτονομία στη λήψη αποφάσεων και η συχνά μαύρη φύση των αλγοριθμικών διαδικασιών, δημιουργούν νέα ηθικά διλήμματα που αφορούν τόσο την ουσία της δικαιοσύνης όσο και τη διασφάλιση των θεμελιωδών αξιών της κοινωνίας. Οι βασικές αυτές διαστάσεις αφορούν τον μετριασμό της αλγοριθμικής μεροληψίας (bias mitigation), την προστασία της ιδιωτικότητας, την αρχή της ισότητας πρόσβασης, καθώς και τα πολιτισμικά και κοινωνικά ζητήματα που εγείρονται με τη χρήση συστημάτων ΤΝ στη δικαιοσύνη[88][89]

Τα LLMs, λόγω της εκπαίδευσής τους σε τεράστιους όγκους δεδομένων που προέρχονται από ετερογενείς πηγές, είναι ευάλωτα στην εμφάνιση και ενίσχυση συστηματικών μορφών προκαταλήψεων bias. Η μεροληψία μπορεί να έχει πολλές μορφές, από την ενίσχυση στερεοτύπων στο output του συστήματος, έως την άνιση μεταχείριση κοινωνικών ομάδων ή τη συστηματική υποεκπροσώπηση συγκεκριμένων κατηγοριών δεδομένων κατά την εκπαίδευση. Μελέτες δείχνουν ότι τα LLMs μπορεί να ενσωματώνουν προκαταλήψεις που υπάρχουν στα δεδομένα εκπαίδευσης, αναπαράγοντας έτσι ιστορικές ανισότητες και σε ορισμένες περιπτώσεις, οδηγώντας σε εσφαλμένες ή άδικες αλγοριθμικές συστάσεις που δύνανται να επηρεάσουν τη δικαστική κρίση ή την έκδοση υποθέσεων [90], [91]. Η διαχείριση της αλγοριθμικής μεροληψίας αποτελεί έτσι κεντρικό ζήτημα ηθικής ευθύνης για τους developers και τους φορείς υλοποίησης τέτοιων συστημάτων, καθώς η αδυναμία ελέγχου των προκαταλήψεων μπορεί να οδηγήσει σε παραβίαση της αρχής της ισότητας ενώπιον του νόμου και της αμεροληψίας της δικαιοσύνης.

Παράλληλα, το ζήτημα της ιδιωτικότητας αναδεικνύεται ως εξαιρετικά σύνθετο στην εποχή της νομικής τεχνητής νοημοσύνης. Τα LLMs, κατά την εκπαίδευση και λειτουργία τους, ενδέχεται να επεξεργάζονται ή ακόμη και να αποθηκεύουν προσωπικά δεδομένα, ευαίσθητες πληροφορίες ή νομικά προστατευόμενα στοιχεία, δημιουργώντας αυξημένους κινδύνους παραβίασης της ιδιωτικότητας και απώλειας του ελέγχου επί των δεδομένων. Η λήψη επαρκών τεχνικών και οργανωτικών μέτρων για την ανωνυμοποίηση, την ασφαλή διαχείριση και την προστασία των δεδομένων αποτελεί ηθικό και νομικό ζητούμενο πρώτης γραμμής, καθώς η αποτυχία διασφάλισης της ιδιωτικότητας δύναται να υπονομεύσει την κοινωνική εμπιστοσύνη στη χρήση της τεχνητής νοημοσύνης στο δικαστικό περιβάλλον.[89][92]

Ένα ακόμη σημαντικό ηθικό ζήτημα αποτελεί η διασφάλιση της δίκαιης πρόσβασης και της μη αποκλειστικότητας στη χρήση των LLMs στο δικαστικό σύστημα[89]. Ο κίνδυνος δη-

μιουργίας ενός 'τεχνολογικού χάσματος' (digital divide) μεταξύ διαφορετικών κοινωνικών ή γεωγραφικών ομάδων, είτε λόγω ανισοκατανομής των υποδομών είτε λόγω έλλειψης τεχνικής κατάρτισης, μπορεί να εντείνει τις ήδη υπάρχουσες κοινωνικές ανισότητες. Η ηθική ευθύνη των developers και των φορέων λήψης αποφάσεων επεκτείνεται επομένως και στη σχεδίαση συστημάτων που προάγουν την καθολική προσβασιμότητα (universal accessibility), τη γλωσσική και πολιτισμική συμπερίληψη (linguistic and cultural inclusion) και τον σεβασμό της διαφορετικότητας.[93][91]

Τέλος, η ανάπτυξη και λειτουργία των LLMs στο νομικό πεδίο συνδέεται με τη γενικότερη κοινωνική ευθύνη για πρόβλεψη, πρόληψη και αντιμετώπιση των ακούσιων συνεπειών που μπορεί να προκύψουν από τη μαζική αλγοριθμική λήψη αποφάσεων. Οι πρακτικές αξιολόγησης και ελέγχου auditing, η ενσωμάτωση μηχανισμών παρακολούθησης (monitoring) και η συνεχής αξιολόγηση των συστημάτων αποτελούν ουσιώδη εργαλεία ηθικής διακυβέρνησης (ethical governance). Οι προκλήσεις που αφορούν τη 'μαύρη' φύση των LLMs, την αδυναμία πλήρους εξηγήσιμης λήψης αποφάσεων (explainable decision-making) και τον κίνδυνο ακούσιων σφαλμάτων ή παραισθήσεων (hallucinations), καθιστούν επιτακτική τη διαρκή εγρήγορση και τη διεπιστημονική συνεργασία (interdisciplinary collaboration) όλων των εμπλεκομένων φορέων.

5.3.1 Τεχνικές Μετριασμού Αλγοριθμικών Προκαταλήψεων

Η αλγοριθμική μεροληψία (algorithmic bias) αποτελεί ίσως τον σημαντικότερο ηθικό και τεχνικό κίνδυνο για τα συστήματα νομικής τεχνητής νοημοσύνης, καθώς η ανεπαρκής διαχείρισή της μπορεί να οδηγήσει σε συστημικές ανισότητες, αδικίες και αμφισβήτηση της θεσμικής νομιμότητας των αποτελεσμάτων των LLMs. Οι πηγές μεροληψίας στα LLMs είναι πολυπαραγοντικές και ξεκινούν ήδη από το στάδιο της επιλογής και προετοιμασίας των δεδομένων εκπαίδευσης, συνεχίζονται στην αρχιτεκτονική και τις παραμετρικές επιλογές του μοντέλου, και εκτείνονται ως τη φάση χρήσης και αλληλεπίδρασης με το ανθρώπινο περιβάλλον.[91]

Η ενσωμάτωση τεχνικών μετριασμού της μεροληψίας αποτελεί απαραίτητο εργαλείο για κάθε προγραμματιστή που στοχεύει σε αξιόπιστες, αμερόληπτες και κοινωνικά αποδεκτές νομικές εφαρμογές AI. Οι τεχνικές αυτές μπορούν να ταξινομηθούν σε τρεις βασικές κατηγορίες: τεχνικές προ-επεξεργασίας δεδομένων (pre-processing), τεχνικές ενδο-μοντέλου (in-processing) και τεχνικές μετά-επεξεργασίας αποτελεσμάτων (post-processing).

Στο στάδιο της προ-επεξεργασίας[88], έμφαση δίνεται στην αναλυτική διερεύνηση των δεδομένων εκπαίδευσης με σκοπό την ταυτοποίηση και απομάκρυνση προκαταλήψεων, την ενίσχυση της εκπροσώπησης μειονοτικών ομάδων και τη χρήση τεχνικών data balancing ή data augmentation. Επιπλέον, εφαρμόζονται εργαλεία ανίχνευσης συστηματικών αποκλίσεων (bias detection metrics) καθώς και τεχνικές anonymization και de-biasing των εισροών, ώστε να περιοριστεί η είσοδος συστημικών στρεβλώσεων στο μοντέλο[88],[90].

Κατά την εκπαίδευση και τη βελτιστοποίηση των LLMs, αξιοποιούνται τεχνικές in-processing, όπως η ενσωμάτωση fairness constraints στους αλγορίθμους μάθησης, η χρήση adversarial training για τον εντοπισμό και τη διόρθωση bias, καθώς και η ανάπτυξη εξειδικευμένων loss functions που επιβραβεύουν την αμεροληψία και τιμωρούν τις συστηματικές

αποκλίσεις στις προβλέψεις του μοντέλου. Η σύγχρονη έρευνα προτείνει επίσης την ενσωμάτωση regularization terms που λειτουργούν ως φραγμοί στην ενίσχυση υπαρχόντων στερεοτύπων, ενώ τεχνικές όπως οι μετρικές αμεροληψίας fairness-aware embeddings και η ενσωμάτωση βαθμονομημένων επιπέδων εξόδων (calibrated output layers) ενισχύουν τη σταθερότητα του συστήματος έναντι προκαταλήψεων[88],[91].

Στο τελικό στάδιο παραγωγής των αποτελεσμάτων, εφαρμόζονται τεχνικές post-processing που εστιάζουν στον εντοπισμό και τη διόρθωση προκαταλήψεων στα αποτελέσματα outputs. Οι μέθοδοι αυτοί περιλαμβάνουν την αυτόματη αναγνώριση μεροληπτικών προτύπων στην έξοδο του μοντέλου, την αναπροσαρμογή των τελικών προβλέψεων βάσει μετρικών δικαιοσύνης (fairness metrics) και την αξιοποίηση εξωτερικών ενοτήτων ελέγχων auditing modules που επιτρέπουν τη δυναμική ανατροφοδότηση του συστήματος με πραγματικά δεδομένα χρήσης. Οι πρακτικές αυτές συνοδεύονται συχνά από συμμετοχή ανθρώπινου παράγοντα (human-in-the-loop auditing), ενισχύοντας τη δυνατότητα ανίχνευσης ανεπιθύμητων αποκλίσεων και επιτρέποντας την άμεση παρέμβαση σε περιπτώσεις αποτυχίας[85],[91].

Η αποτελεσματικότητα των τεχνικών μετριασμού μεροληψίας εξαρτάται σε μεγάλο βαθμό από τη συστηματική εφαρμογή ολοκληρωμένων μετρικών αξιολόγησης, όπως οι δίκτες άνισης επίπτωσης (disparate impact ratios), η δημογραφική ισότητα (demographic parity), οι εξοσιμμένες πιθανότητες (equalized odds) και οι καμπύλες βαθμονόμησης calibration curves. Η διαρκής παρακολούθηση της συμπεριφοράς του LLM σε διαφορετικά υποσύνολα δεδομένων και η χρήση πραγματικών σημείων αναφοράς (real-world benchmarks) αποτελούν προϋποθέσεις για τη διασφάλιση της συνεχούς αμεροληψίας του συστήματος σε πραγματικές νομικές συνθήκες[90],[92].

Αναμφίβολα, η διαχείριση της αλγοριθμικής μεροληψίας δεν αποτελεί στατικό ζήτημα, αλλά διαρκή ηθική και τεχνική πρόκληση που απαιτεί συνδυασμό εξειδικευμένων εργαλείων, οργανωτικών πολιτικών και διεπιστημονικής συνεργασίας. Μόνο μέσα από τη συνεχή και ενεργητική ενασχόληση των προγραμματιστών, των νομικών και των ανεξάρτητων φορέων αξιολόγησης μπορεί να διασφαλιστεί ότι τα LLMs που ενσωματώνονται στο δικαστικό και νομικό οικοσύστημα λειτουργούν ως πολλαπλασιαστές δικαιοσύνης και όχι ως ενισχυτές κοινωνικών ανισοτήτων.

5.3.2 Προστασία της Ιδιωτικότητας στη Νομική ΤΝ

Η ιδιωτικότητα συνιστά μία από τις πιο σύνθετες και κρίσιμες προκλήσεις στον σχεδιασμό και την υλοποίηση μεγάλων γλωσσικών μοντέλων για τον νομικό τομέα. Τα LLMs, λόγω του μεγάλου όγκου και της φύσης των δεδομένων με τα οποία εκπαιδεύονται, συχνά έρχονται σε επαφή με προσωπικά και ευαίσθητα νομικά δεδομένα, γεγονός που εγείρει σημαντικά ζητήματα τόσο ως προς την προστασία των πληροφοριών των υποκειμένων όσο και ως προς την ασφάλεια των ιδίων των διαδικασιών απονομής δικαιοσύνης. Όπως αναδεικνύεται και στη σχετική βιβλιογραφία[89] η ανάγκη για ενσωμάτωση ισχυρών μηχανισμών προστασίας της ιδιωτικότητας είναι πλέον επιτακτική σε κάθε στάδιο ανάπτυξης και εφαρμογής νομικών συστημάτων τεχνητής νοημοσύνης.[89][91]

Στο νομικό οικοσύστημα, όπου διακυβεύονται θεμελιώδη δικαιώματα και ευαίσθητες πληροφορίες για πολίτες, διαδίκους, δικηγόρους και θεσμικούς φορείς, η παραβίαση της

ιδιωτικότητας ενδέχεται να έχει ανεπανόρθωτες συνέπειες, τόσο για την προσωπική ζωή όσο και για την αξιοπιστία του ίδιου του συστήματος απονομής δικαιοσύνης[89]. Οι προκλήσεις αυτές καθίστανται εντονότερες καθώς τα LLMs μπορούν να λειτουργήσουν ως αποθήκες γνώσης (knowledge repositories) που διατηρούν, συνδυάζουν ή ακόμα και διαρρέουν πληροφορίες μέσω των αποτελεσμάτων τους, ιδίως σε περιβάλλοντα όπου δεν υπάρχουν αυστηροί μηχανισμοί ελέγχου και διαχείρισης πρόσβασης (access control)[65].

Η προστασία της ιδιωτικότητας στα συστήματα νομικής τεχνητής νοημοσύνης (AI) απαιτεί πολυεπίπεδη προσέγγιση, η οποία ξεκινά από το στάδιο της επιλογής και επεξεργασίας των δεδομένων εκπαίδευσης. Οι προγραμματιστές καλούνται να εφαρμόζουν τεχνικές ανωνυμοποίησης (anonymization), ψευδωνυμοποίησης (pseudonymization) και αρχής ελαχιστοποίησης των δεδομένων (data minimization), διασφαλίζοντας ότι οι αλγοριθμικές διαδικασίες δεν διατηρούν, δεν επεξεργάζονται και δεν αναπαράγουν αναγνωρίσιμα προσωπικά δεδομένα χωρίς τη συγκατάθεση των υποκειμένων των δεδομένων ή την ύπαρξη σαφούς νομικής βάσης.

Όπως επισημαίνεται και στη σχετική βιβλιογραφία[89],[91], η χρήση διαφοροποιημένης ιδιωτικότητας (differential privacy), ομοσπονδιακής μάθησης (federated learning) και μηχανικής μάθησης με διαφύλαξη της ιδιωτικότητας (privacy-preserving machine learning) αποτελεί σύγχρονη πρακτική που μειώνει τον κίνδυνο διαρροής δεδομένων (data leakage), χωρίς να θυσιάζεται η αποδοτικότητα του συστήματος.

Εξίσου σημαντική είναι η ενσωμάτωση μηχανισμών ελέγχου πρόσβασης (access control) και καταγραφής (logging) σε κάθε στάδιο χρήσης του LLM. Η διαφάνεια ως προς το ποιος, πότε και πώς έχει πρόσβαση σε δεδομένα και αποτελέσματα του συστήματος, ενισχύει την ιχνηλασιμότητα (traceability) και τη λογοδοσία (accountability)[89]. Παράλληλα, η τακτική αξιολόγηση κινδύνου (risk assessment) σε πραγματικές συνθήκες χρήσης, η διενέργεια penetration tests και η συνεχής εκπαίδευση των χρηστών συμβάλλουν στην πρόληψη παραβιάσεων και στη διαμόρφωση κουλτούρας προστασίας της ιδιωτικότητας[65].

Το ρυθμιστικό πλαίσιο (όπως ο GDPR στην Ευρώπη και τα επιμέρους εθνικά νομικά πρότυπα) ενισχύει τον ρόλο της ιδιωτικότητας, θεσπίζοντας υποχρεώσεις για "privacy by design and by default", λογοδοσία (accountability), δικαίωμα ενημέρωσης και διαγραφής ("right to be forgotten") και αυστηρές κυρώσεις για παραβάσεις. Οι προγραμματιστές καλούνται όχι μόνο να συμμορφώνονται με τις νομικές απαιτήσεις, αλλά και να σχεδιάζουν συστήματα που εμπνέουν εμπιστοσύνη στους χρήστες, διασφαλίζοντας την ισορροπία ανάμεσα στη λειτουργικότητα και στη διασφάλιση των προσωπικών δεδομένων.

Είναι σαφές πως, καθώς τα LLMs ενσωματώνονται όλο και περισσότερο σε ευαίσθητες διαδικασίες του δικαστικού και διοικητικού οικοσυστήματος, η προστασία της ιδιωτικότητας δεν αποτελεί απλώς νομική υποχρέωση, αλλά βασική ηθική αρχή και τεχνική προτεραιότητα που καθορίζει την αποδοχή, την αξιοπιστία και την αποτελεσματικότητα της τεχνητής νοημοσύνης στο δίκαιο.

5.3.3 Δίκαιη Πρόσβαση και Πολιτισμικές Διαστάσεις

Η ανάπτυξη και ενσωμάτωση μεγάλων γλωσσικών μοντέλων στο δικαστικό και νομικό οικοσύστημα φέρνει στο προσκήνιο τη ζωτική ανάγκη για διασφάλιση της δίκαιης πρόσβασης

στις τεχνολογίες τεχνητής νοημοσύνης και για αναγνώριση της πολιτισμικής και γλωσσικής ποικιλομορφίας. Το ζήτημα της δίκαιης πρόσβασης υπερβαίνει το παραδοσιακό πρόβλημα της τεχνολογικής άνισότητας (digital divide) και αγγίζει θεμελιώδεις αρχές της δικαιοσύνης, της κοινωνικής συμπερίληψης και της προστασίας της πολυφωνίας[94].

Στο πεδίο του δικαίου, η χρήση LLMs και συναφών εργαλείων τεχνητής νοημοσύνης (AI) για την παροχή νομικών υπηρεσιών ή τη διαμεσολάβηση σε διαδικασίες απονομής δικαιοσύνης μπορεί να δημιουργήσει νέες μορφές αποκλεισμού. Πρώτον, η διαθεσιμότητα και ποιότητα των υπηρεσιών τεχνητής νοημοσύνης (AI services) συχνά εξαρτάται από το γλωσσικό, κοινωνικό και οικονομικό υπόβαθρο των χρηστών, οδηγώντας σε προνομιακή μεταχείριση εκείνων που διαθέτουν τους απαραίτητους τεχνολογικούς πόρους ή γνώσεις. Η πρόσβαση σε νομικές πληροφορίες, σε 'έξυπνες' πλατφόρμες δικαστικής βοήθειας ή σε αυτοματοποιημένα εργαλεία διαμεσολάβησης δεν είναι παντού ισότιμη· ομάδες με περιορισμένη τεχνολογική υποδομή, ηλικιωμένοι, άτομα με αναπηρίες ή πληθυσμοί σε αγροτικές/μειονεκτούσες περιοχές διατρέχουν μεγαλύτερο κίνδυνο αποκλεισμού. Όπως επισημαίνεται στη σχετική βιβλιογραφία [94], η δίκαιη και καθολική πρόσβαση παραμένει ένα σύνθετο και ανοιχτό ζητούμενο για την τεχνητή νοημοσύνη στον νομικό τομέα[93].

Επιπλέον, τα LLMs τείνουν να είναι περισσότερο ακριβή και χρήσιμα για τις κυρίαρχες γλώσσες και τα πολιτισμικά πρότυπα που κυριαρχούν στο corpus εκπαίδευσής τους. Η ελλιπής εκπροσώπηση γλωσσών, τοπικών διαλέκτων ή νομικών παραδόσεων δημιουργεί τον κίνδυνο πολιτισμικής ομογενοποίησης και αόρατης 'μεροληψίας' υπέρ συγκεκριμένων κοινωνικών ή γλωσσικών ομάδων [93]. Το φαινόμενο αυτό δεν επηρεάζει μόνο την ποιότητα των νομικών υπηρεσιών που παρέχονται μέσω τεχνητής νοημοσύνης (AI), αλλά υπονομεύει και τη θεσμική ισότητα ενώπιον του νόμου, καθώς τα LLMs δύνανται έστω και ακούσια να ενισχύσουν την περιθωριοποίηση πολιτισμικών ή γλωσσικών μειονοτήτων.

Η διασφάλιση της δίκαιης πρόσβασης και της πολιτισμικής συμπερίληψης απαιτεί στρατηγικές πολλαπλών επιπέδων. Σε τεχνικό επίπεδο, αυτό σημαίνει επένδυση σε data augmentation και προσαρμοστικές τεχνικές εκπαίδευσης που ενισχύουν την παρουσία ποικίλων γλωσσών και πολιτισμικών προτύπων στο training corpus των LLMs. Σε οργανωτικό επίπεδο, απαιτείται διαρκής αξιολόγηση της προσβασιμότητας των τεχνολογικών λύσεων, με έμφαση στη σχεδίαση φιλικών διεπαφών (user-friendly interfaces), τη χρήση πολυγλωσσικών εργαλείων (multilingual tools) και την ανάπτυξη υποστηρικτικών μηχανισμών για ομάδες με περιορισμένη τεχνολογική κατάρτιση ή με ειδικές ανάγκες. Παράλληλα, η συνεχής διαβούλευση με κοινωνικούς, πολιτισμικούς και γλωσσικούς φορείς ενισχύει την ενσωμάτωση των αρχών της συμπερίληψης (inclusion) και της ισότητας στο τελικό σύστημα [93], [65].

Η πολιτισμική ποικιλομορφία (cultural diversity) αποτελεί κρίσιμο παράγοντα για τη νομιμοποίηση της τεχνητής νοημοσύνης (AI) στο δίκαιο. Η ανεπάρκεια ή αδιαφορία ως προς τις τοπικές παραδόσεις, τις αξίες και τις κοινωνικές ευαισθησίες μπορεί να οδηγήσει σε καχυποψία, αντίσταση στην υιοθέτηση των συστημάτων και τελικά σε υπονόμευση της αποτελεσματικότητάς τους. Η ενσωμάτωση της πολιτισμικής ευαισθησίας (cultural sensitivity) στη σχεδίαση, τη λειτουργία και την αξιολόγηση των LLMs οφείλει να είναι διαρκής και προσαρμοζόμενη στις κοινωνικές και θεσμικές εξελίξεις, διασφαλίζοντας ότι η τεχνητή νοημοσύνη (AI) λειτουργεί ως φορέας κοινωνικής συνοχής και όχι ως παράγοντας ενίσχυσης των υπαρχουσών ανισοτήτων.

Συνοψίζοντας, η δίκαιη πρόσβαση (fair access) και η πολιτισμική συμπερίληψη (cultural inclusion) δεν αποτελούν δευτερεύουσες παραμέτρους, αλλά θεμέλιο της θεσμικής και κοινωνικής αποδοχής των LLMs στο νομικό πεδίο, απαιτώντας συνεχή εγρήγορση, τεχνική καινοτομία και ουσιαστική κοινωνική διαβούλευση.

5.4 Συνεργασία Ανθρώπου-TN στο Νομικό Πλαίσιο

Η σύμπραξη ανθρώπου και τεχνητής νοημοσύνης στον νομικό τομέα αναδεικνύεται ως μία από τις πιο ενδιαφέρουσες και σύνθετες πτυχές της ψηφιακής μετάβασης του δικαίου. Η σταδιακή ενσωμάτωση των LLMs και των ευρύτερων συστημάτων τεχνητής νοημοσύνης (AI) δεν περιορίζεται πλέον σε βοηθητικούς ή αυτοματοποιημένους ρόλους, αλλά διαμορφώνει νέες δυναμικές στη λήψη αποφάσεων, στη διαχείριση δικαστικών υποθέσεων, στη νομική έρευνα και στη διαμεσολάβηση μεταξύ διαδίκων και θεσμών.

Η σχέση ανθρώπου και τεχνητής νοημοσύνης στον νομικό χώρο χαρακτηρίζεται από μια συνεχή διαπραγμάτευση μεταξύ αυτοματισμού και ανθρώπινης κρίσης. Από τη μία πλευρά, τα LLMs και τα συναφή εργαλεία μπορούν να ενισχύσουν δραστικά την αποτελεσματικότητα, την ταχύτητα και την ακρίβεια της επεξεργασίας μεγάλου όγκου νομικών πληροφοριών. Η δυνατότητα αδιάλειπτης ανάλυσης νομολογίας, κατασκευής προγνωστικών μοντέλων και παραγωγής προτάσεων επί σύνθετων νομικών ερωτημάτων μετασχηματίζει τον τρόπο εργασίας των νομικών επαγγελματιών και των δικαστηρίων, διευκολύνοντας τη λήψη ενημερωμένων και τεκμηριωμένων αποφάσεων. Από την άλλη πλευρά, η απόλυτη αυτοματοποίηση της νομικής λήψης αποφάσεων παραμένει ουτοπική – και πιθανώς επικίνδυνη. Η πολυπλοκότητα των νομικών ερωτημάτων, η ανάγκη ερμηνείας του πλαισίου, η στάθμιση της αρχής της επιείκειας (equity) και ο χειρισμός των ιδιαίτερων κοινωνικών, πολιτισμικών και προσωπικών διαστάσεων κάθε υπόθεσης, απαιτούν διαρκή ανθρώπινη εμπλοκή. Η εμπειρία, το ηθικό κριτήριο και η ικανότητα ενσυναίσθησης του ανθρώπινου παράγοντα δεν μπορούν να αναπαραχθούν πλήρως από αλγοριθμικά συστήματα, όσο εξελιγμένα κι αν είναι αυτά.[94]

Από την άλλη πλευρά, η απόλυτη αυτοματοποίηση της νομικής λήψης αποφάσεων παραμένει ουτοπική – και πιθανώς επικίνδυνη. Η πολυπλοκότητα των νομικών ερωτημάτων, η ανάγκη ερμηνείας του πλαισίου, η στάθμιση της αρχής της επιείκειας (equity) και ο χειρισμός των ιδιαίτερων κοινωνικών, πολιτισμικών και προσωπικών διαστάσεων κάθε υπόθεσης, απαιτούν διαρκή ανθρώπινη εμπλοκή. Η εμπειρία, το ηθικό κριτήριο και η ικανότητα ενσυναίσθησης του ανθρώπινου παράγοντα δεν μπορούν να αναπαραχθούν πλήρως από αλγοριθμικά συστήματα, όσο εξελιγμένα κι αν είναι αυτά [93], [65].

Το πιο ώριμο και αποτελεσματικό μοντέλο που προκύπτει από τη διεθνή εμπειρία είναι αυτό της υβριδικής συνεργασίας (hybrid collaboration)[65], όπου τα LLMs λειτουργούν ως υποστηρικτικά εργαλεία και όχι ως τελικοί κριτές. Ο ρόλος της τεχνητής νοημοσύνης (AI) επικεντρώνεται στη συλλογή, οργάνωση και αρχική επεξεργασία πληροφορίας, στην ανίχνευση προτύπων και στην υποβολή εναλλακτικών προτάσεων, ενώ η τελική κρίση, η ερμηνεία των στοιχείων και η στάθμιση των συνεπειών παραμένουν στον άνθρωπο. Αυτό το μοντέλο ενισχύει τη διαφάνεια, την αξιοπιστία και την κοινωνική αποδοχή των αποτελεσμάτων, ενώ ταυτόχρονα αξιοποιεί το πλήρες δυναμικό της τεχνητής νοημοσύνης (AI). Ωστόσο, η συνεργασία ανθρώπου-τεχνητής νοημοσύνης (human-AI collaboration) δεν στερείται προκλήσεων[93].

Η ελλιπής κατανόηση του τρόπου λειτουργίας των LLMs από τους τελικούς χρήστες, η υπέρμετρη εμπιστοσύνη στα outputs των συστημάτων ή, αντίθετα, η δυσπιστία και η άρνηση υιοθέτησης νέων τεχνολογιών, μπορούν να οδηγήσουν είτε σε τεχνοκρατικό φορμαλισμό είτε σε αδικαιολόγητη απόρριψη χρήσιμων εργαλείων. Οι φορείς του δικαίου καλούνται να επενδύσουν τόσο στην εκπαίδευση όσο και στην ενίσχυση της τεχνολογικής κουλτούρας, εξασφαλίζοντας ότι η ανθρώπινη κρίση παραμένει ο τελικός εγγυητής της δικαιοσύνης.

Η θεσμική θωράκιση της συνεργασίας ανθρώπου-τεχνητής νοημοσύνης (human-AI collaboration) περιλαμβάνει μηχανισμούς διαφάνειας (transparency), δυνατότητα εξηγήσιμης παρέμβασης (explainable intervention), σαφείς διαδικασίες ελέγχου και λογοδοσίας (accountability), καθώς και την αναγνώριση της δυνατότητας challenge και επανεξέτασης των αλγοριθμικών αποτελεσμάτων. Παράλληλα, απαιτείται συνεχής αξιολόγηση των κοινωνικών και ηθικών συνεπειών, με στόχο τη διαρκή ανανέωση του πλαισίου και τη διασφάλιση της θεσμικής ισορροπίας ανάμεσα στην καινοτομία και στη δικαιοσύνη. Συνολικά, η συνεργασία ανθρώπου και τεχνητής νοημοσύνης (AI) στο νομικό πεδίο δεν είναι απλώς ένα τεχνικό ή διαδικαστικό ζήτημα, αλλά θεμελιώδης παράμετρος της μετάβασης στη 'νέα δικαιοσύνη' (new justice), όπου η τεχνολογία καλείται να υπηρετήσει –και όχι να υποκαταστήσει– την ανθρώπινη κρίση και τις αξίες του κράτους δικαίου, όπως επισημαίνεται και στη σχετική βιβλιογραφία[93],[94],[55].

5.4.1 Υποστήριξη της Απόφασης έναντι Αυτοματοποίησης της Απόφασης

Η διάκριση μεταξύ της χρήσης των LLMs ως εργαλεία υποστήριξης της απόφασης (decision support systems) και της αυτοματοποιημένης λήψης αποφάσεων (decision automation) αποτελεί θεμελιώδη ζήτημα για τον ρόλο της τεχνητής νοημοσύνης (AI) στο δίκαιο. Η χρήση της τεχνητής νοημοσύνης (AI) ως εργαλείου υποστήριξης της ανθρώπινης απόφασης δίνει έμφαση στη συμβουλευτική και ενισχυτική λειτουργία της τεχνολογίας, όπου οι δικαστές, οι δικηγόροι και τα διοικητικά όργανα διατηρούν τον τελικό έλεγχο και την ευθύνη της κρίσης. Αντίθετα, στα συστήματα αυτοματοποίησης, οι αλγόριθμοι μπορούν να παράγουν αποτελέσματα χωρίς ουσιώδη ανθρώπινη παρέμβαση, με όλες τις τεχνικές, θεσμικές και ηθικές συνέπειες που αυτό συνεπάγεται [65], [55].

Στα συστήματα υποστήριξης, τα LLMs προσφέρουν εξαιρετικά εργαλεία για την ταχεία ανάλυση νομολογίας, την ανεύρεση προτύπων σε μεγάλα νομικά δεδομένα και την παροχή τεκμηριωμένων εναλλακτικών προτάσεων. Το ανθρώπινο υποκείμενο της απόφασης έχει τη δυνατότητα να αξιολογήσει τα outputs, να ερμηνεύσει το νομικό και κοινωνικό πλαίσιο κάθε υπόθεσης και να σταθμίσει μη ποσοτικοποιήσιμους παράγοντες όπως η αρχή της επιείκειας (equity) και οι κοινωνικές επιπτώσεις της απόφασης. Τα συστήματα υποστήριξης ενισχύουν τον επαγγελματία της δικαιοσύνης, χωρίς να υποκαθιστούν τη θεσμική και ηθική ευθύνη του [93], [94].

Αντίθετα, στα συστήματα αυτοματοποίησης, τα LLMs αναλαμβάνουν μεγαλύτερο μέρος – ή και το σύνολο – της διαδικασίας λήψης απόφασης. Σε τέτοια περιβάλλοντα, όπως παρατηρείται σε εφαρμογές έξυπνων δικαστηρίων (smart courts) στην Κίνα ή σε αλγοριθμικές αξιολογήσεις κινδύνου τύπου COMPAS στις ΗΠΑ, η ανθρώπινη παρέμβαση περιορίζεται συχνά σε τυπικό ρόλο επικύρωσης ή άρσης ακραίων σφαλμάτων [87], [85], [65]. Αυτό δημιουργ-

γεί κινδύνους συστημικής προκατάληψης bias, αδιαφάνειας (opacity), 'διαχυμένης ευθύνης' (diffused liability) και αδυναμίας challenge των αποφάσεων, με συνέπειες για τη θεσμική νομιμοποίηση και την κοινωνική αποδοχή της τεχνητής νοημοσύνης (AI) στη δικαιοσύνη [82], [65].

Η διεθνής εμπειρία καταδεικνύει τα όρια και τα πλεονεκτήματα κάθε προσέγγισης. Στην υπόθεση *Loomis v. Wisconsin* (2016) στις ΗΠΑ, το εργαλείο COMPAS χρησιμοποιήθηκε για την εκτίμηση κινδύνου υποτροπής κρατουμένων, επηρεάζοντας την ποινική απόφαση χωρίς να μπορεί ο κατηγορούμενος να αμφισβητήσει ή να κατανοήσει τα αλγοριθμικά κριτήρια. Το ανώτατο δικαστήριο αναγνώρισε πως η έλλειψη διαφάνειας (transparency) και η αδυναμία challenge των αλγοριθμικών outputs προσβάλλει βασικές αρχές δίκαιης δίκης και θεσμικής λογοδοσίας [85], [65].

Στην Κίνα, η αυτοματοποίηση της απόφασης εφαρμόζεται πιλοτικά σε 'έξυπνα δικαστήρια' (smart courts), κυρίως για την αυτοματοποιημένη επίλυση μικροδιαφορών (π.χ. traffic fines, απλές εμπορικές υποθέσεις). Αν και διασφαλίζεται ταχύτητα και αποσυμφόρηση, παραμένει το ζήτημα της εξατομικευμένης δικαιοσύνης και της δυνατότητας ουσιαστικής παρέμβασης του ανθρώπου, ενώ οι αποφάσεις γίνονται συχνά αντικείμενο κριτικής για την τυποποίηση και τον περιορισμό της θεσμικής ευελιξίας [87], [65].

Η Εσθονία προχωρά ένα βήμα πιο πέρα με το πρόγραμμα robot judge[85], όπου LLMs αξιολογούν μικρές διαφορές ως πρώτη βαθμίδα, αλλά υπάρχει πάντοτε η δυνατότητα ανατροπής της απόφασης από άνθρωπο, καθώς προβλέπονται audit trails και δυνατότητα ένστασης. Αυτό το υβριδικό μοντέλο ενσωματώνει πρακτικά τις ευρωπαϊκές αρχές διαμεσολάβησης ανθρώπινου παράγοντα human-in-the-loop και το δικαίωμα εξηγησιμότητας right to explanation" που επιβάλλει ο EU AI Act [84], [85].

Από τεχνική και θεσμική σκοπιά, η βέλτιστη ενσωμάτωση των μεγάλων γλωσσικών μοντέλων στο νομικό οικοσύστημα προϋποθέτει την υιοθέτηση συγκεκριμένων εγγυήσεων (safeguards) και τεχνικών υλοποίησης. Η εφαρμογή εργαλείων επεξηγήσιμης τεχνητής νοημοσύνης (explainable AI), όπως τα LIME, SHAP και η οπτικοποίηση της προσοχής (attention visualization), σε συνδυασμό με την πλήρη καταγραφή των εισροών και των παραμέτρων λειτουργίας του συστήματος (prompt archives, metadata logs), συμβάλλει ουσιαστικά στη διαφάνεια (transparency) και τη λογοδοσία (accountability). Επιπλέον, η επιβολή υποχρεωτικής ανθρώπινης ανασκόπησης (human-in-the-loop) και η δυνατότητα υποβολής ένστασης (challenge) σε κάθε αλγοριθμική απόφαση διασφαλίζουν την προστασία των δικαιωμάτων των διαδίκων και την ενίσχυση της εμπιστοσύνης στο σύστημα δικαιοσύνης[82],[55]. Παράλληλα, η συστηματική διενέργεια αξιολογήσεων κινδύνου και επιπτώσεων (risk and impact assessments), η δυναμική αξιολόγηση της ισοτιμίας (fairness) και της ανθεκτικότητας (robustness) των αλγορίθμων, καθώς και η τακτική επικαιροποίηση των πολιτικών ελέγχου (auditing policies), διασφαλίζουν ότι το σύστημα ανταποκρίνεται στις εξελισσόμενες απαιτήσεις της κοινωνίας και των θεσμικών πλαισίων[90],[91].

5.4.2 Επαγγελματική Ευθύνη

Η ενσωμάτωση της τεχνητής νοημοσύνης (artificial intelligence, AI) στο δικαστικό και νομικό περιβάλλον δημιουργεί νέα, σύνθετα ζητήματα επαγγελματικής ευθύνης για όλους

τους εμπλεκόμενους φορείς: δικαστές, δικηγόρους, διοικητικά στελέχη, αλλά και τεχνικούς, επιστήμονες δεδομένων και οργανισμούς παροχής υπηρεσιών τεχνητής νοημοσύνης (AI service providers). Σε ένα τέτοιο περιβάλλον, η παραδοσιακή λογική της επαγγελματικής ευθύνης—που βασιζόταν στην άμεση και ελέγξιμη σχέση μεταξύ πράξης και αποτελέσματος—αντιμετωπίζει τον κίνδυνο διαχυμένης ευθύνης (diffused liability) και συστημικής αβεβαιότητας [82], [93], [65].

Καταρχάς, η χρήση LLMs στο πλαίσιο της νομικής πρακτικής θέτει το ερώτημα της ευθύνης για τυχόν λανθασμένες, ελλιπείς ή μεροληπτικές αλγοριθμικές συστάσεις. Παρότι ο επαγγελματίας του δικαίου—δικαστής ή δικηγόρος—διατηρεί τυπικά τον τελικό λόγο στη λήψη της απόφασης, η αυξανόμενη πολυπλοκότητα και ο όγκος των αποτελεσμάτων των LLMs καθιστούν ολοένα δυσκολότερη την πλήρη αξιολόγηση και επαλήθευση κάθε προτεινόμενης λύσης. Αυτό εγκυμονεί τον κίνδυνο μετακύλισης της ευθύνης σε αλγοριθμικά συστήματα ή στους προγραμματιστές, ιδίως σε περιβάλλοντα αυτοματοποίησης, όπου η παρέμβαση του ανθρώπινου παράγοντα περιορίζεται [82], [65].

Οι οργανισμοί και οι πάροχοι τεχνητής νοημοσύνης (AI providers) καλούνται να διαμορφώσουν σαφείς πολιτικές ευθύνης και διαφάνειας, με βάση τις αρχές της τεκμηρίωσης, της auditability και της ανιχνευσιμότητας (traceability) κάθε αλγοριθμικής διαδικασίας. Η υποχρέωση τήρησης αναλυτικών audit trails, η επαρκής ενημέρωση του τελικού χρήστη για τη φύση και τους περιορισμούς του LLM, καθώς και η δυνατότητα challenge κάθε αλγοριθμικής εισήγησης, αποτελούν κρίσιμες ασφαλιστικές δικλίδες για την κατανομή της ευθύνης [85], [65]. Το ευρωπαϊκό πλαίσιο (EU AI Act) και η διεθνής πρακτική απαιτούν οι τελικοί χρήστες να έχουν πάντα τη δυνατότητα τεκμηριωμένης ανασκόπησης των outputs και να μην εξαρτώνται αποκλειστικά από αλγοριθμικές εισηγήσεις, ειδικά σε υποθέσεις με σημαντικές κοινωνικές ή οικονομικές επιπτώσεις [84], [85].

Στο πεδίο της επαγγελματικής δεοντολογίας, οι δικηγόροι και οι δικαστές οφείλουν να διασφαλίζουν τη διαφύλαξη της ανεξαρτησίας και της αμεροληψίας τους, ακόμα και όταν αξιοποιούν εργαλεία τεχνητής νοημοσύνης (AI tools) στην καθημερινή πρακτική. Αυτό συνεπάγεται ενεργή αξιολόγηση της αξιοπιστίας, της τεχνικής επάρκειας και των πιθανών προκαταλήψεων (biases) των εργαλείων που χρησιμοποιούνται, καθώς και διαρκή εκπαίδευση σε νέες τεχνολογίες και στις τεχνικές εξηγησιμότητας και ελέγχου [82], [93]. Παράλληλα, θεσμικά κείμενα (όπως οι κώδικες δεοντολογίας δικηγόρων και δικαστών σε ΗΠΑ, ΕΕ και Κίνα) εμπλουτίζονται πλέον με διατάξεις που προβλέπουν ευθύνη για την αλόγιστη ή ανεπαρκώς τεκμηριωμένη χρήση αλγοριθμικών συστημάτων [87], [93].

Ειδικό ζήτημα αποτελεί η ευθύνη των τεχνικών (προγραμματιστών, επιστημόνων δεδομένων, οργανισμών). Η λανθασμένη παραμετροποίηση, η ανεπαρκής τεκμηρίωση ή η αδυναμία διασφάλισης της διαφάνειας (transparency) και της αμεροληψίας (impartiality) του συστήματος μπορεί να επιφέρει νομικές συνέπειες, ιδίως σε περιπτώσεις "λογισμικού υψηλού ρίσκου" (high-risk software). Η υιοθέτηση πρακτικών responsible AI, η τήρηση προτύπων auditing (ISO/IEC, IEEE), η ενσωμάτωση πρακτικών human-in-the-loop και η τεχνική συμμόρφωση με το εκάστοτε ρυθμιστικό πλαίσιο αποτελούν πλέον αναγκαίες προϋποθέσεις επαγγελματικής ευθύνης και συμμόρφωσης [85], [55].

Η διεθνής εμπειρία δείχνει ότι η έλλειψη σαφών πλαισίων ευθύνης ενέχει σοβαρούς κινδύνους στη γνωστή υπόθεση COMPAS όπως έχουμε αναφέρει, η αδυναμία ελέγχου και

διαφάνειας (transparency) του αλγορίθμου οδήγησε σε δυσκολίες απόδοσης ευθύνης τόσο στους developers όσο και στους τελικούς χρήστες, αναδεικνύοντας την ανάγκη θεσμικής θωράκισης της επαγγελματικής ευθύνης σε όλα τα στάδια του κύκλου ζωής της τεχνητής νοημοσύνης (AI lifecycle) [85], [65]. Αντίστοιχα, τα "smart courts" στην Κίνα δείχνουν πως η μετατόπιση της ευθύνης στον κρατικό φορέα ή στον προγραμματιστή γείρει σημαντικά ερωτήματα ως προς τη λογοδοσία (accountability) και τον σεβασμό των δικαιωμάτων του πολίτη [87], [65].

Συνοψίζοντας, η επαγγελματική ευθύνη στην εποχή της νομικής τεχνητής νοημοσύνης είναι σύνθετη, πολυεπίπεδη και διεπιστημονική. Απαιτεί σαφή κατανομή ρόλων και ευθυνών μεταξύ νομικών επαγγελματιών και τεχνικών φορέων, συνεχή εκπαίδευση, διάχυση της διαφάνειας (transparency) και των auditing πρακτικών, και διαρκή συμμόρφωση με το ρυθμιστικό και ηθικό πλαίσιο κάθε έννομης τάξης.

5.4.3 Μοντέλα Αλληλεπίδρασης Δικαστή-Αλγορίθμου

Η εξέλιξη των μοντέλων αλληλεπίδρασης μεταξύ δικαστή και τεχνητής νοημοσύνης δεν συνιστά απλώς μια τεχνολογική μετάβαση, αλλά αναδεικνύει θεμελιώδη ερωτήματα για το ρόλο, την ανεξαρτησία, και τη θεσμική λογοδοσία του ανθρώπινου παράγοντα στο νέο, υβριδικό οικοσύστημα της απονομής δικαιοσύνης. Η διεθνής εμπειρία αποδεικνύει ότι η υιοθέτηση κάθε μοντέλου σχετίζεται με διαφορετικό βαθμό καινοτομίας, θεσμικής ασφάλειας και κοινωνικής αποδοχής.

Σε αυτό το μοντέλο, ο αλγόριθμος παρέχει πληθώρα πληροφοριών: αναλύει δικαστικά προηγούμενα, υποδεικνύει σχετικά νομικά πρότυπα, συνθέτει και συγκρίνει διαφορετικές εκδοχές νομολογίας ή αναγνωρίζει μοτίβα σε μεγάλα σύνολο δεδομένων. Το παραγόμενο αποτέλεσμα παρουσιάζεται στον δικαστή ως 'πρόταση' ή 'υπόδειξη', χωρίς δεσμευτική ισχύ. Ο άνθρωπος παράγοντας παραμένει κυρίαρχος:

- Αξιολογεί, αμφισβητεί ή αγνοεί το αλγοριθμικό αποτέλεσμα βάσει της δικής του εμπειρίας, κρίσης και κατανόησης της υπόθεσης
- Ενισχύει την ταχύτητα και πληρότητα της νομικής έρευνας, επιτρέποντας την κάλυψη πολύ μεγαλύτερης πληροφορίας από ό,τι ήταν μέχρι πρότινος εφικτό.
- Εξασφαλίζει τη διατήρηση του ανθρώπινου κριτηρίου ως "τελικού φραγμού" (φινάλ σαφεγυαρδ) απέναντι σε τεχνικά λάθη ή βίας των αλγορίθμων
- Προάγει τη θεσμική διαφάνεια και εμπιστοσύνη, καθώς ο δικαστής μπορεί να ελέγξει πλήρως τον τρόπο ενσωμάτωσης των αλγοριθμικών συστάσεων στη δικανική του σκέψη

Στις ΗΠΑ και την Ευρώπη, η πλειοψηφία των εργαλείων νομικής βοήθειας (legal research tools) λειτουργεί σε αυτό το μοντέλο (π.χ. Westlaw Edge, Lexis+ με AI-powered search). Ο δικαστής χρησιμοποιεί το LLMs ως έξυπνο βοηθό για να εντοπίσει προηγούμενα ή ερμηνείες, αλλά η τελική ερμηνεία και η ευθύνη παραμένουν δικές του [93], [55].

Το υβριδικό μοντέλο σηματοδοτεί μια πιο στενή συνεργασία: ο αλγόριθμος λαμβάνει προκαταρκτικές αποφάσεις ή παράγει προσχέδια αποφάσεων, τα οποία ο δικαστής οφείλει

να ανασκοπήσει, να σχολιάσει ή να εγκρίνει. Σε αυτή τη διαδικασία, ο δικαστής λειτουργεί ως αξιολογητής τόσο της ποιότητας του αλγοριθμικού συλλογισμού (reasoning) όσο και της ορθότητας του αποτελέσματος. Ενσωματώνονται μηχανισμοί ανιχνευσιμότητας (traceability), ιστορικού ελέγχου (audit trails) και δυνατότητα αμφισβήτησης (challenge) όχι μόνο από τον ίδιο τον δικαστή αλλά και, υπό προϋποθέσεις, από τους διαδίκους.

Τα πλεονεκτήματα του υβριδικού μοντέλου είναι πολλαπλά. Πρώτον, αυξάνει την αποδοτικότητα της δικαστικής λειτουργίας, καθώς αυτοματοποιεί μεγάλο μέρος των επαναλαμβανόμενων διαδικασιών, διατηρώντας όμως τον ουσιαστικό ανθρώπινο έλεγχο. Επιπλέον, υλοποιεί στην πράξη την αρχή του human-in-the-loop, όπως απαιτεί ο ευρωπαϊκός νομοθέτης (AI Act), διασφαλίζοντας ότι ο άνθρωπος έχει τον τελικό λόγο σε κρίσιμες αποφάσεις [85]. Τέλος, διευκολύνει τη συνεχή βελτίωση των αλγοριθμικών συστημάτων, μέσω βρόχων ανατροφοδότησης (feedback loops) και εποπτευόμενης μάθησης (supervised learning) από τις αποφάσεις και τις διορθώσεις των ίδιων των δικαστών.

Η Εσθονία[95] αποτελεί πρωτοπόρο περίπτωση, καθώς το πρόγραμμα ρομποτικός δικαστής (robot judge) προβλέπει την αυτόματη εκδίκαση μικροδιαφορών. Ωστόσο, κάθε απόφαση υπόκειται σε έλεγχο και μπορεί να ανατραπεί από ανθρώπινο δικαστή, ενώ κάθε αλγοριθμικό αποτέλεσμα (output) καταγράφεται με ιστορικό ελέγχου (audit trail) για εκ των υστέρων (ex post) επανεξέταση

Στο μοντέλο της αυτοματοποιημένης απόφασης (Automated/Default Decision), η απόφαση λαμβάνεται σχεδόν αποκλειστικά από τα LLMs ή το ευρύτερο σύστημα τεχνητής νοημοσύνης, κυρίως σε τυποποιημένες, χαμηλού ρίσκου υποθέσεις, όπως είναι τα traffic fines ή οι microclaims. Ο ρόλος του δικαστή περιορίζεται στη δυνατότητα παρέμβασης μόνο σε περίπτωση άσκησης ένστασης (challenge) ή αν ανακύψει ζήτημα νομιμότητας ή σφάλματος.

Τα βασικά πλεονεκτήματα του μοντέλου περιλαμβάνουν την τεράστια εξοικονόμηση χρόνου και πόρων, καθώς και τη σημαντική μείωση του φόρτου εργασίας για τα δικαστήρια, ιδιαίτερα στις επαναλαμβανόμενες και τυποποιημένες υποθέσεις. Επιπλέον, η διαδικασία γίνεται σαφέστερη και περισσότερο τυποποιημένη, διευκολύνοντας την ταχεία διεκπεραίωση.

Ωστόσο, το συγκεκριμένο μοντέλο συνοδεύεται από σοβαρούς κινδύνους και προκλήσεις. Η κύρια ανησυχία αφορά την απομάκρυνση από την αρχή της εξατομικευμένης κρίσης και το δικαίωμα ακρόασης κάθε διάδικου, ενώ η δυσκολία απόδοσης ευθύνης σε περιπτώσεις σφαλμάτων ή προκαταλήψεων παραμένει άλυτο ζήτημα. Επίσης, σε πολλά συστήματα περιορίζεται η δυνατότητα ένστασης, ειδικά αν δεν υπάρχουν επαρκείς θεσμικοί μηχανισμοί προστασίας [87], [65].

Χαρακτηριστικό παράδειγμα αποτελεί η Κίνα, όπου τα έξυπνα δικαστήρια διαχειρίζονται αυτόματα χιλιάδες υποθέσεις μικρών διαφορών, με τη συμμετοχή του δικαστή να περιορίζεται σε περιπτώσεις ένστασης ή όταν ανακύπτει σοβαρό νομικό ερώτημα [87], [65].

Για την ενίσχυση της αλληλεπίδρασης ανθρώπου και συστημάτων τεχνητής νοημοσύνης, υιοθετούνται μια σειρά από εξειδικευμένες τεχνικές και πρακτικές. Πρώτα, η επεξηγήσιμη τεχνητή νοημοσύνη (Explainable AI) δίνει έμφαση στην παροχή οπτικοποιήσεων (visualizations), ανάλυσης βήμα-βήμα (step-by-step reasoning) και συστηματικής παρακολούθησης των πηγών (citation tracking), ώστε ο δικαστής να μπορεί να κατανοεί το σκεπτικό της αλγοριθμικής σύστασης [55].

Επιπλέον, εφαρμόζονται μητρώα ελέγχου και αρχείων καταγραφής (audit trails and logs), τα οποία διασφαλίζουν την αυτόματη καταγραφή όλων των εισροών, των παραμέτρων και των ληφθέντων αποφάσεων. Αυτό επιτρέπει τον μετέπειτα έλεγχο της νομιμότητας και της αιτιολόγησης των συστάσεων [65]. Σημαντικό ρόλο έχουν επίσης οι κύκλοι ανατροφοδότησης (feedback loops), όπου οι επιλογές και τα σχόλια των δικαστών τροφοδοτούν το σύστημα, επιτρέποντας την εποπτευόμενη προσαρμογή (supervised fine-tuning) με διαρκή νομική επίβλεψη.

Συνολικά, η ανάλυση των μοντέλων αλληλεπίδρασης (interaction models) επιβεβαιώνει ότι το μέλλον της δικαιοσύνης δεν ανήκει ούτε στην πλήρη αυτοματοποίηση (full automation) ούτε στη διατήρηση της παραδοσιακής πρακτικής, αλλά στη δυναμική συνεργασία ανθρώπινης κρίσης και αλγοριθμικής καινοτομίας. Η επιλογή μοντέλου εξαρτάται τόσο από το θεσμικό και κοινωνικό πλαίσιο όσο και από το είδος της υπόθεσης.

Τα συμβουλευτικά και υβριδικά μοντέλα, με ενισχυμένη ανθρώπινη εποπτεία (human oversight), διασφαλίζουν μεγαλύτερη διαφάνεια (transparency), λογοδοσία (accountability) και κοινωνική αποδοχή, αξιοποιώντας ταυτόχρονα τις δυνατότητες της τεχνητής νοημοσύνης (AI) για ανάλυση δεδομένων και ενίσχυση της αποδοτικότητας (efficiency). Αντίθετα, τα πλήρως αυτοματοποιημένα μοντέλα (fully automated models) απαιτούν αυστηρούς περιορισμούς ως προς τη χρήση και τη δυνατότητα ελέγχου.

Η πραγματική πρόκληση για τα μελλοντικά δικαστικά συστήματα είναι να ενισχύσουν τη συνεργασία ανθρώπου και τεχνητής νοημοσύνης (AI) μέσα από διαφανείς, επαληθεύσιμους (verifiable) και θεσμικά προστατευμένους μηχανισμούς αλληλεπίδρασης, διασφαλίζοντας τη δικαιοσύνη και τον σεβασμό των θεμελιωδών δικαιωμάτων.

5.5 Προκλήσεις Εφαρμογής Μεγάλων Γλωσσικών Μοντέλων στο Ελληνικό Νομικό Πλαίσιο

Η εφαρμογή Μεγάλων Γλωσσικών Μοντέλων στο ελληνικό νομικό σύστημα παρουσιάζει μία σειρά από πολυεπίπεδες προκλήσεις, που σχετίζονται τόσο με τη γλωσσική ιδιαιτερότητα όσο και με τις θεσμικές και υποδομικές συνθήκες της χώρας. Παρά την πρόοδο της τεχνητής νοημοσύνης στο πεδίο της νομικής πληροφορικής, η ελληνική περίπτωση παραμένει εν μέρει απομονωμένη, καθώς χαρακτηρίζεται από σημαντικές ελλείψεις σε πόρους, δεδομένα και τυποποιημένα θεσμικά πλαίσια.

Η ελληνική γλώσσα κατατάσσεται στις γλώσσες χαμηλών πόρων (low-resource languages) για τα περισσότερα παγκόσμια μοντέλα, με αποτέλεσμα η εκπαίδευση, η ρύθμιση (fine-tuning) και η εφαρμογή τους να παρουσιάζουν αισθητά χαμηλότερη απόδοση σε σύγκριση με τα αγγλικά ή άλλες κυρίαρχες γλώσσες [96]. Η περιορισμένη διαθεσιμότητα νομικών δεδομένων υψηλής ποιότητας στα ελληνικά, σε συνδυασμό με την απουσία εθνικών πρωτοβουλιών για ανοικτά και καθαρά νομικά σύνολα δεδομένων (legal datasets), εμποδίζει τη δημιουργία εξειδικευμένων μοντέλων για το ελληνικό δίκαιο.

Επιπλέον, όπως επιβεβαιώνεται σε πρόσφατες μελέτες, ακόμα και βασικές εργασίες όπως η αυτόματη περίληψη δικαστικών αποφάσεων παρουσιάζουν ιδιαίτερες δυσκολίες στα ελληνικά, με τα υπάρχοντα μοντέλα να αποτυγχάνουν στη διατήρηση της ουσίας και της τεκμη-

ρίωσης του συλλογισμού [97]. Η ανάγκη για ανάπτυξη στοχευμένων εργαλείων επεξεργασίας νομικού κειμένου στα ελληνικά είναι κρίσιμη, καθώς η γλωσσική ακρίβεια αποτελεί προϋπόθεση για κάθε είδους νομική εφαρμογή.

Μία ακόμα σημαντική πρόκληση αφορά την απουσία δημόσιας και θεσμικά εγκεκριμένης πρόσβασης σε εκτεταμένα σύνολα νομικών δεδομένων, καθώς μεγάλο μέρος της ελληνικής νομολογίας παραμένει είτε μη δημοσιευμένο είτε μη δομημένο με κατάλληλα μεταδεδομένα [98]. Το γεγονός αυτό δυσχεραίνει τόσο την ποιοτική εκπαίδευση των LLMs όσο και τη δυνατότητα ελέγχου, διαφάνειας και επαναληψιμότητας στις αποφάσεις τους.

Η έλλειψη επαρκούς σχολιασμένου υλικού (annotated datasets) σε συνδυασμό με τη μη τυποποιημένη γλωσσική μορφή των αποφάσεων – που συχνά περιέχουν αποσπασματικές ή ιδιωματικές εκφράσεις – δημιουργεί εμπόδια στη γενίκευση των μοντέλων, καθώς και στην ασφαλή χρήση τους για υποβοήθηση στη νομική πρακτική ή στη σύνταξη δικαστικών εγγράφων.

Το ελληνικό νομικό σύστημα δεν ακολουθεί δογματικά το αγγλοσαξονικό πρότυπο δεσμευτικού δεδικασμένου (binding precedent) αλλά ενσωματώνει στοιχεία από τη ρωμαϊκή-ηπειρωτική παράδοση. Αυτό σημαίνει ότι η νομολογία έχει περιορισμένη δεσμευτικότητα και μεγαλύτερη διακύμανση, γεγονός που δημιουργεί δυσκολίες για μοντέλα που στηρίζονται σε προβλεπτικότητα, συνέπεια και προηγούμενες αποφάσεις [99].

Επιπλέον, δεν υπάρχει επίσημα αναγνωρισμένο πρότυπο για τη μορφή και τη δομή των αποφάσεων, με αποτέλεσμα να παρατηρείται έλλειψη συνέπειας στη διατύπωση των πραγματικών περιστατικών, των νομικών επιχειρημάτων και των διατάξεων που εφαρμόζονται. Αυτό δυσχεραίνει την εξαγωγή συμπερασμάτων από τα LLMs και μειώνει την αξιοπιστία τους σε πρακτικές εφαρμογές [5].

Αν και η Ελλάδα οδεύει προς την εφαρμογή του Κανονισμού Τεχνητής Νοημοσύνης της ΕΕ (AI Act), σε εθνικό επίπεδο παραμένει κενό ως προς τη θέσπιση ειδικών ρυθμίσεων για την αξιοποίηση των LLMs σε δικαστικά περιβάλλοντα. Δεν υπάρχουν ακόμη μηχανισμοί πιστοποίησης, διαφάνειας ή ευθύνης που να καλύπτουν την ειδική χρήση αυτών των μοντέλων για τη σύνταξη νομικών γνωμοδοτήσεων, για την αυτοματοποιημένη επεξεργασία υποθέσεων ή για τη σύνταξη προσχεδίων αποφάσεων.

Η ανάγκη θεσμικής προσαρμογής αναδεικνύεται ως επείγουσα, όχι μόνο για λόγους διαφάνειας και λογοδοσίας αλλά και για την ενίσχυση της εμπιστοσύνης των επαγγελματιών του δικαίου απέναντι σε εργαλεία τεχνητής νοημοσύνης. Η εμπειρία άλλων χωρών δείχνει ότι η απουσία σαφούς πλαισίου οδηγεί είτε σε υπερεξάρτηση από αδιαφανείς τεχνολογίες είτε σε συντηρητική απόρριψη της καινοτομίας.

5.6 Συμπεράσματα

Βλέποντας τη μεγάλη εικόνα, γίνεται φανερό ότι η συζήτηση για το ρυθμιστικό τοπίο της τεχνητής νοημοσύνης στο δίκαιο ξεπερνά κατά πολύ τα στενά όρια των νομοθετικών πινακίδων και των θεσμικών ανταγωνισμών. Ναι, η Ευρώπη επενδύει σε αυστηρούς θεσμούς, οι ΗΠΑ στηρίζονται περισσότερο στο γρσε λαω και στην αγορά, ενώ η Κίνα υιοθετεί κρατικά κεντρικό έλεγχο και εποπτεία. Καθεμία από αυτές τις προσεγγίσεις φανερώνει τόσο τις πολιτισμικές όσο και τις θεσμικές προτεραιότητες κάθε κοινωνίας, αλλά και τα δομικά τους

ελλείμματα.

Το διακύβευμα όμως δεν είναι απλώς ποιος θα θέσει τους όρους του παιχνιδιού στην επόμενη μέρα της τεχνητής νοημοσύνης, αλλά πώς θα διασφαλιστεί η ουσία: να παραμείνει ο άνθρωπος στο κέντρο της δικαιοσύνης και η τεχνολογία να λειτουργεί ως ενισχυτής της ισότητας και της διαφάνειας, όχι ως πηγή νέων αδικιών ή κοινωνικών ρωγμών.

Η πρόκληση που προβάλλει είναι διπλή: αφενός, να σχεδιαστούν κανόνες που δεν πνίγουν την καινοτομία, αλλά θέτουν σαφή όρια στις "σκοτεινές ζώνες" της αυτοματοποίησης· αφετέρου, να καλλιεργηθεί κουλτούρα λογοδοσίας και διαφάνειας, έτσι ώστε τα ίδια τα συστήματα να μπορούν να ελεγχθούν, να αιτιολογήσουν τις αποφάσεις τους και να διορθώσουν τα λάθη τους — πριν αυτά γίνουν κοινωνικά ή θεσμικά μη αναστρέψιμα. Στο φινάλε, η μάχη δεν είναι μεταξύ Ε.Ε., ΗΠΑ και Κίνας. Η πραγματική μάχη είναι για μια τεχνητή νοημοσύνη που θα σέβεται τον άνθρωπο και θα υπηρετεί τη δικαιοσύνη με σύγχρονα εργαλεία, μετρήσιμη λογοδοσία και διαρκή δυνατότητα ελέγχου. Όσο πιο γρήγορα το κατανοήσουμε — και το απαιτήσουμε — τόσο πιο κοντά θα βρεθούμε σε ένα σύστημα όπου οι κανόνες, οι άνθρωποι και η τεχνολογία συνυπάρχουν με ουσιαστικό, δημοκρατικό και αποτελεσματικό τρόπο.

Μελλοντικές Κατευθύνσεις

Η ταχεία ενσωμάτωση των μεγάλων γλωσσικών μοντέλων (LLMs) και της τεχνητής νοημοσύνης στο νομικό οικοσύστημα μετασχηματίζει ραγδαία τα θεσμικά, τεχνικά και κοινωνικά πλαίσια του δικαίου. Καθώς οι εφαρμογές επεκτείνονται από την αυτοματοποίηση της νομικής έρευνας μέχρι τη λήψη δικαστικών αποφάσεων και τη βελτίωση της πρόσβασης στη δικαιοσύνη, αναδύονται νέες ευκαιρίες αλλά και σύνθετες προκλήσεις για επιστήμονες, επαγγελματίες και θεσμικούς φορείς [7], [100].

Το κεφάλαιο αυτό εστιάζει στις κύριες κατευθύνσεις που θα διαμορφώσουν το μέλλον της τεχνητής νοημοσύνης στο δίκαιο, δίνοντας έμφαση στη μεταμόρφωση της νομικής εκπαίδευσης και στην ανάπτυξη διεπιστημονικών δεξιοτήτων, στις τεχνικές εξελίξεις που σχετίζονται με την ανθεκτικότητα, την εξηγήσιμη λειτουργία, την προσαρμογή στον τομέα και την πολυγλωσσία των LLMs, στην ανάπτυξη υβριδικών συστημάτων και στη συνεργασία ανθρώπου και μηχανής σε δικαστικές λειτουργίες, καθώς και στις νέες προσεγγίσεις για τη διασφάλιση της προσβασιμότητας, της δικαιοσύνης και της συμπερίληψης σε ένα ψηφιακά διαρκώς εξελισσόμενο περιβάλλον.

Αυτές οι τάσεις συνθέτουν το ρυθμιστικό, κοινωνικό και τεχνολογικό πλαίσιο που θα καθορίσει τον ρόλο της τεχνητής νοημοσύνης στη νομική πράξη τα επόμενα χρόνια, τονίζοντας την ανάγκη για συνεχή διαβούλευση, τεχνική προσαρμοστικότητα και ηθική εγρήγορση.

6.1 Αντίκτυπος στην νομική εκπαίδευση και τις επαγγελματικές δεξιότητες

Η ενσωμάτωση της τεχνητής νοημοσύνης και των μεγάλων γλωσσικών μοντέλων στον νομικό τομέα φέρνει ριζικές αλλαγές στις δομές και τα περιεχόμενα της νομικής εκπαίδευσης, αλλά και στις απαιτήσεις για τη δια βίου μάθηση και την επαγγελματική προσαρμοστικότητα των νομικών επαγγελματιών. Η διεθνής επιστημονική συζήτηση πλέον αναγνωρίζει ότι ο παραδοσιακός δικηγόρος του 20ού αιώνα δεν επαρκεί στο νέο νομικό περιβάλλον που διαμορφώνεται από τις εξελίξεις στην τεχνολογία, τη διαχείριση δεδομένων, τη ρομποτική και την αλγοριθμική λήψη αποφάσεων. Η σύγχρονη νομική εκπαίδευση καλείται να εφοδιάσει τους φοιτητές όχι μόνο με εμβάθυνση στη νομική ανάλυση, αλλά και με γνώσεις μηχανικής μάθησης, αλγοριθμικής λογικής, διαχείρισης δεδομένων, ψηφιακής δεοντολογίας και κριτικής αξιολόγησης της τεχνητής νοημοσύνης [7],[100], [93], [101].

Αυτή η μετατόπιση της έμφασης συνοδεύεται από την υιοθέτηση νέων παιδαγωγικών με-

θόδων και τη σταδιακή εισαγωγή μαθημάτων τεχνολογίας στα βασικά προγράμματα σπουδών. Όπως επισημαίνει η Ryan και οι συνεργάτες της [7], η πρακτική εξοικείωση με εργαλεία τεχνητής νοημοσύνης, όπως το ChatGPT, Grokx, οδηγεί στην ανάπτυξη νέων δεξιοτήτων ανάλυσης, αυτονομίας και υπευθυνότητας, και βοηθά τους φοιτητές να κατανοούν έμπρακτα τόσο τις δυνατότητες όσο και τα όρια των συστημάτων της τεχνητής νοημοσύνης. Σε πολλά πανεπιστήμια, ιδίως στην Ευρώπη και την Αμερική, παρατηρείται στροφή σε προγράμματα Law and Technology, με διατμηματικά σεμινάρια, εργαστήρια νομικής πληροφορικής καθώς και διαδραστικά προθερς που φέρνουν τους φοιτητές σε επαφή με πραγματικά σύνολα δεδομένων (Datasets), ψηφιακές πλατφόρμες και αλγοριθμικά εργαλεία [100], [102].

Παράλληλα, η αυξανόμενη χρήση της τεχνητής νοημοσύνης στη νομική κλινική εκπαίδευση μεταμορφώνει την πρακτική προσέγγιση των φοιτητών στο δίκαιο. Οι νομικές κλινικές, όπως δείχνουν case studies σε Ηνωμένο Βασίλειο, Ισραήλ, Ινδία και Αυστραλία, αξιοποιούν την τεχνητή νοημοσύνη όχι μόνο για την επίλυση πραγματικών υποθέσεων (legal problem-solving), αλλά και για τη διδασκαλία δεξιοτήτων ψηφιακής έρευνας, διαχείρισης δεδομένων πελατών, αξιολόγησης αυτοματοποιημένων νομικών γνωμοδοτήσεων και ελέγχου ποιότητας των αποτελεσμάτων των LLMs [7], [100], [101]. Οι φοιτητές εξασκούνται στην αναγνώριση προκαταλήψεων (bias), στην ερμηνεία και στον έλεγχο των αποτελεσμάτων που παράγονται από αλγοριθμικά συστήματα, αναπτύσσοντας έτσι μια διττή προσέγγιση: ψηφιακή δεξιότητα αλλά και διαρκή κριτική εγρήγορση για τις ηθικές και νομικές διαστάσεις της τεχνητής νοημοσύνης.

Η βιβλιογραφία υπογραμμίζει ότι η απλή χρήση των μεγάλων γλωσσικών μοντέλων δεν αρκεί. Απαιτείται η εσωτερική κατανόηση των βασικών αρχών λειτουργίας των συστημάτων αυτών, ώστε οι μελλοντικοί νομικοί να μπορούν να αξιολογούν ρεαλιστικά τα πλεονεκτήματα και τους κινδύνους κάθε τεχνολογικής λύσης [93], [102]. Έτσι, η εκπαίδευση επικεντρώνεται πλέον στη διαμόρφωση ενός 'ψηφιακά αυτόνομου νομικού', ικανού να ανταποκριθεί στις μεταβαλλόμενες απαιτήσεις των νομικών επαγγελματιών διεθνώς.

Ένα επιπλέον κρίσιμο σημείο είναι η ανάγκη για δια βίου μάθηση. Η ταχύτατη πρόοδος της τεχνολογίας και η διαρκής αλλαγή των εργαλείων και μεθόδων δημιουργούν την απαίτηση οι επαγγελματίες να παραμένουν διαρκώς ενημερωμένοι για τις εξελίξεις στη νομική πληροφορική, να παρακολουθούν εκπαιδευτικά προγράμματα, συνέδρια και διεθνή συνέδρια, να συμμετέχουν σε ομάδες εργασίας και να υιοθετούν ένα μινδσετ συνεχούς προσαρμογής [7]. Πολλές νομικές σχολές ενσωματώνουν micro-credentials, ειδικά MOOCs, training σε real-world πλατφόρμες (π.χ. document automation, case management), ενώ οι επαγγελματικές ενώσεις προσφέρουν πιστοποιήσεις σε legal tech, data privacy, digital forensics κ.ά.

Αξιοσημείωτη είναι και η επίδραση της τεχνητής νοημοσύνης στα soft skills που θεωρούνται πλέον ουσιώδεις: επικοινωνία μέσω ψηφιακών μέσων, διαχείριση ψηφιακών ομάδων, προσαρμοστικότητα, διαπολιτισμική κατανόηση, καθώς και η ικανότητα διαπραγμάτευσης με 'ψηφιακά εργαλεία' (π.χ. chatbots, online mediation). Τα μεγάλα δικηγορικά γραφεία και οι διεθνείς οργανισμοί δίνουν ιδιαίτερη έμφαση στη συνδυαστική κατάρτιση, με τις αγγελίες εργασίας να απαιτούν πλέον, εκτός από άριστη γνώση δικαίου, τεχνικές δεξιότητες, εξοικείωση με βάσεις δεδομένων νομικού περιεχομένου legal databases, NLP εργαλεία, AI review software και νομική τεκμηρίωση που αξιοποιεί αυτοματοποιημένα συστήματα [102],

[101], [103].

Τέλος, ο κοινωνικός και θεσμικός αντίκτυπος της ενσωμάτωσης της τεχνητής νοημοσύνης στη νομική εκπαίδευση είναι ιδιαίτερα έντονος σε χώρες με περιορισμένη πρόσβαση σε πόρους, όπου η αξιοποίηση της τεχνητής νοημοσύνης μπορεί να γεφυρώσει το εκπαιδευτικό χάσμα, να παρέχει πρόσβαση σε ποιοτική γνώση και να ενισχύσει την κοινωνική κινητικότητα [7], [101].

6.1.1 Ανάπτυξη διεπιστημονικών δεξιοτήτων

Η ανάγκη για διεπιστημονικές δεξιότητες στο νομικό επάγγελμα αποτυπώνεται έμπρακτα στην ποικιλία των νέων, τεχνολογικά προσανατολισμένων εκπαιδευτικών προγραμμάτων που υλοποιούν πανεπιστήμια σε όλο τον κόσμο [7], [100]. Πέρα από τη θεωρητική εξοικείωση, η πραγματική μετάβαση προς τον νομικό της ψηφιακής εποχής επιτυγχάνεται μέσα από πρακτική μάθηση, διατμηματικά ερευνητικά projects και ενσωμάτωση εργαλείων αιχμής στη διδακτική διαδικασία.

Συγκεκριμένα, σχολές νομικής όπως το Stanford, το University College London (UCL), το Harvard και το KU Leuven έχουν εντάξει στα προγράμματά τους εργαστήρια όπου οι φοιτητές εργάζονται με πραγματικά σύνολα δεδομένων, δοκιμάζουν αλγορίθμους ανάλυσης νομικών εγγράφων (legal document classification), εφαρμόζουν τεχνικές επεξεργασίας φυσικής γλώσσας (Natural Language Processing) και συμμετέχουν σε έργα που αφορούν την ανάπτυξη εργαλείων λεγαλ τεξι συμμαριζατιον, αυτόματης ανάκτησης νομικής πληροφορίας, και ανάλυσης δικαστικών αποφάσεων με χρήση LLMs[93], [102]. Αυτή η πρακτική προσέγγιση τους διδάσκει να κατανοούν όλη την αλυσίδα παραγωγής και χρήσης των συστημάτων τεχνητής νοημοσύνης – από τη συλλογή δεδομένων (data collection) και ξεκαθάρισμα δεδομένων (data cleaning), μέχρι το fine-tuning μοντέλων, τη μέτρηση μετρικών απόδοσης και τη διαχείριση θεμάτων προκαταλήψεων [102].

Επιπλέον, τα πιο σύγχρονα προγράμματα σπουδών δίνουν ιδιαίτερη έμφαση στη χρήση συγκεκριμένων εργαλείων που προέρχονται από τη μηχανική μάθηση και την τεχνητή νοημοσύνη. Για παράδειγμα, στο UCL LawTech Lab και στο CodeX του Πανεπιστημίου Stanford, φοιτητές εξοικειώνονται με σύγχρονες γλώσσες προγραμματισμού, χρήση βιβλιοθηκών όπως spaCy, NLTK και HuggingFace Transformers, υλοποίηση embeddings (Word2Vec, BERT), σχεδιασμό pipelines ανάλυσης νομικών δεδομένων και ανάπτυξη demo εφαρμογών που συνδέουν user interface (π.χ. legal search platforms) με backend LLMs [100], [102], [101].

Παράλληλα, πολλά πανεπιστήμια προωθούν capstone projects στα οποία μικτές ομάδες φοιτητών πληροφορικής και νομικής αναλαμβάνουν να υλοποιήσουν end-to-end λύσεις – π.χ. αυτοματοποιημένη κατηγοριοποίηση υποθέσεων, ζηατβοτς νομικής πληροφόρησης, συστήματα recommendation για νομολογία ή frameworks για explainable legal prediction. Στα συγκεκριμένα projects, ιδιαίτερη βαρύτητα δίνεται στα ζητήματα ηθικής (π.χ. data anonymization, responsible use, fairness auditing), αλλά και στην παραγωγή πλήρους τεχνικής τεκμηρίωσης, η οποία απαιτεί εξοικείωση με version control, καταγραφή audit trails και αξιολόγηση αποτελεσμάτων με real-world legal benchmarks [7], [103].

Η συνεργασία με επαγγελματικούς φορείς και η συμμετοχή σε open source projects

(π.χ. συμμετοχή στο Harvard Caselaw Access Project, στα Legal Data Science Workshops του Cornell ή στην ανάπτυξη legal ontologies για το European Case Law Identifier) παρέχουν στους φοιτητές εμπειρία σε κανονιστικά/προτυποποιημένα σύνολα δεδομένων, API design, semantic search, και αξιολόγηση λειτουργίας εργαλείων σε διαφορετικά νομικά συστήματα [93], [101].

Αξίζει τέλος να τονιστεί ότι η διεπιστημονική εκπαίδευση περιλαμβάνει και πιο κλασικά τεχνολογικά μαθήματα, όπως βασικές αρχές πληροφορικής, data privacy, cloud computing και cyber security, συχνά υπό το πρίσμα της συμμόρφωσης με ρυθμιστικά πλαίσια (GDPR, AI Act). Το προφίλ του σύγχρονου νομικού ενισχύεται έτσι όχι μόνο ως προς την τεχνική κατανόηση, αλλά και ως προς την ικανότητα ελέγχου (auditing), διασφάλισης ποιότητας (quality assurance) και ηθικής λογοδοσίας (ethical accountability) σε ένα διασυνδεδεμένο, ψηφιακό οικοσύστημα.

Συνολικά, το μοντέλο της τεχνικής-νομικής διεπιστημονικότητας που προωθείται σε διεθνές επίπεδο αναβαθμίζει το επίπεδο της νομικής εκπαίδευσης και διαμορφώνει επαγγελματίες που μπορούν να ελέγχουν, να υλοποιούν και να ερμηνεύουν αλγοριθμικές λύσεις στην πράξη – με πλήρη επίγνωση των κινδύνων, των ορίων και των προοπτικών της τεχνητής νοημοσύνης στο δίκαιο.

6.1.2 Ενσωμάτωση τεχνικών και νομικών προγραμμάτων σπουδών

Η ενσωμάτωση τεχνολογικών και τεχνικών αντικειμένων στα παραδοσιακά προγράμματα σπουδών των νομικών σχολών συνιστά στρατηγική επιλογή για τη διαμόρφωση αποφοίτων που είναι κατάλληλα εξοπλισμένοι για να ανταποκριθούν στις απαιτήσεις της ψηφιακής εποχής. Η ανάγκη αυτή εντείνεται από την ευρεία διάδοση των μεγάλων γλωσσικών μοντέλων, της τεχνητής νοημοσύνης αλλά και των εργαλείων αυτοματοποίησης νομικών εργασιών, τα οποία διεισδύουν πλέον σε όλο το φάσμα της δικαστικής και διοικητικής πράξης [7], [93].

Σύμφωνα με τα διεθνή πρότυπα, η ενσωμάτωση τεχνικών αντικειμένων σε νομικά προγράμματα πραγματοποιείται με πολλαπλούς τρόπους:

- α) εισαγωγή υποχρεωτικών ή επιλεγόμενων μαθημάτων πληροφορικής, επιστήμης δεδομένων, βασικών αρχών μηχανικής μάθησης και τεχνητής νοημοσύνης
- β) σχεδιασμός διατμηματικών προγραμμάτων που φέρνουν σε συνεργασία σχολές νομικής με τμήματα μηχανικών, πληροφορικής και κοινωνικών επιστημών
- γ) δημιουργία πρακτικών εργαστηρίων, εκπαιδευτικών σεμιναρίων και ανάπτυξη εφαρμογών σε πραγματικά νομικά δεδομένα [100], [102], [101].

Χαρακτηριστικά παραδείγματα συνδυαστικών προγραμμάτων συναντώνται στο Stanford CodeX (Κέντρο Νομικής Πληροφορικής), το Harvard Law Lab, το LegalTech Lab του KU Leuven και το University College London (UCL), όπου οι φοιτητές έχουν τη δυνατότητα να διδαχθούν βασικές τεχνικές επεξεργασίας φυσικής γλώσσας, ανάπτυξη και αξιολόγηση αλγοριθμικών συστημάτων καθώς και πρακτική άσκηση σε εργαλεία ανάλυσης νομικών κειμένων και αυτοματοποιημένης επεξεργασίας δικαστικών εγγράφων [100], [102]. Σε αυτά τα προγράμματα, συχνά ενσωματώνονται έργα όπου μικτές ομάδες φοιτητών αναπτύσσουν από

κοινού νομικές εφαρμογές βασισμένες στη τεχνητή νοημοσύνη, εξετάζοντας τόσο την τεχνική λειτουργία όσο και τις κανονιστικές/ηθικές διαστάσεις των λύσεων τους.

Επιπρόσθετα, σε πολλές σχολές νομικής της Ευρώπης, της Βόρειας Αμερικής και της Ασίας, εισάγονται modules για τη διαχείριση της ιδιωτικότητας (privacy management), τις αρχές προστασίας δεδομένων (data protection) και τη συμμόρφωση με θεσμικά πλαίσια όπως ο Γενικός Κανονισμός για την Προστασία Δεδομένων (GDPR) και ο Ευρωπαϊκός Κανονισμός Τεχνητής Νοημοσύνης (AI Act) [93], [103]. Η πρακτική εκπαίδευση εμπλουτίζεται με σενάρια χρήσης που καλύπτουν πραγματικές υποθέσεις νομικής πληροφορικής (legal informatics), ανάλυση μεγάλων όγκων δεδομένων (big data analysis) και προσομοιώσεις αυτόματης παραγωγής νομικών γνωμοδοτήσεων μέσω αλγοριθμικών εργαλείων.

Η εμπειρική ανάλυση δείχνει ότι η επιτυχής ενσωμάτωση των τεχνικών αντικειμένων προϋποθέτει τη στενή συνεργασία μεταξύ καθηγητών νομικής και μηχανικών, την ανάπτυξη κοινών μαθημάτων, τη δημιουργία κοινών βιβλιογραφιών και τη διεπιστημονική προσέγγιση στη διδασκαλία [100], [101], [103]. Σε αρκετές περιπτώσεις, υλοποιούνται κοινές διαλέξεις (joint lectures), συνέδρια και ερευνητικά έργα, όπου οι φοιτητές καλούνται να επιλύσουν πρακτικά προβλήματα, να αναλύσουν σύνολα δεδομένων, να αναπτύξουν και να τεκμηριώσουν λύσεις τόσο σε τεχνικό όσο και σε νομικό επίπεδο.

Κρίσιμο στοιχείο της σύγχρονης προσέγγισης είναι και η ένταξη της εκπαίδευσης σε ηθικές και κανονιστικές αρχές, ώστε οι φοιτητές να αποκτούν όχι μόνο τεχνικές αλλά και δεοντολογικές δεξιότητες. Αυτό επιτρέπει στους αποφοίτους να λειτουργούν με υπευθυνότητα σε ρόλους compliance, auditing, και αξιολόγησης της τεχνητής νοημοσύνης στην πράξη [7], [93].

Τέλος, η δυναμική ενσωμάτωση τεχνικών αντικειμένων συμβάλλει αποφασιστικά στη γέφυρα του ψηφιακού χάσματος, προσφέροντας ίσες ευκαιρίες σε φοιτητές με διαφορετικά τεχνολογικά υπόβαθρα και διευκολύνοντας τη μετάβαση προς μια σχολή νομικής που ανταποκρίνεται στις απαιτήσεις της εποχής της τεχνητής νοημοσύνης. Η σταδιακή θεσμική υιοθέτηση τέτοιων πρακτικών καταγράφεται τόσο σε εθνικές στρατηγικές (π.χ. European E-Justice Strategy 2024–2028) όσο και σε διεθνή ερευνητικά προγράμματα, καθιστώντας τη διασύνδεση νομικών και τεχνικών προγραμμάτων σπουδών όχι απλώς επιλογή, αλλά αναγκαιότητα για το μέλλον του δικαίου.

6.2 Τεχνικές Εξελίξεις

Η διαρκώς επιταχυνόμενη πρόοδος της τεχνητής νοημοσύνης έχει οδηγήσει στη ριζική μεταμόρφωση του νομικού οικοσυστήματος, δημιουργώντας σύνθετες τεχνικές και θεσμικές προκλήσεις αλλά και ανεκτίμητες δυνατότητες καινοτομίας. Σήμερα, το φάσμα των συστημάτων AI περιλαμβάνει τόσο παραδοσιακές συμβολικές προσεγγίσεις, όσο και σύγχρονα αλγοριθμικά μοντέλα μηχανικής μάθησης και μεγάλα γλωσσικά μοντέλα. Παράλληλα, υβριδικές και πολυτροπικές αρχιτεκτονικές επιτρέπουν τη συνδυαστική αξιοποίηση διαφορετικών τεχνολογιών για την κάλυψη εξειδικευμένων νομικών αναγκών [92], [104].

Οι βασικές τεχνικές προκλήσεις στο νομικό οικοσύστημα της τεχνητής νοημοσύνης αφορούν τέσσερις πυλώνες:

- **Ανθεκτικότητα (robustness):** Η ανάγκη τα συστήματα να παραμένουν αξιόπιστα, ακριβή και ασφαλή, ακόμα και όταν αντιμετωπίζουν ακραία, άγνωστα ή παραμορφωμένα δεδομένα εισόδου [92], [90].
- **Εξηγησιμότητα (explainability):** Η ικανότητα των συστημάτων να παρέχουν κατανοητές, επαληθεύσιμες και διαφανείς εξηγήσεις για τα outputs τους – προϋπόθεση για τη λογοδοσία και την αποδοχή τους από τη δικαστική και ρυθμιστική κοινότητα [104].
- **Προσαρμογή σε τομείς (domain adaptation):** Η τεχνολογική δυνατότητα των μοντέλων να μεταφέρονται και να εξειδικεύονται σε επιμέρους νομικούς τομείς, καλύπτοντας θεματικές και εθνικές ιδιαιτερότητες [92], [55].
- **Πολυγλωσσία & διασυνοριακές δυνατότητες (multilinguality, cross-jurisdictional capabilities):** Η ανάγκη για αξιόπιστη υποστήριξη πολλαπλών γλωσσών και διακρατικών εφαρμογών, με στόχο την ισότιμη πρόσβαση και την αντιμετώπιση του ψηφιακού χάσματος [105].

Η διεθνής βιβλιογραφία τονίζει ότι για την επίλυση των παραπάνω προκλήσεων απαιτείται συνδυασμός προηγμένων αλγοριθμικών τεχνικών (όπως adversarial training, interpretable machine learning, transfer learning), αποτελεσματικών μηχανισμών ελέγχου και αξιολόγησης (auditing, evaluation frameworks) και ρυθμιστικής συμμόρφωσης (compliance με (AI Act, GDPR) [104], [106].

Το παρόν τμήμα αναλύει διεξοδικά τα παραπάνω ζητήματα, παρουσιάζοντας ενδεικτικές τεχνικές λύσεις, εμπειρικά παραδείγματα από διεθνή συστήματα και τα ερευνητικά πεδία που αναδεικνύονται ως κρίσιμα για το μέλλον της τεχνητής νοημοσύνης στο δίκαιο.

6.2.1 Βελτιώσεις στην ανθεκτικότητα

Η ανθεκτικότητα (robustness) αποτελεί θεμελιώδες τεχνικό ζητούμενο στην υλοποίηση συστημάτων τεχνητής νοημοσύνης στον τομέα του δικαίου, καθώς επηρεάζει άμεσα την αξιοπιστία, τη διαφάνεια και την κοινωνική αποδοχή τους. Ο όρος αναφέρεται στην ικανότητα ενός αλγοριθμικού συστήματος να διατηρεί σταθερή και ακριβή συμπεριφορά ακόμα και όταν εκτίθεται σε απρόβλεπτα, ασαφή, ή κακόβουλα παραμορφωμένα δεδομένα εισόδου (adversarial examples), καθώς και σε μεταβολές του περιβάλλοντος χρήσης (distribution shifts) [92], [90].

Από την πρώτη περίοδο της συμβολικής τεχνητής νοημοσύνης (symbolic AI & expert systems), η ανθεκτικότητα διασφαλιζόταν μέσω της προσεκτικής χειροκίνητης καταγραφής όλων των δυνατών καταστάσεων σε κανόνες και βάσεις γνώσης. Αυτή η προσέγγιση, αν και εξαιρετικά διαφανής, περιοριζόταν από την αδυναμία να αντιμετωπίσει περιπτώσεις πέρα από τις προκαθορισμένες, οδηγώντας έτσι σε ευθραυστότητα (fragility) όταν τα συστήματα εκτίθεντο σε νέα ή μη προβλέψιμα σενάρια [55].

Η μετάβαση στα σύγχρονα συστήματα μηχανικής μάθησης και ιδιαίτερα στα μεγάλα γλωσσικά μοντέλα, έχει φέρει μια νέα διάσταση στην έννοια της ανθεκτικότητας. Τα συστήματα αυτά, αν και παρέχουν σημαντικά μεγαλύτερη ευελιξία και ικανότητα γενίκευσης σε μη προγραμματισμένες περιπτώσεις, αντιμετωπίζουν σοβαρά ζητήματα ευαισθησίας σε

αλλαγές εισόδου. Χαρακτηριστικά παραδείγματα αποτελούν οι παραισθήσεις (hallucinations), δηλαδή η παραγωγή φαινομενικά ορθών αλλά στην πραγματικότητα ανυπόστατων και εσφαλμένων απαντήσεων, καθώς και η ευαισθησία σε μικρές, κακόβουλες μεταβολές στις εισόδους (adversarial examples), που μπορεί να οδηγήσουν σε εντελώς λανθασμένες αποφάσεις [92], [90], [55].

Για την αντιμετώπιση αυτών των τεχνικών προκλήσεων, η πρόσφατη βιβλιογραφία έχει προτείνει και εξετάσει πληθώρα μεθόδων και τεχνικών λύσεων:

- **Εκπαίδευση με αντίπαλα παραδείγματα (Adversarial Training):** Τα μοντέλα εκπαιδεύονται σε ένα σύνολο δεδομένων που περιλαμβάνει όχι μόνο ‘κανονικά’ δεδομένα, αλλά και τροποποιημένες εκδοχές τους που έχουν δημιουργηθεί για να προκαλέσουν εσκεμμένα αστοχίες. Αυτή η τεχνική ενισχύει την ικανότητα του συστήματος να αναγνωρίζει και να αποφεύγει παρόμοιες μελλοντικές επιθέσεις [92], [90].
- **Εμπλουτισμός δεδομένων εκπαίδευσης (Data Augmentation):** Με τη μέθοδο αυτή αυξάνεται η ποικιλομορφία των δεδομένων εκπαίδευσης, συμπεριλαμβάνοντας πολλές παραλλαγές πραγματικών νομικών κειμένων, ακραία σενάρια (εδγε-ζασες), αλλά και διαφορετικές γλωσσικές ή πολιτισμικές εκδοχές κειμένων. Αυτό βελτιώνει τη στατιστική γενίκευση και μειώνει το ρίσκο σφαλμάτων όταν το σύστημα εκτεθεί σε πρωτόγνωρα ή μη προβλέψιμα δεδομένα [90], [105].
- **Τεχνικές τακτικής κανονικοποίησης (Regularization techniques):** Μέθοδοι όπως το dropout, το weight decay και το label smoothing ενσωματώνονται κατά την εκπαίδευση των μοντέλων, προκειμένου να μειωθεί η υπερπροσαρμογή (overfitting), αυξάνοντας έτσι την ανθεκτικότητα και τη δυνατότητα γενίκευσης σε νέα, άγνωστα δεδομένα εισόδου [92], [55].
- **Αρχιτεκτονικές υβριδικών συστημάτων (Hybrid architectures):** Η συνδυασμένη χρήση συμβολικών προσεγγίσεων (symbolic AI, expert systems) με μοντέλα μηχανικής μάθησης (μασηνине λεαρνινγ) δημιουργεί συστήματα που μπορούν να εκμεταλλευτούν τα πλεονεκτήματα της διαφάνειας και της προβλεψιμότητας των κανόνων με την ευελιξία των στατιστικών μοντέλων, ενισχύοντας συνολικά την ανθεκτικότητα σε διάφορες νομικές εφαρμογές [55].

Στην πράξη, τέτοιες τεχνικές έχουν εφαρμοστεί επιτυχώς σε διάφορες περιπτώσεις, όπως:

- Η χρήση εκπαίδευσης με αντίπαλα παραδείγματα (adversarial training) σε συστήματα κατηγοριοποίησης δικαστικών αποφάσεων και αξιολόγησης κινδύνου (risk assessment systems), όπου παρατηρήθηκε αύξηση της ακρίβειας και μείωση των σφαλμάτων ακόμα και όταν αντιμετωπίζουν εσκεμμένα παραμορφωμένα δεδομένα εισόδου [92], [90].
- Η υλοποίηση υβριδικών συστημάτων σε πλατφόρμες αυτοματοποίησης νομικής έρευνας και διαχείρισης συμβολαίων (automated contract analysis), συνδυάζοντας συμβολικούς κανόνες για κρίσιμα νομικά ζητήματα με βαθιά μοντέλα μηχανικής μάθησης για ευελιξία και προσαρμοστικότητα [55].

- Η αξιοποίηση τεχνικών κανονικοποίησης (regularization) και εμπλουτισμού δεδομένων (data augmentation) σε μεγάλα γλωσσικά μοντέλα εξειδικευμένα στο δίκαιο (π.χ. LegalBERT, LawGPT), με σκοπό την αποτελεσματικότερη αντιμετώπιση των παραισθήσεων (hallucinations) και τη βελτίωση της σταθερότητας σε εξειδικευμένα νομικά σύνολα δεδομένων. [92], [90], [55].

Η συνεχιζόμενη πρόοδος στον τομέα της ανθεκτικότητας δεν περιορίζεται μόνο στην τεχνική διάσταση. Η εφαρμογή των συστημάτων αυτών σε πραγματικά δικαστικά και διοικητικά περιβάλλοντα απαιτεί επίσης αυστηρούς μηχανισμούς ελέγχου, συνεχή αξιολόγηση σε ρεαλιστικά σενάρια χρήσης (real-world evaluation) και αυστηρή συμμόρφωση με κανονιστικά πλαίσια, ώστε να διασφαλιστεί ότι τα συστήματα ΑΙ παραμένουν ασφαλή, αξιόπιστα και κοινωνικά αποδεκτά [104].

6.2.2 Βελτιώσεις στην εξηγήσιμη λειτουργία

Η εξηγήσιμη λειτουργία αποτελεί μία από τις πλέον θεμελιώδεις τεχνικές απαιτήσεις στα σύγχρονα συστήματα τεχνητής νοημοσύνης στον τομέα του δικαίου. Η δυνατότητα ενός συστήματος τεχνητής νοημοσύνης να εξηγήει τις αποφάσεις ή τις προβλέψεις που παράγει αποτελεί απαραίτητη προϋπόθεση τόσο για τη νομική συμμόρφωση όσο και για τη διατήρηση της εμπιστοσύνης των χρηστών και των θεσμικών φορέων στο σύστημα. Η εξηγησιμότητα συνδέεται άμεσα με ζητήματα διαφάνειας (transparency), λογοδοσίας (accountability) και δίκαιης αντιμετώπισης (fairness), και αποτελεί κεντρική προτεραιότητα στον Ευρωπαϊκό Κανονισμό για την Τεχνητή Νοημοσύνη [104], [107]. Τα πρώτα συστήματα τεχνητής νοημοσύνης στο δίκαιο, τα συμβολικά συστήματα (symbolic/expert systems), παρείχαν εκ φύσεως υψηλή εξηγησιμότητα, καθώς οι αποφάσεις τους βασίζονταν σε σαφώς διατυπωμένους και ερμηνεύσιμους κανόνες. Αντίθετα, τα σύγχρονα συστήματα μηχανικής μάθησης, ιδιαίτερα τα βαθιά νευρωνικά δίκτυα και τα μεγάλα γλωσσικά μοντέλα, αν και παρουσιάζουν εντυπωσιακή ευελιξία και αποτελεσματικότητα, χαρακτηρίζονται συχνά από τη λεγόμενη συμπεριφορά 'μαύρου κουτιού' (black box), όπου οι διαδικασίες λήψης αποφάσεων δεν είναι άμεσα κατανοητές ή ερμηνεύσιμες από τον άνθρωπο [107], [89]. Για την αντιμετώπιση αυτής της κρίσιμης πρόκλησης, η επιστημονική βιβλιογραφία προτείνει διάφορες τεχνικές και προσεγγίσεις:

- **Μοντέλα με εγγενή εξηγησιμότητα (intrinsically explainable models):** Μοντέλα όπως τα decision trees, τα συμβολικά μοντέλα και τα Bayesian networks παρέχουν διαφανή και κατανοητή δομή αποφάσεων, κάτι ιδιαίτερα χρήσιμο για νομικές εφαρμογές που απαιτούν απόλυτη σαφήνεια [107], [89].
- **Εκ των υστέρων τεχνικές ερμηνείας (post-hoc explainability):** Τεχνικές όπως το LIME (Local Interpretable Model-Agnostic Explanations), το SHAP (Shapley Additive Explanations), και τα attention-based visualizations επιτρέπουν την εκ των υστέρων ερμηνεία των αποφάσεων πολύπλοκων μοντέλων. Αυτές οι τεχνικές παρέχουν τη δυνατότητα να κατανοηθεί ποιες παράμετροι και δεδομένα εισόδου είχαν το μεγαλύτερο βάρος στις αποφάσεις [106], [107], [89].

- **Εφαρμογή υβριδικών συστημάτων (hybrid approaches):** Συνδυάζουν τη συμβολική προσέγγιση και τους κανόνες με μοντέλα βαθιάς μάθησης, ώστε η διαδικασία λήψης αποφάσεων να παραμένει κατανοητή χωρίς να θυσιάζεται η ευελιξία και η δυνατότητα γενίκευσης [55].

Η βιβλιογραφία επισημαίνει ότι η ενίσχυση της εξηγησιμότητας δεν αποτελεί απλώς τεχνική απαίτηση, αλλά προϋπόθεση θεσμικής και κοινωνικής αποδοχής των συστημάτων τεχνητής νοημοσύνης στο δίκαιο. Σε πρόσφατες μελέτες υπογραμμίζεται ότι οι χρήστες (δικαστές, δικηγόροι, πολίτες) εμπιστεύονται περισσότερο συστήματα που παρέχουν σαφείς εξηγήσεις και μπορούν να αιτιολογήσουν με διαφανή τρόπο τις αποφάσεις τους [89].

Επιπλέον, η Ευρωπαϊκή στρατηγική για την Ηλεκτρονική Δικαιοσύνη και ο Ευρωπαϊκός Κανονισμός για την Τεχνητή Νοημοσύνη απαιτούν από τις νομικές εφαρμογές τεχνητής νοημοσύνης αυστηρούς μηχανισμούς διαφάνειας και εξηγησιμότητας, με υποχρεωτική καταγραφή και παρουσίαση των κριτηρίων αποφάσεων σε κάθε στάδιο της διαδικασίας, ώστε να διασφαλίζεται η διαφάνεια, η λογοδοσία και η προστασία των δικαιωμάτων των πολιτών [104]. Καθώς οι απαιτήσεις για εξηγησιμότητα αυξάνονται, η έρευνα συνεχίζει να επικεντρώνεται σε καινοτόμες τεχνικές όπως η ενσωμάτωση causal inference σε μοντέλα ΑΙ, η ανάπτυξη semantic explainability frameworks και η ενίσχυση της διαφάνειας μέσω της τυποποίησης των audit trails και της συστηματικής αξιολόγησης της ερμηνευσιμότητας μέσω προτύπων όπως το ISO/IEC και το IEEE[89].

6.2.3 Τεχνικές προσαρμογής στον τομέα

Η προσαρμογή στον τομέα (domain adaptation) αποτελεί ένα κρίσιμο τεχνικό ζήτημα κατά την ανάπτυξη συστημάτων τεχνητής νοημοσύνης στο δίκαιο, δεδομένου ότι η επιτυχία αυτών των συστημάτων εξαρτάται άμεσα από την ικανότητά τους να προσαρμόζονται αποτελεσματικά στα ιδιαίτερα χαρακτηριστικά διαφορετικών νομικών κλάδων και περιβαλλόντων. Σε αντίθεση με γενικά μοντέλα τεχνητής νοημοσύνης, τα οποία εκπαιδεύονται σε μεγάλα, γενικού τύπου σύνολα δεδομένων, οι νομικές εφαρμογές τεχνητής νοημοσύνης απαιτούν υψηλό βαθμό εξειδίκευσης για την επίτευξη αποδεκτών επιδόσεων [92], [55]. Τα συστήματα τεχνητής νοημοσύνης στο δίκαιο συχνά καλούνται να λειτουργήσουν σε ποικίλα και εξειδικευμένα νομικά πλαίσια, όπως το ποινικό, το αστικό, το διοικητικό ή το εμπορικό δίκαιο. Κάθε ένας από αυτούς τους τομείς χαρακτηρίζεται από ιδιαίτερη ορολογία, διαδικασίες και πρότυπα που απαιτούν την προσαρμογή των μοντέλων, έτσι ώστε να παρέχουν αξιόπιστες και ακριβείς αποφάσεις ή προτάσεις [55], [107].

Η βιβλιογραφία έχει αναγνωρίσει μια σειρά από αποτελεσματικές τεχνικές για την προσαρμογή των μοντέλων ΑΙ σε νομικούς τομείς, οι οποίες περιλαμβάνουν:

- **Εξειδικευμένη εκπαίδευση (fine-tuning):** Τα μοντέλα αρχικά εκπαιδεύονται σε μεγάλα γενικά σύνολα δεδομένων (π.χ. Wikipedia, γενικές νομικές βάσεις) και στη συνέχεια υποβάλλονται σε δευτερογενή εκπαίδευση (fine-tuning) σε εξειδικευμένα νομικά σύνολα δεδομένων, ώστε να ενσωματώνουν τη συγκεκριμένη νομική ορολογία και τις ιδιαιτερότητες του αντίστοιχου νομικού κλάδου [92], [107].

- **Μεταφορά μάθησης (transfer learning):** Χρησιμοποιώντας προϋπάρχουσα γνώση από εκπαιδευμένα μοντέλα (π.χ. GPT, BERT), η μεταφορά μάθησης επιτρέπει την αποτελεσματική προσαρμογή του συστήματος σε νέες νομικές υποπεριοχές με περιορισμένα σύνολα δεδομένων, μειώνοντας τον απαιτούμενο χρόνο και κόστος προσαρμογής [92], [55], [107].
- **Μάθηση με λίγα παραδείγματα (few-shot learning):** Ιδιαίτερα χρήσιμη στις νομικές υποπεριοχές όπου δεν υπάρχουν εκτενή σύνολα δεδομένων. Με την τεχνική αυτή, τα μοντέλα μπορούν να προσαρμόζονται με ελάχιστα εξειδικευμένα παραδείγματα, αξιοποιώντας την υπάρχουσα γνώση από γενικότερη εκπαίδευση [55], [107].
- **Μηχανισμοί συνεχούς μάθησης (continuous learning):** Με τη συνεχή ενσωμάτωση νέων δεδομένων και τη δυναμική επικαιροποίηση των μοντέλων, επιτυγχάνεται αποτελεσματική και διαρκής προσαρμογή στα μεταβαλλόμενα νομικά περιβάλλοντα, π.χ. μεταβολές νομοθεσίας, νέες δικαστικές αποφάσεις ή αλλαγές στον τρόπο ερμηνείας των νομικών εννοιών [92], [55].

Τέλος, η εφαρμογή συστημάτων προσαρμογής στον τομέα απαιτεί αυστηρή συμμόρφωση με θεσμικά και κανονιστικά πρότυπα, όπως ο Ευρωπαϊκός Κανονισμός για την Τεχνητή Νοημοσύνη και ο Γενικός Κανονισμός για την Προστασία Δεδομένων (GDPR), ώστε να διασφαλίζεται ότι τα προσαρμοσμένα μοντέλα λειτουργούν εντός θεσμικά αποδεκτών πλαισίων, διατηρώντας υψηλά επίπεδα διαφάνειας, λογοδοσίας και προστασίας των δικαιωμάτων των πολιτών [104].

6.2.4 Πολυγλωσσία και Διασυννοριακές Δυνατότητες

Η πολυγλωσσία και οι διασυννοριακές δυνατότητες αποτελούν κρίσιμα χαρακτηριστικά των σύγχρονων συστημάτων τεχνητής νοημοσύνης που εφαρμόζονται στον νομικό τομέα. Καθώς η παγκοσμιοποίηση ενισχύει τις διεθνείς νομικές συναλλαγές, η ανάγκη για συστήματα που μπορούν να λειτουργήσουν αποτελεσματικά σε πολλαπλά γλωσσικά και πολιτισμικά περιβάλλοντα καθίσταται όλο και πιο έντονη. Στο πλαίσιο αυτό, η ανάπτυξη μοντέλων και τεχνικών που επιτρέπουν αξιόπιστες και υψηλής ποιότητας πολυγλωσσικές εφαρμογές στο δίκαιο, καθίσταται πρωταρχικής σημασίας [92], [55], [107], [74].

Η ανάγκη για εξειδικευμένα δεδομένα εκπαίδευσης γίνεται ακόμα πιο έντονη σε πολύγλωσσα κράτη ή διεθνή οργανισμούς, όπου η νομική ακρίβεια μεταξύ των γλωσσών πρέπει να διατηρείται με απόλυτη αυστηρότητα [107], [74]. Ένα πολύ χαρακτηριστικό παράδειγμα είναι η Ελβετία, μια χώρα με τέσσερις επίσημες γλώσσες (Γερμανικά, Γαλλικά, Ιταλικά και Ρετορομανικά), στην οποία η ανάγκη για πολυγλωσσικές νομικές μεταφράσεις είναι ιδιαίτερα αυξημένη.

Η ανάλυση και αξιολόγηση των συστημάτων που διεξήχθη στο πλαίσιο του SwiLTraBench ανέδειξε αρκετά σημαντικά ευρήματα[74]:

- Τα frontier μοντέλα, όπως το Claude-3.5-Sonnet[108], επιτυγχάνουν πολύ υψηλές επιδόσεις σε μεταφράσεις νομικών κειμένων σε πολλαπλές γλώσσες, ενώ ειδικευμένα συστήματα μετάφρασης, όπως το MADLAD-400, παρουσιάζουν εξαιρετικές επιδόσεις

σε συγκεκριμένους τύπους νομικών κειμένων, αλλά δυσκολεύονται περισσότερο σε άλλους τύπους

- Η τεχνική της λεπτομερούς προσαρμογής σε εξειδικευμένα νομικά δεδομένα βελτιώνει σημαντικά την ποιότητα των μεταφράσεων, ωστόσο, ακόμα και αυτά τα μοντέλα υπολείπονται των κορυφαίων «frontier» μοντέλων που αξιολογούνται σε ζερο-σηοι σενάρια .
- Τα συστήματα παρουσιάζουν ομοιόμορφη απόδοση στις περισσότερες από τις γλώσσες, αλλά αντιμετωπίζουν μεγαλύτερες προκλήσεις στη μετάφραση από ή προς γλώσσες με λιγότερους διαθέσιμους πόρους (όπως η Ρετορομανική), γεγονός που τονίζει την ανάγκη για περαιτέρω ανάπτυξη και βελτίωση των συστημάτων στις λιγότερο ομιλούμενες γλώσσες.

Τέλος, οι τεχνικές πολυγλωσσικότητας και διασυνοριακής προσαρμογής δεν περιορίζονται στην απλή μετάφραση κειμένων, αλλά μπορούν να επεκταθούν σε πολύ πιο σύνθετες διαδικασίες όπως η νομική ανάλυση και η αυτοματοποίηση διασυνοριακών νομικών διαδικασιών (cross-jurisdictional automation). Στο πλαίσιο αυτό, η συνεχής συνεργασία μεταξύ νομικών και μηχανικών, καθώς και η συμμόρφωση με διεθνείς κανονισμούς, αποτελούν βασικές προϋποθέσεις για την επιτυχημένη ανάπτυξη συστημάτων που θα είναι πραγματικά αποτελεσματικά και αποδεκτά στο διεθνές νομικό οικοσύστημα [104].

6.3 Ανάπτυξη υβριδικών συστημάτων

Η ανάπτυξη υβριδικών συστημάτων τεχνητής νοημοσύνης αποτελεί μια στρατηγική προσέγγιση που συνδυάζει διαφορετικά μοντέλα, τεχνολογίες και τεχνικές, με σκοπό τη βέλτιστη αξιοποίηση των πλεονεκτημάτων τους και την αντιμετώπιση των περιορισμών που παρουσιάζει κάθε επιμέρους τεχνολογία ξεχωριστά. Στο πεδίο του δικαίου, τα υβριδικά συστήματα αποκτούν ιδιαίτερη σημασία, καθώς η πολυπλοκότητα των νομικών ζητημάτων απαιτεί τον συνδυασμό πολλαπλών και συμπληρωματικών τεχνολογικών μεθόδων για τη διασφάλιση της ακρίβειας, της αξιοπιστίας και της διαφάνειας των αποτελεσμάτων [104], [55], [107].

Ο όρος υβριδικά συστήματα αναφέρεται κυρίως στη συνδυαστική αξιοποίηση των συμβολικών προσεγγίσεων (symbolic AI), όπως τα συστήματα βασισμένα σε κανόνες (rule-based systems), μαζί με σύγχρονα μοντέλα μηχανικής μάθησης (machine learning), βαθιά νευρωνικά δίκτυα (deep neural networks), μεγάλα γλωσσικά μοντέλα (LLMs), και τεχνικές επεξεργασίας φυσικής γλώσσας. Ειδικότερα στον τομέα του δικαίου, τα υβριδικά συστήματα στοχεύουν στην επίτευξη μιας ισορροπίας μεταξύ της διαφάνειας και της εξηγησιμότητας των συμβολικών συστημάτων από τη μία πλευρά, και της υψηλής ευελιξίας και αποτελεσματικότητας που προσφέρουν τα μοντέλα μηχανικής μάθησης από την άλλη [92], [107].

Η ανάπτυξη υβριδικών συστημάτων προσφέρει μια σειρά από σημαντικά πλεονεκτήματα:

- **Βελτίωση της εξηγησιμότητας (explainability):** Τα υβριδικά συστήματα αξιοποιούν συμβολικούς κανόνες και διαφανείς δομές σε κρίσιμα σημεία της διαδικασίας λήψης αποφάσεων, καθιστώντας ευκολότερη την κατανόηση και την αιτιολόγηση των τελικών

αποτελεσμάτων τους. Ταυτόχρονα, αξιοποιούν τις δυνατότητες γενίκευσης και ευελιξίας των μοντέλων μηχανικής μάθησης όπου η ερμηνεία είναι λιγότερο κρίσιμη [107], [89].

- **Αύξηση της ανθεκτικότητας (robustness):** Η συνδυαστική χρήση πολλαπλών μοντέλων μειώνει τον κίνδυνο των σφαλμάτων, των «παραισθήσεων» (ηαλλυσινατιονς) και των προβλημάτων που σχετίζονται με adversarial examples, καθώς τα συμβολικά και στατιστικά μοντέλα λειτουργούν συμπληρωματικά ως προς την αντιμετώπιση αδυναμιών ή λαθών [92], [90].
- **Εξειδικευμένη προσαρμογή (domain adaptation):** Τα υβριδικά συστήματα επιτρέπουν πιο εύκολη και αποτελεσματική προσαρμογή σε συγκεκριμένους νομικούς τομείς, αφού μπορούν να συνδυάζουν γενικά στατιστικά μοντέλα με εξειδικευμένους κανόνες ή τοπική γνώση, παρέχοντας έτσι μεγαλύτερη ακρίβεια και ευελιξία σε διαφορετικά νομικά πλαίσια [55], [107].
- **Βελτίωση της διασυννοριακής λειτουργίας και πολυγλωσσίας (multilinguality):** Με τη συνδυαστική χρήση πολυγλωσσικών νευρωνικών μοντέλων και συμβολικών λεξικογραφικών κανόνων, τα υβριδικά συστήματα μπορούν να διατηρήσουν υψηλή ποιότητα και νομική ακρίβεια σε πολύπλοκες μεταφραστικές εφαρμογές, εξυπηρετώντας έτσι καλύτερα τις διεθνείς και πολυγλωσσικές ανάγκες του νομικού τομέα [55], [74].

6.3.1 Πλαίσια συνεργασίας μεταξύ ανθρώπων και τεχνητής νοημοσύνης

Η αποτελεσματική συνεργασία μεταξύ ανθρώπων και συστημάτων τεχνητής νοημοσύνης (Human-AI Collaboration) αποτελεί μια στρατηγική προσέγγιση που στοχεύει στην αξιοποίηση των καλύτερων χαρακτηριστικών τόσο του ανθρώπου όσο και της τεχνολογίας, για τη βελτιστοποίηση των νομικών διαδικασιών και τη διασφάλιση της υψηλότερης δυνατής ποιότητας στις νομικές υπηρεσίες. Στο πλαίσιο αυτό, η τεχνητή νοημοσύνη δεν αντιμετωπίζεται ως πλήρης αντικαταστάτης των ανθρώπινων δικηγόρων ή δικαστών, αλλά ως ένα ισχυρό εργαλείο υποστήριξης της ανθρώπινης κρίσης, ικανό να ενισχύσει την ταχύτητα, την ακρίβεια και την αξιοπιστία των αποφάσεων [104], [107], [89]. Η βιβλιογραφία υπογραμμίζει την ανάγκη ανάπτυξης σαφών πλαισίων και αρχιτεκτονικών που να επιτρέπουν την αποτελεσματική ενσωμάτωση της τεχνητής νοημοσύνης σε διαδικασίες που καθοδηγούνται από ανθρώπους. Τα πλαίσια αυτά ορίζουν με σαφήνεια τους ρόλους των ανθρώπων και της τεχνητής νοημοσύνης, τα σημεία παρέμβασης και αλληλεπίδρασης, καθώς και τους μηχανισμούς για τον έλεγχο και τη διασφάλιση της αξιοπιστίας των τεχνητών αποφάσεων [92], [104]. Μεταξύ των κύριων αρχών που καθορίζουν τα πλαίσια συνεργασίας ανθρώπου-AI στο νομικό πεδίο, περιλαμβάνονται:

- **Αρχή της ανθρώπινης εποπτείας (human oversight):** Ο άνθρωπος διατηρεί τον τελικό έλεγχο των αποφάσεων που παράγονται από τα συστήματα AI, έχοντας τη δυνατότητα να εγκρίνει, να απορρίψει ή να τροποποιήσει τις προτάσεις της τεχνητής νοημοσύνης. Αυτό διασφαλίζει τη διατήρηση της ανθρώπινης ευθύνης, η οποία είναι απαραίτητη για τη συμμόρφωση με κανονιστικά πλαίσια, όπως ο Ευρωπαϊκός Κανονισμός Τεχνητής Νοημοσύνης [104].

- **Αρχή της διαφανούς κατανομής ρόλων (transparent role allocation):** Τα υβριδικά συστήματα απαιτούν σαφή και διαφανή διαχωρισμό αρμοδιοτήτων ανάμεσα στους ανθρώπους και τα μοντέλα AI, με συγκεκριμένα όρια, ώστε να αποφεύγεται η σύγχυση και να διασφαλίζεται η λογοδοσία. Η ανθρώπινη κρίση αξιοποιείται εκεί όπου η ερμηνεία και η αξιολόγηση είναι κρίσιμες, ενώ η τεχνητή νοημοσύνη χρησιμοποιείται για αυτοματοποιημένες διαδικασίες μεγάλης κλίμακας και αναλύσεις υψηλής ακρίβειας [107], [89].
- **Αρχή της συνεργατικής ανατροφοδότησης (collaborative feedback):** Οι άνθρωποι δεν είναι μόνο παθητικοί καταναλωτές των αποτελεσμάτων των συστημάτων AI, αλλά συμμετέχουν ενεργά στην ανατροφοδότηση του συστήματος, διορθώνοντας και βελτιώνοντας τα μοντέλα, μέσω της αξιολόγησης των προτάσεων, της επισήμανσης λαθών και της παροχής νέων δεδομένων εκπαίδευσης [107], [89].

Η μελλοντική έρευνα στον τομέα της συνεργασίας ανθρώπου-AI συνεχίζει να εστιάζει στην ανάπτυξη ακόμη πιο εξελιγμένων αρχιτεκτονικών και πλαισίων, που να επιτρέπουν μεγαλύτερη ευελιξία, διαφάνεια και συνεργασία. Στο πλαίσιο αυτό, είναι απαραίτητη η συνεχής διεπιστημονική συνεργασία μεταξύ μηχανικών, νομικών και επιστημόνων κοινωνικών επιστημών, ώστε τα συστήματα AI να σχεδιάζονται και να υλοποιούνται με τρόπο που να είναι όχι μόνο τεχνικά αποτελεσματικά, αλλά και κοινωνικά, ηθικά και θεσμικά αποδεκτά [104], [107], [89].

6.3.2 Εξέλιξη των έξυπνων δικαστηρίων

Η τεχνολογική ωρίμανση και η θεσμική αποδοχή των έξυπνων δικαστηρίων (smart courts) τα τελευταία χρόνια σηματοδοτούν μια νέα φάση δικαστικής καινοτομίας που υπερβαίνει την απλή αυτοματοποίηση διαδικασιών. Η σύγχρονη εξέλιξη των smart courts εστιάζει πλέον σε διαλειτουργικά, πολυγλωσσικά και διαφανή συστήματα που λειτουργούν εντός ενός ευρύτερου διεθνούς και τεχνολογικά δυναμικού οικοσυστήματος [74].

Η αυξανόμενη πολυπλοκότητα των διασυννοριακών υποθέσεων και η ευρωπαϊκή στρατηγική για την ηλεκτρονική δικαιοσύνη European E-Justice Strategy 2024–2028) προωθούν την υιοθέτηση διαλειτουργικών πλατφορμών (interoperability platforms) που διασυνδέουν τα εθνικά δικαστικά συστήματα, επιτρέποντας την ανταλλαγή πληροφοριών, τη διαχείριση διεθνών διαφορών και την ενιαία πρόσβαση σε νομολογία και διαδικασίες. Οι πλατφόρμες αυτές υιοθετούν διεθνή πρότυπα (π.χ. e-CODEX) και APIs που επιτρέπουν την αυτόματη αλληλεπίδραση μεταξύ συστημάτων με διαφορετικά νομικά και τεχνικά χαρακτηριστικά, μειώνοντας δραστικά τον χρόνο και το κόστος επίλυσης διεθνών υποθέσεων.

Τα τελευταία χρόνια, η χρήση πολυγλωσσικών συστημάτων AI και η πρόοδος στη νευρωνική μηχανική μετάφραση (Neural Machine Translation, NMT) αλλάζουν ριζικά τη δυνατότητα αυτόματης επεξεργασίας νομικών κειμένων σε πολλές γλώσσες. Η ανάπτυξη βενζηνμαρκε[74], που καλύπτει πολυγλωσσικά νομικά κείμενα για τη σύγκριση της ποιότητας διαφόρων NMT και LLM μοντέλων, αποδεικνύει ότι τα νέα frontier LLMs μπορούν πλέον να παρέχουν αξιόπιστες μεταφράσεις ακόμα και σε γλώσσες με περιορισμένους πόρους.

Μια ακόμη σημαντική πρόκληση είναι η βαθμιαία θεσμική αποδοχή των νέων τεχνολογιών από το σύνολο των εμπλεκόμενων φορέων της δικαιοσύνης. Παρόλο που η αυτοματοπο-

ίηση υπόσχεται σημαντικά οφέλη ως προς την ταχύτητα και τη διαφάνεια των διαδικασιών, εξακολουθούν να καταγράφονται επιφυλάξεις από τους δικαστικούς λειτουργούς και τους νομικούς επαγγελματίες ως προς τον βαθμό αξιοπιστίας, τη δυνατότητα ελέγχου και την πραγματική ερμηνευσιμότητα των αλγοριθμικών αποφάσεων. Η αποτελεσματική υιοθέτηση των smart courts προϋποθέτει όχι μόνο την τεχνική ωριμότητα των συστημάτων, αλλά και τη διαρκή κατάρτιση και ενεργό εμπλοκή των ανθρώπινων φορέων, ώστε να διασφαλίζεται η ισορροπία μεταξύ ανθρώπινης κρίσης και αλγοριθμικής υποστήριξης [104] [89].

Τέλος, το μέλλον των smart courts διαγράφεται ως πολυτροπικό, εξωστρεφές και διασυνδεδεμένο, με διαλειτουργικά οικοσυστήματα που υποστηρίζουν real-time διαχείριση υποθέσεων, αυτοματοποιημένη ανάλυση ποικίλων αποδεικτικών μέσων (κειμένων, εικόνων, βίντεο), ενσωματωμένους ελέγχους ποιότητας και ρισκ ασσεσμεντ μοδουλές για κάθε δικαστική διαδικασία. Οι προκλήσεις αυτές απαιτούν συνεχή τεχνική και θεσμική καινοτομία, καθώς και διεπιστημονική συνεργασία, ώστε να διαμορφωθούν συστήματα που είναι όχι μόνο τεχνολογικά άρτια, αλλά και κοινωνικά αποδεκτά, θεσμικά συμβατά και ικανά να ενισχύσουν την εμπιστοσύνη στο δικαστικό σύστημα της ψηφιακής εποχής [104], [107], [89], [74].

6.3.3 Αρχιτεκτονικές ενσωμάτωσης

Η θεμελιώδης αρχή μιας σύγχρονης αρχιτεκτονικής ενσωμάτωσης είναι η πολυστρωματική προσέγγιση (multi-layered architecture), η οποία οργανώνει το σύστημα σε διακριτά επίπεδα με σαφείς αρμοδιότητες. Πρώτα, το επίπεδο παρουσίασης (presentation/UI layer) αποτελεί το σημείο αλληλεπίδρασης των χρηστών — δικαστών, δικηγόρων και διαδίκων — με το σύστημα, μέσω ασφαλών web-βασεδ ή natie εφαρμογών. Ακολουθεί το επίπεδο λογικής/επεξεργασίας (application/service layer), όπου υλοποιούνται οι ροές εργασίας (workflows), η επιχειρησιακή λογική και οι ενσωματωμένοι αλγόριθμοι τεχνητής νοημοσύνης, συμπεριλαμβανομένων των LLMs, συστημάτων NLP και συμβολικών μοντέλων. Το τρίτο επίπεδο, αυτό των δεδομένων (data layer), είναι υπεύθυνο για τη διαχείριση πολυτροπικών δεδομένων (structured/unstructured), βάσεων δεδομένων, legal document stores και knowledge graphs [92].

Σε επίπεδο τεχνικής υλοποίησης, κυρίαρχο ρόλο διαδραματίζει η διαμοιρασμένη αρχιτεκτονική υπηρεσιών (microservices), όπου κάθε λειτουργία υλοποιείται ως ανεξάρτητη υπηρεσία (API-first approach), διευκολύνοντας την επεκτασιμότητα, τη συντήρηση και την ασφάλεια κατά τη διασύνδεση με τρίτα συστήματα [74]. Η αποτελεσματική διασύνδεση legacy και νέων συστημάτων επιτυγχάνεται μέσω integration layers και adapters — όπως message queues και ETL pipelines — τα οποία επιτρέπουν στα υβριδικά συστήματα να γεφυρώνουν παλαιά δικαστικά πληροφοριακά συστήματα με σύγχρονα εργαλεία NLP/LLMs και πλατφόρμες cloud. Επιπλέον, οι υποδομές cloud/hybrid cloud καθίστανται απαραίτητες για την εκτέλεση μοντέλων με υψηλές απαιτήσεις, όπως η διεξαγωγή LLM inference σε cloud environments, αξιοποιώντας μηχανισμούς load balancing και fault tolerance για τη διασφάλιση υψηλής διαθεσιμότητας και ανθεκτικότητας [89]. Για την προστασία ευαίσθητων δεδομένων — όπως δικαστικές αποφάσεις και προσωπικά δεδομένα — εφαρμόζονται υβριδικά σενάρια με on-premises encryption και storage [89], [74]. Ιδιαίτερη βαρύτητα έχει και ο σχεδιασμός ροών δεδομένων (data pipelines) που υποστηρίζουν streaming, batch process-

ing και asynchronous integration. Με αυτόν τον τρόπο, τα δικαστικά συστήματα δύνανται να επεξεργάζονται δεδομένα σε πραγματικό χρόνο από πολλαπλές πηγές (e-filing, evidence upload, real-time translation), εφαρμόζοντας διαδοχικά φίλτρα όπως data cleansing, entity recognition, topic extraction [92], [55].

Η διατήρηση της συνοχής και της διαφάνειας στη ροή δεδομένων, η διαχείριση της ταυτότητας και των δικαιωμάτων πρόσβασης (identity & access management), καθώς και η υλοποίηση end-to-end monitoring και logging, αποτελούν κρίσιμους τομείς τεχνικής έρευνας. Παράλληλα, η συμμόρφωση με το GDPR/AI Act επιβάλλει την ενσωμάτωση συστημάτων audit trail, αυτόματης δημιουργίας αναφορών συμμόρφωσης (automatic compliance reporting) και τεχνικών επεξηγησιμότητας (explainability) σε κάθε στάδιο της ροής [107], [89].

Τέλος, η τεχνολογική εξέλιξη εστιάζει στην υιοθέτηση federated/hybrid AI infrastructures, που περιλαμβάνουν κατανεμημένη εκπαίδευση μοντέλων (distributed model training), υπολογισμούς με διαφύλαξη της ιδιωτικότητας (privacy-preserving computation) και ασφαλή διαμοιρασμό δεδομένων (data-sharing) μεταξύ οργανισμών, προκειμένου να διασφαλιστεί η αποτελεσματικότητα, η ασφάλεια και η κυριαρχία των δεδομένων τόσο σε εθνικό όσο και σε διακρατικό επίπεδο [89], [74].

6.4 Προσβασιμότητα και δικαιοσύνη

Η είσοδος της τεχνητής νοημοσύνης και των τεχνολογιών legal tech στον χώρο της δικαιοσύνης συνιστά μια ριζική τομή με διπλό πρόσημο: από τη μια, ανοίγει νέους ορίζοντες για την αυτοματοποίηση, την ταχύτητα και την αποδοτικότητα, ενώ από την άλλη, αναδεικνύει εντονότερα από ποτέ τα ζητήματα της προσβασιμότητας, της ισότητας και της κοινωνικής δικαιοσύνης. Το ερώτημα που τίθεται πλέον με ιδιαίτερη οξύτητα είναι κατά πόσο οι τεχνολογίες αυτές λειτουργούν ως εργαλεία απελευθέρωσης και ενίσχυσης της δημοκρατικής πρόσβασης στη δικαιοσύνη ή, αντίθετα, ως παράγοντες αναπαραγωγής υφιστάμενων και δημιουργίας νέων ανισοτήτων [89], [74], [10].

Σε τεχνικό επίπεδο, η προσβασιμότητα αφορά τόσο την υποδομή και την τεχνολογική δυνατότητα πρόσβασης (π.χ. γρήγορο διαδίκτυο, λειτουργικοί υπολογιστές/κινητά, προσβάσιμα user interfaces) όσο και την ικανότητα χρήσης των συστημάτων αυτών από πολίτες με διαφορετικά επίπεδα ψηφιακού αλφαριθμητισμού. Ιδιαίτερη πρόκληση αποτελούν οι πληθυσμιακές ομάδες που ήδη βιώνουν κοινωνικό ή οικονομικό αποκλεισμό (π.χ. ηλικιωμένοι, άτομα με αναπηρία, φτωχά νοικοκυριά, πληθυσμοί σε αγροτικές ή απομακρυσμένες περιοχές), οι οποίες κινδυνεύουν να μείνουν ακόμη πιο πίσω στην ψηφιακή εποχή της δικαιοσύνης. Το φαινόμενο αυτό περιγράφεται στη βιβλιογραφία ως ψηφιακό χάσμα δεύτερης γενιάς (second digital divide), που αφορά όχι απλώς την ύπαρξη τεχνολογίας, αλλά και τη δυνατότητα αποτελεσματικής, συνεχούς και αξιοπρεπούς χρήσης της [10].

Η τεχνολογική καινοτομία δύναται να προσφέρει σημαντικά οφέλη στη δικαιοσύνη. Πλατφόρμες online dispute resolution (ODR), virtual courtrooms, αυτοματοποιημένα νομικά chatbots και έξυπνα συστήματα ενημέρωσης επιτρέπουν την επίλυση διαφορών και την παροχή νομικής πληροφόρησης χωρίς τα παραδοσιακά εμπόδια της γεωγραφικής απόστασης ή του υψηλού κόστους πρόσβασης σε δικηγόρους και δικαστήρια. Η χρήση AI μπορεί να επιταχύνει δραστικά διαδικασίες όπως η κατάθεση εγγράφων, η λήψη δικαστικών

αποφάσεων για απλές υποθέσεις, και η μετάφραση νομικών κειμένων σε πολλαπλές γλώσσες [74], [10].

Ωστόσο, η εφαρμογή αυτών των τεχνολογιών αναδεικνύει και σημαντικούς κινδύνους. Μελέτες δείχνουν ότι τα πιο ευάλωτα κοινωνικά στρώματα όχι μόνο έχουν λιγότερη πρόσβαση σε σύγχρονες συσκευές ή σταθερό ίντερνετ, αλλά συχνά στερούνται και της τεχνολογικής αυτοπεποίθησης ή των απαραίτητων γνώσεων για να επωφεληθούν ουσιαστικά από τις υπηρεσίες legal tech. Σε δικαστικές διαδικασίες μέσω Zoom ή άλλων virtual platforms, πολίτες με φτωχότερες συσκευές, ασταθή σύνδεση ή ελλιπή εξοικείωση εμφανίζονται ως μαύρα κουτιά, δεν έχουν ισοτίμη παρουσία και συχνά υφίστανται έμμεσες διακρίσεις ακόμη και από τον ίδιο τον τρόπο που διεξάγεται η διαδικασία [10]. Επιπλέον, η αυτοματοποίηση πολλών σταδίων της δικαστικής διαδικασίας μέσω AI και ODR μπορεί να οδηγήσει σε τυποποίηση της κρίσης, απώλεια της δυνατότητας εξατομικευμένης δικαστικής προστασίας, ή ακόμα και ενίσχυση των bias, αν τα συστήματα δεν έχουν εκπαιδευτεί κατάλληλα για να αναγνωρίζουν τις ανάγκες των διαφορετικών κοινωνικών ομάδων [107], [10].

Η ανάλυση κόστους-οφέλους αναδεικνύει ότι, ενώ η μαζική υιοθέτηση legal tech μειώνει το λειτουργικό κόστος για το κράτος και τους χρήστες και επιταχύνει σημαντικά τις διαδικασίες, ενδέχεται να επιφέρει αόρατα κόστη για όσους αδυνατούν να προσαρμοστούν ή να συμμετάσχουν αποτελεσματικά. Αυτά τα κόστη περιλαμβάνουν τον κίνδυνο αδυναμίας πλήρους εκπροσώπησης, την ελλιπή πληροφόρηση, και τη σταδιακή αποξένωση από το δικαστικό σύστημα. Η βιβλιογραφία προτείνει ως αναγκαίες τις επενδύσεις σε inclusive design, στην τεχνική εκπαίδευση πολιτών, στη διαρκή αξιολόγηση των κοινωνικών συνεπειών, αλλά και στη δημιουργία off-line ή υποστηρικτικών λύσεων για όσους δεν μπορούν να καλύψουν το ψηφιακό χάσμα.

Τέλος, το μέλλον της προσβασιμότητας και της δικαιοσύνης στην ψηφιακή εποχή θα κριθεί από το αν τα τεχνικά και θεσμικά συστήματα θα μπορέσουν να συνδυάσουν την καινοτομία με την καθολική συμπερίληψη, εξασφαλίζοντας ότι κανείς δεν αποκλείεται από τα οφέλη της τεχνητής νοημοσύνης στη δικαιοσύνη εξαιτίας κοινωνικών, τεχνικών ή οικονομικών εμποδίων. Το στοιχείο της legal tech δεν είναι απλώς η αποδοτικότητα, αλλά η ουσιαστική ισότητα στην πρόσβαση στα δικαιώματα και την προστασία της αξιοπρέπειας κάθε πολίτη [89], [10].

6.4.1 Ανάλυση κόστους-οφέλους

Η ανάλυση κόστους-οφέλους στην ενσωμάτωση τεχνολογιών legal tech και τεχνητής νοημοσύνης στο δικαστικό σύστημα αποτελεί μια πολυεπίπεδη και διεπιστημονική διαδικασία, που περιλαμβάνει όχι μόνο τους στενούς λογαριασμούς κόστους λειτουργίας ή εγκατάστασης, αλλά και τις ευρύτερες κοινωνικές, θεσμικές και τεχνικές συνέπειες της τεχνολογικής καινοτομίας [10].

Η υιοθέτηση legal tech παρέχει ισχυρές δυνατότητες αυτοματοποίησης βασικών διαδικασιών, όπως η ηλεκτρονική κατάθεση δικογράφων (e-filing), η αυτοματοποιημένη δρομολόγηση και ειδοποίηση των μερών, η διαχείριση εγγράφων, οι online συνεδριάσεις και η χρήση εργαλείων αυτοματοποιημένης ανάλυσης υποθέσεων. Αυτά μειώνουν δραστικά το λειτουργικό κόστος για τα δικαστήρια και το κράτος, ενώ διευκολύνουν τους πολίτες στην

εξοικονόμηση χρόνου και πόρων. Για παράδειγμα, η ευρωπαϊκή στρατηγική e-Justice και οι μεγάλης κλίμακας εφαρμογές ODR (Online Dispute Resolution) έχουν οδηγήσει σε σημαντική μείωση του χρόνου διεκπεραίωσης και του κόστους πρόσβασης, ειδικά σε υποθέσεις μικροδιαφορών ή επαναλαμβανόμενης δικαστικής ύλης [10].

Σημαντικά τεχνικά οφέλη προκύπτουν από την εφαρμογή advanced AI [109]: η αυτοματοποίηση του classification εγγράφων, η νομική μετάφραση (με NMT ή LLMs), η αναζήτηση νομολογίας, και η χρήση legal chatbots αυξάνουν την αποτελεσματικότητα και μειώνουν το εργασιακό φορτίο για επαγγελματίες και πολίτες [74], [107]. Ιδιαίτερη σημασία έχει το γεγονός ότι τα αυτοματοποιημένα συστήματα διευκολύνουν τον μαζικό χειρισμό υποθέσεων που θα ήταν αδύνατο να διεκπεραιωθούν έγκαιρα από το υπάρχον προσωπικό, ενώ σε περιπτώσεις κρίσεων (όπως η πανδημία) επιτρέπουν τη συνέχιση της δικαιοσύνης χωρίς φυσική παρουσία [10].

Πέραν των φανερών οικονομικών ωφελημάτων, η τεχνολογική καινοτομία επιφέρει κρυφά κόστη σε πολλαπλά επίπεδα. Σε τεχνικό επίπεδο, η ανάπτυξη, συντήρηση, αναβάθμιση και διασύνδεση των πληροφοριακών συστημάτων απαιτεί διαρκείς επενδύσεις, εξειδικευμένο προσωπικό και τεχνογνωσία, καθώς και συμβατότητα με συστήματα legacy ή διεθνή πρότυπα. Η ανάγκη για cloud-based ή υβριδικές αρχιτεκτονικές, ενσωμάτωση μηχανισμών διαφάνειας (auditability, explainability), και συμμόρφωση με AI Act & GDPR προσθέτουν νέες απαιτήσεις σε χρόνο, εξοπλισμό και ανθρώπινους πόρους [89], [74], [10].

Σε κοινωνικό επίπεδο, η άνιση πρόσβαση στην τεχνολογία [10] και η απουσία επαρκούς ψηφιακού αλφαριθμητισμού ενισχύουν το ψηφιακό χάσμα. Η βιβλιογραφία υπογραμμίζει πως πολίτες με ασταθή πρόσβαση στο internet, χαμηλή τεχνολογική εξοικείωση, ή χωρίς δυνατότητα αγοράς νέου εξοπλισμού αδυνατούν να συμμετέχουν ισότιμα σε ODR, virtual hearings, ή αυτοματοποιημένες διαδικασίες. Έτσι, το κόστος μετατοπίζεται αόρατα στους πιο ευάλωτους με μορφή χαμένων υποθέσεων, ανεπαρκούς εκπροσώπησης ή δυσκολίας διεκπεραίωσης βασικών δικαστικών αναγκών. Κόστος και ποιότητα δικαστικής προστασίας

Ένα ακόμα αόρατο κόστος σχετίζεται με τη τυποποίηση των διαδικασιών και την αυτοματοποίηση της κρίσης: όσο αυξάνει η εξάρτηση από αλγοριθμικά ή rule-based συστήματα για υποθέσεις χαμηλής οικονομικής αξίας, τόσο μεγαλύτερη η απειλή για απώλεια της εξατομικευμένης δικαστικής προστασίας, της ουσιαστικής ακρόασης και της αξιολόγησης ιδιαίτερων περιστάσεων κάθε υπόθεσης. Ταυτόχρονα, η πιθανότητα εμφάνισης προκατάληψης (bias), ακούσιων σφαλμάτων ή ακόμη και αλγοριθμικού αποκλεισμού (algorithmic exclusion) παραμένει υπαρκτή, όπως αναδεικνύεται στις διεθνείς πρακτικές της διαδικτυακής επίλυσης διαφορών (ODR) και της αυτοματοποιημένης διαχείρισης υποθέσεων (automated case management).

6.4.2 Εφαρμογές πρόσβασης στη δικαιοσύνη: Μελλοντικές κατευθύνσεις και προκλήσεις

Η περαιτέρω εξέλιξη των εφαρμογών πρόσβασης στη δικαιοσύνη (Access to Justice) θα εξαρτηθεί από τον βαθμό στον οποίο οι τεχνολογικές καινοτομίες θα ενταχθούν σε ένα συνολικό πλαίσιο συμπερίληψης, θεσμικής διακυβέρνησης και κοινωνικής λογοδοσίας. Οι διεθνείς μελέτες υπογραμμίζουν ότι για να επιτευχθεί πραγματική και καθολική πρόσβαση

στη δικαιοσύνη, απαιτούνται πολιτικές και τεχνικές παρεμβάσεις που υπερβαίνουν την απλή παροχή ονλινε εργαλείων ή πλατφορμών [89], [10].

Πρώτον, κεντρικός στόχος των μελλοντικών εφαρμογών πρέπει να είναι η ενσωμάτωση πρακτικών *inclusive design*, δηλαδή ο σχεδιασμός πλατφορμών με τη συμμετοχή και τη διαρκή ανατροφοδότηση των ίδιων των ευάλωτων χρηστών, ώστε να προσαρμόζονται οι υπηρεσίες στις πραγματικές ανάγκες τους. Η εξέλιξη των *legal chatbots* και των αυτοματοποιημένων βοηθών θα πρέπει να συνοδεύεται από μηχανισμούς συνεχούς βελτίωσης, τακτικό *αυδιτίνγκ* για κοινωνικό *bias* και προγράμματα εκπαίδευσης ψηφιακού αλφαριθμητισμού, ώστε να διασφαλίζεται η αξιοποίηση των εργαλείων από τους πολίτες που το έχουν μεγαλύτερη ανάγκη [10].

Δεύτερον, η υιοθέτηση υβριδικών μοντέλων εξυπηρέτησης (*blended justice models*), που συνδυάζουν τεχνικά εργαλεία με φυσική υποστήριξη (π.χ. *legal clinics*, *community justice hubs*, υποστηρικτικά *call centers*), μπορεί να λειτουργήσει ως αντίβαρο στις εγγενείς αδυναμίες της πλήρους αυτοματοποίησης. Οι πολιτικές προσβασιμότητας του μέλλοντος θα απαιτούν σταθερή χρηματοδότηση τέτοιων δομών, θεσμική συνεργασία μεταξύ κρατικών, τοπικών και μη κερδοσκοπικών φορέων και ανοιχτή διακυβέρνηση.

Επιπλέον, προτείνεται η ανάπτυξη συστημάτων αξιολόγησης κοινωνικού αντικτύπου (*social impact assessment*), *benchmarking* των εφαρμογών με δείκτες όπως η πραγματική συμμετοχή ευάλωτων ομάδων, ο ρυθμός επιτυχούς αυτοεκπροσώπησης (*self-representation success rate*), και η διατήρηση της εμπιστοσύνης των πολιτών στο σύστημα. Χωρίς διαρκή αξιολόγηση και ρύθμιση, υπάρχει ο κίνδυνος οι "καινοτομίες" να παραμείνουν προσχηματικές ή να διευρύνουν το χάσμα αντί να το γεφυρώνουν.

Τέλος, στο ευρωπαϊκό και διεθνές πλαίσιο, η εξέλιξη των *A2J* εφαρμογών σχετίζεται όλο και περισσότερο με τις διαλειτουργικές, πολυγλωσσικές και πολυτροπικές πλατφόρμες (*multilingual/interoperable justice platforms*), με έμφαση στη διασυννοριακή αναγνώριση αποφάσεων, την αυτόματη μετάφραση και την αλληλοσύνδεση πληροφοριακών συστημάτων. Η μελλοντική προσβασιμότητα στη δικαιοσύνη θα εξαρτηθεί από τη διασφάλιση της ασφάλειας, της διαφάνειας και της πολιτισμικής προσαρμογής των τεχνολογικών λύσεων, καθώς και από τη συνεχή, ενεργή συμμετοχή της κοινωνίας των πολιτών στη διαμόρφωση και την εποπτεία τους [74], [10].

Συνοψίζοντας, η πραγματική πρόκληση για το μέλλον των εφαρμογών πρόσβασης στη δικαιοσύνη δεν είναι απλώς η τεχνολογική καινοτομία, αλλά η ολιστική ενσωμάτωσή της σε ένα οικοσύστημα συμπερίληψης, εμπιστοσύνης και κοινωνικής δικαιοσύνης, που δεν αφήνει κανέναν πίσω στην ψηφιακή εποχή του δικαίου.

6.4.3 Ζητήματα ψηφιακού χάσματος

Το ψηφιακό χάσμα αναδεικνύεται συστηματικά στη βιβλιογραφία ως μια σημαντική πρόκληση για την ισότιμη πρόσβαση στη δικαιοσύνη στη σύγχρονη εποχή της τεχνητής νοημοσύνης. Η διεθνής εμπειρία και η σχετική βιβλιογραφία δείχνουν ότι το ψηφιακό χάσμα δεν περιορίζεται στην απλή ύπαρξη ή μη τεχνολογικού εξοπλισμού ή σύνδεσης στο διαδίκτυο. Αντίθετα, αποτελεί ένα πολυπαραγοντικό φαινόμενο που περιλαμβάνει τεχνικές, κοινωνικές, οικονομικές και ψυχολογικές διαστάσεις [10], [89].

Σε πολλές χώρες και ιδιαίτερα σε ευάλωτες κοινωνικές ομάδες, η πρόσβαση σε γρήγορο και σταθερό ίντερνετ, σύγχρονες συσκευές και ασφαλή λογισμικά παραμένει ανεπαρκής. Οι πολίτες που ζουν σε αγροτικές ή απομακρυσμένες περιοχές, άτομα με χαμηλό εισόδημα ή ηλικιωμένοι συχνά στερούνται της δυνατότητας να συμμετάσχουν ισότιμα σε διαδικασίες online δικαιοσύνης ή τηλεδιασκέψεις δικαστηρίων.

Η τεχνική πρόσβαση δεν συνεπάγεται αυτόματα και ουσιαστική συμμετοχή. Η έννοια του ψηφιακού χάσματος δεύτερης γενιάς (second digital divide) αναφέρεται στην αδυναμία αξιοποίησης των τεχνολογικών δυνατοτήτων λόγω έλλειψης ψηφιακών δεξιοτήτων, αυτοπεποίθησης, ή γνώσης των βασικών αρχών προστασίας δεδομένων, διαδικτυακής ασφάλειας και χρηστικής πλοήγησης στις δικαστικές πλατφόρμες. Οι μελέτες δείχνουν ότι άτομα με χαμηλή τεχνολογική εξοικείωση εμφανίζουν σημαντικά χαμηλότερα ποσοστά επιτυχούς διεκπεραίωσης υποθέσεων ή ακόμα και συμμετοχής σε ονλινε ακροάσεις [10], [89].

Το ψηφιακό χάσμα εντείνεται από κοινωνικούς και ψυχολογικούς παράγοντες, όπως ο φόβος απέναντι στην τεχνολογία, η καχυποψία ως προς τη διαφάνεια και την αξιοπιστία των αυτοματοποιημένων διαδικασιών, και η πολιτισμική ή γλωσσική αποξένωση που βιώνουν μειονότητες και μη φυσικοί ομιλητές της κυρίαρχης γλώσσας[74]. Η δυσπιστία αυτή ενδέχεται να λειτουργήσει αποτρεπτικά ακόμη και για πολίτες που διαθέτουν επαρκή τεχνολογικά μέσα, οδηγώντας σε εθελούσιο ή ακούσιο αποκλεισμό τους από διαδικασίες νομικής τεχνολογίας [89][10].

Οι κοινωνικές επιπτώσεις του ψηφιακού χάσματος είναι εμφανείς στη μειωμένη εκπροσώπηση, στη χαμηλότερη επιτυχία αυτοεκπροσώπησης (self-representation), στη διατήρηση του φόβου απέναντι στα δικαστήρια και στη σταδιακή αποξένωση μεγάλων τμημάτων του πληθυσμού από την έννοια της δικαιοσύνης ως δημόσιο αγαθό. Η σύγχρονη βιβλιογραφία και οι policy studies προτείνουν ως αντίδοτο μια πολυεπίπεδη στρατηγική: επενδύσεις σε υλικές και τεχνικές υποδομές, διαρκή προγράμματα εκπαίδευσης ψηφιακού αλφαριθμητισμού, θεσμική ενίσχυση των τοπικών legal hubs ή δημόσιων σημείων πρόσβασης, και ενεργή συμμετοχή των ίδιων των ευάλωτων ομάδων στο σχεδιασμό και την αξιολόγηση των νέων τεχνολογιών [10].

Τέλος, είναι απαραίτητη η συνεχής μέτρηση και αξιολόγηση των επιπέδων πρόσβασης και συμμετοχής ώστε να ανιχνεύονται έγκαιρα οι νέες μορφές αποκλεισμού και να λαμβάνονται διορθωτικά μέτρα με κοινωνικά υπεύθυνο τρόπο [10].

6.5 Συμπεράσματα

Στον 21ο αιώνα, το στοίχημα της δικαιοσύνης δεν είναι μόνο τεχνικό ή θεσμικό, αλλά βαθιά ανθρώπινο και συλλογικό. Η τεχνητή νοημοσύνη και η νομική τεχνολογία έχουν τη δύναμη να γκρεμίσουν τείχη αλλά και να υψώσουν νέα. Η πρόοδος δεν αρκεί αν δεν συνοδεύεται από ευαισθησία απέναντι στις ανισότητες, ανοιχτότητα στη διαφωνία και αφοσίωση στη διαφάνεια.

Όσα αναλύθηκαν αποδεικνύουν πως καμία τεχνολογική καινοτομία δεν μπορεί να αποδώσει καρπούς χωρίς την ενεργή συμμετοχή όλων των εμπλεκομένων: μηχανικών, νομικών, κοινωνικών επιστημόνων, θεσμικών φορέων και – κυρίως – της κοινωνίας των πολιτών. Η μετάβαση σε ένα δίκαιο, ανθεκτικό και ανοιχτό οικοσύστημα δικαιοσύνης θα επιτευχθεί

μόνο μέσα από συνεχή διάλογο, στοχευμένη εκπαίδευση, στρατηγικές συμπερίληψης και συστηματική αξιολόγηση των κοινωνικών συνεπειών της τεχνολογίας.

Το μέλλον παραμένει ανοιχτό: τα ρίσκα είναι υπαρκτά αλλά και οι δυνατότητες τεράστιες. Η επιλογή είναι δική μας, να επιμείνουμε στην τεχνολογική αριστεία ως όχημα για περισσότερη ισότητα, διαφάνεια και προσβασιμότητα στη δικαιοσύνη για όλους.

7.1 Περίληψη των ευρημάτων

Η παρούσα εργασία πραγματοποίησε μια εις βάθος επισκόπηση και τεχνική ανάλυση της επίδρασης των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στο δίκαιο, καλύπτοντας το φάσμα από τη θεωρητική θεμελίωση, τα βασικά αρχιτεκτονικά χαρακτηριστικά και τις μηχανιστικές διαφορές τους, έως και την εφαρμογή τους σε πραγματικές νομικές εργασίες, συστήματα και δικαστικά περιβάλλοντα. Από την ανάλυση προκύπτει ότι τα LLMs προσφέρουν σημαντικές δυνατότητες στην ανάλυση φυσικής γλώσσας, αναγνώριση προτύπων και σύνθεση απαντήσεων, αξιοποιώντας σύγχρονες νευρωνικές αρχιτεκτονικές όπως οι μετασχηματιστές (transformers) και εμβεδδινγς υψηλής διαστατικότητας. Σε πρακτικό επίπεδο, έχουν ενισχύσει ριζικά βασικές εργασίες της νομικής επιστήμης, όπως η νομική έρευνα και η ανάκτηση πληροφορίας (legal research & information retrieval), η οποία μετασχηματίστηκε από παραδοσιακά συστήματα βασισμένα σε λέξεις-κλειδιά σε πολυδιάστατα πιπελινες σημαντις ρετριάλα, διανυσματικής αναζήτησης (vector search) και RAG. Επιπλέον, η κατηγοριοποίηση και η αυτόματη ανάλυση νομικών εγγράφων πραγματοποιούνται πλέον με ακρίβεια, ταχύτητα και δυνατότητα ανίχνευσης θεματικών ή θεσμικών συσχετισμών, ακόμα και σε μεγάλης κλίμακας corpora. Επιπρόσθετα, η χρήση των LLMs για πρόβλεψη δικαστικών αποφάσεων (judicial decision prediction), ανάλυση νομολογίας και αυτόματη παραγωγή αιτιολογίας έχει αποδώσει εντυπωσιακά αποτελέσματα σε πιλοτικές και παραγωγικές εφαρμογές σε πολλές χώρες, όπως η Κίνα, η Τουρκία, οι ΗΠΑ, η Ελβετία και οι Φιλιππίνες. Τα pipelines που συνδυάζουν σημαντις ρετριάλα, ιεραρχική επεξεργασία πραγματικών περιστατικών και συμβολικές μεθόδους αυξάνουν τη διαφάνεια και τη συνέπεια των προβλέψεων, ενώ εφαρμόζονται δομημένα φραμεворκς για αξιολόγηση και διασταύρωση αποτελεσμάτων (benchmarking, gold standard annotation).

Επιπλέον, η εξέλιξη αυτοματοποιημένων εργαλείων νομικής υποστήριξης, όπως τα chatbots, οι virtual assistants και τα εργαλεία αυτόματης παραγωγής εγγράφων, αξιοποιεί τις δυνατότητες των LLMs για παροχή προσαρμοσμένης πληροφορίας, ανάλυση συμβολαίων και υποστήριξη πολιτών ή νομικών επαγγελματιών. Οι σύγχρονες πλατφόρμες ενσωματώνουν RAG pipelines, εξειδικευμένα μοντέλα επεξεργασίας νομικών ερωτημάτων και μηχανισμούς αξιολόγησης της ακρίβειας και της αληθοφάνειας (faithfulness, citation accuracy).

Η επιστημονική και θεσμική αξιολόγηση των συστημάτων αυτών επιβάλλει την υιοθέτηση πολυδιάστατων πλαισίων αξιολόγησης που ξεπερνούν τις κλασικές μετρικές όπως accuracy,

F1 και recall, και ενσωματώνουν ειδικές μετρικές για την αποτίμηση της νομικής τεκμηρίωσης, της φαιτηφυλνεις, της ερμηνευσιμότητας, της ιχνηλασιμότητας (provenance) και της συνέπειας των παραπομπών (citation consistency). Προηγμένα frameworks, όπως το LegalBench-RAG, LeCoDe, RAGAS και LexGLUE, επιτρέπουν τη συστηματική σύγκριση μοντέλων, retrievers και τεχνικών explainability, ενώ datasets όπως το LeCaRD και το CUAD εμπλουτίζουν τις δυνατότητες εκπαίδευσης και αξιολόγησης σε πρακτικά σενάρια.

Παρά τις τεχνολογικές εξελίξεις, παραμένουν σημαντικές προκλήσεις, όπως η αλγοριθμική προκατάληψη (algorithmic bias) και τα φαινόμενα παραισθήσεων (hallucinations), που οδηγούν σε ατεκμηρίωτες ή εσφαλμένες απαντήσεις. Η περιορισμένη διαφάνεια λόγω της λειτουργίας των μοντέλων ως «μαύρα κουτιά» (black box models) δυσχεραίνει την ερμηνεία των συλλογισμών. Επιπλέον, η προσαρμογή διεθνών προτύπων αξιολόγησης σε μη αγγλόφωνες ή περιορισμένων πόρων δικαιοδοσίες, όπως η Ελλάδα, παραμένει δύσκολη, ενώ η πρόσβαση σε ανοιχτά και ποιοτικά νομικά δεδομένα αποτελεί κρίσιμο ζήτημα. Τέλος, η θεσμική πολυμορφία και τα ζητήματα απορρήτου εντείνουν την ανάγκη για συνεχή ρυθμιστική και τεχνική προσαρμογή.

Επιπλέον, η υιοθέτηση καινοτόμων τεχνολογιών διαφέρει σημαντικά μεταξύ χωρών. Η Κίνα ξεχωρίζει για την τεχνοκρατική της προσέγγιση, ενώ η Ευρωπαϊκή Ένωση και η Ελβετία εστιάζουν στον αυστηρό ρυθμιστικό έλεγχο. Αντίθετα, οι ΗΠΑ χαρακτηρίζονται από εμπορική υιοθέτηση, ενώ χώρες όπως η Τουρκία και οι Φιλιππίνες πειραματίζονται με συνεργατικά και πιλοτικά μοντέλα. Η επιτυχής ενσωμάτωση των LLMs προϋποθέτει την ισορροπία μεταξύ τεχνολογικής καινοτομίας, ρυθμιστικής επάρκειας, θεσμικής διαφάνειας και ενεργού συμμετοχής των νομικών επαγγελματιών.

Συνολικά, οι τεχνολογίες LLMs οδηγούν σε ποιοτική αναβάθμιση της νομικής επιστήμης, αυξάνοντας την αποτελεσματικότητα, την ακρίβεια και την προσβασιμότητα στη δικαιοσύνη. Ωστόσο, η πραγματική καινοτομία απαιτεί διαρκή κριτική αξιολόγηση, νομική και τεχνική διαλειτουργικότητα, διεπιστημονική συνεργασία και συνεχή προσαρμογή στα δεδομένα της τεχνητής νοημοσύνης.

7.2 Επιπτώσεις στη νομική πρακτική

Η ραγδαία ενσωμάτωση των Μεγάλων Γλωσσικών Μοντέλων (LLMs) στη νομική επιστήμη και πράξη αναδιαμορφώνει θεμελιακά το τοπίο της τεχνικής καινοτομίας, επηρεάζοντας την εξέλιξη αλγορίθμων, την αρχιτεκτονική των πληροφοριακών συστημάτων, αλλά και τη μεθοδολογία ανάπτυξης και αξιολόγησης νέων τεχνολογιών στον νομικό χώρο.

Στο πρώτο επίπεδο, η ανάπτυξη και προσαρμογή ειδικών LLMs για νομικές εφαρμογές (όπως τα LegalBERT, CaseLawBERT, GreekLegalBERT, LexLM) έχει καταστήσει σαφές ότι τα γενικού σκοπού μοντέλα δεν επαρκούν για τα σύνθετα συμφραζόμενα και τη θεσμική ορολογία του δικαίου. Η ανάγκη για υψηλής ακρίβειας σημασιολογική κατανόηση, αναγνώριση οντοτήτων, αποτύπωση λογικών συσχετίσεων και επεξεργασία πολύπλοκων ρητορικών δομών οδήγησε στην ανάπτυξη τεχνικών fine-tuning, domain adaptation και meta-learning, με έμφαση στην αξιοποίηση θεματικών και γλωσσικά εξειδικευμένων ζορπορα. Το επιστημονικό ενδιαφέρον μετατοπίζεται πλέον προς πιο σύνθετα μοντέλα multitask learning, cross-lingual transfer & prompt engineering, που διευκολύνουν τη χρήση των LLMs σε διαφορετικά νο-

μικά συστήματα και γλώσσες.

Σε αυτό το πλαίσιο, η διαθεσιμότητα δεδομένων αναδεικνύεται ως καθοριστικός παράγοντας τεχνικής προόδου. Η δημιουργία annotated datasets, benchmarking collections (π.χ. LeCaRD, LexGLUE, LeCoDe, ContractNLI) και η επισημείωση πραγματικών νομικών ερωτημάτων, ρητρών ή αποφάσεων είναι πλέον το επίκεντρο της επιστημονικής συνεργασίας μεταξύ νομικών, μηχανικών και επιστημόνων δεδομένων. Η δυσκολία συλλογής δεδομένων σε μη αγγλόφωνες έννομες τάξεις ή τομείς χαμηλής πληροφορίας αποτελεί σήμερα σημαντικό τεχνικό φραγμό, που απαιτεί λύσεις όπως μεταφορά γνώσης (transfer learning), data augmentation και αυτόματη δημιουργία συνθετικών δεδομένων μέσω LLMs.

Η τεχνική αρχιτεκτονική των συστημάτων έχει εξελιχθεί σε πολυστρωματικά, modular και επεκτάσιμα pipelines, όπου τα LLMs συνεργάζονται με συστήματα αναζήτησης (ρετρίε-ερς), βάσεις δεδομένων εστορ εμβεδδινγκς, κνωλεδγε γραπτης, μηχανισμούς ρερανκινγκ και αξιολόγησης faithfulness/citation accuracy. Η αρχιτεκτονική Retrieval-Augmented Generation (RAG) λειτουργεί ως σημείο τομής μεταξύ αναζήτησης πληροφορίας και γλωσσικής παραγωγής, προσφέροντας αξιοπιστία, επαληθευσσιμότητα και δυνατότητα ζιτατιον τραςινγκ που δεν ήταν εφικτές με κλασικά μοντέλα. Τα σύγχρονα πιπελινες ενσωματώνουν multi-hop reasoning, contextual reordering, hybrid dense+sparse retrieval και dynamic prompt assembly, ενώ ταυτόχρονα διαχειρίζονται μεγάλους όγκους δεδομένων με χαμηλό latency και κλιμακούμενη υποδομή.

Ένας ακόμη τομέας με ιδιαίτερο τεχνικό βάθος είναι η αξιολόγηση και benchmarking. Τα νομικά LLM pipelines πλέον απαιτούν πολυδιάστατη αποτίμηση. Η βιβλιογραφία δείχνει ότι οι μετρικές αυτές δεν είναι πάντα συγκλίνοντες, γεγονός που δημιουργεί νέες προκλήσεις στην επιλογή κατάλληλων benchmarks για κάθε εφαρμογή.

Η ανάπτυξη μηχανισμών ελέγχου παραισθήσεων (hallucination detection), αυδιτινγκ και explainability βρίσκεται στο επίκεντρο της έρευνας και της αγοράς. Η τεχνική κοινότητα αναπτύσσει εργαλεία για τη διασταύρωση της παραγόμενης πληροφορίας με gold standards, την ενσωμάτωση citation requirements σε prompts, τη δημιουργία transparency modules που αναδεικνύουν τα intermediate reasoning steps, και την αυτόματη ειδοποίηση για outputs αμφίβολης εγκυρότητας. Παράλληλα, υλοποιούνται εργαλεία real-time feedback, dynamic retraining και active learning με στόχο τη συνεχή βελτίωση της ποιότητας των μοντέλων σε πραγματικές συνθήκες.

Δεν πρέπει να παραβλέπεται το ζήτημα της ασφάλειας, της ιδιωτικότητας και της θεσμικής συμμόρφωσης. Οι απαιτήσεις της GDPR, του AI Act και εθνικών ρυθμιστικών πλαισίων επιβάλλουν την ανάπτυξη privacy-preserving μοντέλων, τη χρήση anonymization pipelines, federated learning για εκπαίδευση χωρίς διαμοιρασμό ευαίσθητων δεδομένων, και ενσωμάτωση auditing trails για κάθε νομικά σημαντική λειτουργία. Η τεχνική ανταπόκριση σε αυτά τα αιτήματα συχνά απαιτεί συμβιβασμούς μεταξύ αποδοτικότητας, κλίμακας, ασφάλειας και κόστους υλοποίησης.

Ένας σημαντικός άξονας τεχνικής εξέλιξης αποτελεί η εκπαίδευση και η συνεργασία μεταξύ μηχανικών και νομικών. Η ανάγκη για διεπιστημονική κατανόηση των περιορισμών, των αδυναμιών και των προσδοκιών των δύο κλάδων οδήγησε στη δημιουργία κοινών ομάδων εργασίας, workshops, legal data hackathons και ανοικτών πλαισίων τεκμηρίωσης, με σκοπό τη βελτίωση της κατανόησης και την ενίσχυση της πρακτικής συνάφειας των εργαλείων που

αναπτύσσονται.

Συνολικά, η τεχνική ανάπτυξη που πυροδότησαν τα LLMs στον νομικό τομέα μετατόπισε το επίκεντρο της καινοτομίας από την "απλή αυτοματοποίηση" σε ένα πολυδιάστατο οικοσύστημα τεχνολογίας αιχμής, όπου η τεχνική αρτιότητα, η πρακτική αξιοπιστία και η θεσμική εγκυρότητα συνυπάρχουν ως ισότιμοι στόχοι. Η προοπτική για το μέλλον περιλαμβάνει τη δημιουργία ακόμη πιο εξειδικευμένων, διαλειτουργικών και explainable LLMs, την αύξηση της χρήσης γνώσης γραφημάτων και ενισχυμένης αναζήτησης, και τη διαμόρφωση κοινών διεθνών προτύπων για την αξιολόγηση και διαλειτουργικότητα των λεγαλ AI συστημάτων. Η διαρκής εστίαση στην επιστημονική μεθοδολογία, τα ανοιχτά δεδομένα, την τεχνική επάρκεια και τη διεπιστημονική συνεργασία θα είναι καθοριστικοί παράγοντες για την επιτυχία και την υπεύθυνη εφαρμογή των LLMs στην υπηρεσία της δικαιοσύνης.

7.3 Μελλοντικές κατευθύνσεις της έρευνας

Η επιστημονική έρευνα στα Μεγάλα Γλωσσικά Μοντέλα και τις εφαρμογές τους στη νομική τεχνητή νοημοσύνη εξελίσσεται ταχύτατα, ανοίγοντας πολυδιάστατες προοπτικές. Βασικός άξονας είναι η ανάπτυξη πολυγλωσσικών και τοπικά προσαρμοσμένων LLMs, καθώς τα γενικά μοντέλα λειτουργούν καλύτερα σε δικαιοδοσίες με εκτενή αγγλόφωνα ή κινεζόφωνα δεδομένα. Για το λόγο αυτό, η έρευνα επικεντρώνεται στην αξιοποίηση τεχνικών όπως το transfer learning, η δημιουργία συνθετικών συνόλων δεδομένων και η διεθνής συνεργασία για κοινή χρήση ανοικτών νομικών δεδομένων, προκειμένου να επεκταθεί η εφαρμοσιμότητα των LLMs σε τοπικό επίπεδο.

Παράλληλα, η ενίσχυση της ερμηνευσιμότητας (explainability) και των μηχανισμών ελέγχου (auditing) αποτελεί προτεραιότητα. Η ανάπτυξη πλαισίων που τεκμηριώνουν τον συλλογισμό των LLMs πέραν των απλών οπτικοποιήσεων προσοχής είναι κρίσιμη για την αξιοπιστία και κοινωνική αποδοχή, με συστήματα που παρέχουν πλήρη διαδρομή συλλογισμού (reasoning path), τεκμηρίωση πηγών και δυναμικές καταγραφές (audit trails). Η αντιμετώπιση των παραισθήσεων (hallucinations) παραμένει κεντρικό θέμα, με έμφαση σε υβριδικές αρχιτεκτονικές retrieval-generation, μηχανισμούς διασταυρούμενης επαλήθευσης και αλγορίθμους αυτοελέγχου και ανίχνευσης αντιφάσεων, προσαρμοσμένους στα νομικά πλαίσια.

Επιπλέον, η θεσμοθέτηση του Κανονισμού για την Τεχνητή Νοημοσύνη της Ευρωπαϊκής Ένωσης (EU AI Act) καθιστά τις παραπάνω τεχνικές όχι απλώς επιθυμητές, αλλά νομικά επιβεβλημένες. Οι απαιτήσεις του Κανονισμού για ερμηνευσιμότητα, ανιχνευσιμότητα και λογική τεκμηρίωση στις αποφάσεις των συστημάτων υψηλού κινδύνου επηρεάζουν άμεσα τις εφαρμογές LLMs στη νομική πράξη. Αν και ο Κανονισμός έχει ήδη τεθεί σε ισχύ, η πλήρης εφαρμογή του προϋποθέτει σημαντική τεχνική και θεσμική προσαρμογή. Η μεταβατική αυτή περίοδος αποτελεί κρίσιμο χρονικό παράθυρο για την ανάπτυξη λύσεων που να διασφαλίζουν συμμόρφωση με το ρυθμιστικό πλαίσιο αλλά και κοινωνική αποδοχή των νομικών AI συστημάτων.

Η καθιέρωση κοινών πλαισίων αξιολόγησης (benchmarking) και η προώθηση της ανοικτής επιστήμης (open science) συμβάλλουν στην αναβάθμιση των συστημάτων μέσω πολυδιάστατων, διαγλωσσικών benchmarks που ενσωματώνουν μετρικές για την πιστότητα, ερμη-

νευσιμότητα και νομική ακρίβεια. Η λογική της επαναληψιμότητας ενισχύει τη διαφάνεια και διευκολύνει τη συνεργασία ανάμεσα σε διαφορετικούς επιστημονικούς και επαγγελματικούς τομείς.

Η θεσμική, κοινωνική και δεοντολογική διάσταση των LLMs αποκτά αυξανόμενη σημασία, με ζητήματα όπως η ισότιμη πρόσβαση, ο σαφής καθορισμός θεσμικών ορίων και η εκπαίδευση των νομικών σε ψηφιακές δεξιότητες να προβάλλουν ως κρίσιμες προκλήσεις. Η καθιέρωση δυναμικών, διεθνώς αναγνωρισμένων ρυθμιστικών πλαισίων, είναι απαραίτητη για τη διασφάλιση ασφαλούς, διαφανούς και υπεύθυνης χρήσης των συστημάτων, ενώ απαιτείται συνεχής προσαρμογή στις τεχνολογικές εξελίξεις και κοινωνικές ανάγκες ώστε να αποτραπούν νέες ανισότητες ή αποκλεισμοί.

Τέλος, η μελλοντική καινοτομία περιλαμβάνει την ανάπτυξη εξελιγμένων αυτοματοποιημένων πρακτόρων (LLM agents) που διαχειρίζονται σύνθετες νομικές ροές εργασίας, ενσωματώνοντας γνώση γραφημάτων και υποστηρίζοντας συνεχή ενημέρωση και προσαρμογή στη νομοθεσία και τη νομική ερμηνεία. Οι τεχνολογικές αυτές εξελίξεις αναμένεται να ενισχύσουν σημαντικά τη διαφάνεια, τη λογοδοσία, την ακρίβεια και την αποδοτικότητα του νομικού συστήματος, συμβάλλοντας στη βελτίωση της ποιότητας και της προσβασιμότητας της δικαιοσύνης προς όφελος της κοινωνίας συνολικά.

Βιβλιογραφία

- [1] Ninareh Mehrabi, Fred Morstatter, Nripsuta Saxena, Kristina Lerman και Aram Galstyan. *A Survey on Bias and Fairness in Machine Learning*. DOI:<http://arxiv.org/abs/1908.09635>, 2019.
- [2] Dnyanesh Panchal, Aaryan Gole, Vaibhav Narute και Raunak Joshi. *LawPal : A Retrieval Augmented Generation Based System for Enhanced Legal Accessibility in India*. DOI:<https://arxiv.org/abs/2502.16573>, 2025.
- [3] Haitao Li, Yunqiu Shao, Yueyue Wu, Qingyao Ai, Yixiao Ma και Yiqun Liu. *LeCaRDv2: A Large-Scale Chinese Legal Case Retrieval Dataset*. DOI:<https://arxiv.org/abs/2310.17609>, 2023.
- [4] Muhammad Rafsan Kabir, Rafeed Mohammad Sultan, Fuad Rahman, Mohammad Ruhul Amin, Sifat Momen, Nabeel Mohammed και Shafin Rahman. *LegalRAG: A Hybrid RAG System for Multilingual Legal Information Retrieval*. DOI: 10.48550/arXiv.2504.16121, 2025.
- [5] Odysseas S. Chlapanis, Dimitrios Galanis, Nikolaos Aletras και Ion Androutsopoulos. *GreekBarBench: A Challenging Benchmark for Free-Text Legal Reasoning and Citations*. DOI: <http://arxiv.org/abs/2505.17267>, 2025.
- [6] European Union. *REGULATION (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act)*.
- [7] Francine Ryan και Liz Hardie. *ChatGPT, I have a Legal Question? The Impact of Generative AI Tools on Law Clinics and Access to Justice*. DOI: 10.19164/ijcle.v3i1i.1401, 31:166-205, 2024.
- [8] Prof Ajit, R Patil, Devdatta Mokashi, Prof Pramod, G Rahate, Prof Prashant Mahadik, Prof Neeraj, Arun Gangurde, Amolkumar N Jadhav, Dr Rajendra και B Mohite. *Legal Rules and Regulations Document Summarizer: Regulatory Compliance with NLP, ML, and LLMs*. URL:www.bpasjournals.com.
- [9] Olaniyi Evans, Emeka Osuji, Olawale Ayoola, W Frank Barton, Olawale Wale-Awe, Osuji Emeka, Olawale O Ayoola, Raymond Alenoghena και Sesan Adeniji. *ChatGPT impacts on access-efficiency, employment, education and ethics: The socio-economics of an AI language model*. URL:<https://bizecons.5profz.com/>, 16:1-17, 2023.

- [10] David Freeman ENngstrom. *Legal Tech and the Future of Civil Justice*. Cambridge: Cambridge University Press, 2023.
- [11] Nguyen Ha Thanh και Ken Satoh. *KRAG Framework for Enhancing LLMs in the Legal Domain*. DOI:<http://arxiv.org/abs/2410.07551>, 2024.
- [12] Michael Benedict L. Virtucio, Jeffrey A. Aborot, John Kevin C. Abonita, Roxanne S. Avinante, Rother Jay B. Copino, Michelle P. Neverida, Vanesa O. Osiana, Elmer C. Peramo, Joanna G. Syjuco και Glenn Brian A. Tan. *Predicting Decisions of the Philippine Supreme Court Using Natural Language Processing and Machine Learning*. DOI: 10.1109/COMPSAC.2018.10348, 2018.
- [13] Stanley Greenstein. *Preserving the rule of law in the era of artificial intelligence (AI)*. DOI:10.1007/s10506-021-09294-4, 2022.
- [14] Irene Benedetto, Luca Cagliero, Michele Ferro, Francesco Tarasconi, Claudia Bernini και Giuseppe Giacalone. *Leveraging large language models for abstractive summarization of Italian legal news*. DOI:10.1007/s10506-025-09431-3, 2025.
- [15] Uchechukwu Kizito, Ogu Phd, Chidiebere Peter και Okechukwu Phd. *ARTIFICIAL INTELLIGENCE AND THE SOCIETY: NAVIGATING THE IMPACT OF AI ON JUSTICE, FREEDOM AND HUMAN RIGHTS*. URL: <https://acjoi.org/index.php/NJP/article/view/6330>, 2024.
- [16] Jacob Devlin, Ming Wei Chang, Kenton Lee, Kristina Toutanova Google και A I Language. *BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding*. DOI:10.18653/v1/N19-1423.
- [17] Chenguang Wang, Mu Li και Alexander J. Smola. *Language Models with Transformers*. DOI: <http://arxiv.org/abs/1904.09408>, 2019.
- [18] Zihang Dai, Zhilin Yang, Yiming Yang, Jaime Carbonell, Quoc V. Le και Ruslan Salakhutdinov. *Transformer-XL: Attentive Language Models Beyond a Fixed-Length Context*. DOI:<http://arxiv.org/abs/1901.02860>, 2019.
- [19] Lazar Peric, Stefan, Mijic, Dominik Stambach, Elliott Ash, Lazar Peric, Stefan Mijic, Dominik Stambach και Eth Zurich. *ETH Library Legal Language Modeling with Transformers Conference Paper Rights / license: Creative Commons Attribution 4.0 International Legal Language Modeling with Transformers*. DOI:<https://doi.org/10.3929/ethz-b-000456079>, 2020.
- [20] Amit Tewari. *LegalPro-BERT: Classification of Legal Provisions by fine-tuning BERT large language model*. DOI:<https://orcid.org/0000-0002-6821-3833>.
- [21] Mi Young Kim, Julianio Rabelo, Housam Khalifa Bashier Babiker, Md Abed Rahman και Randy Goebel. *Legal Information Retrieval and Entailment Using Transformer-based Approaches*. DOI:10.1007/s12626-023-00153-z, 18:101–121, 2024.

- [22] Andrew Vold, Thomson Reuters και Jack G Conrad. *Using Transformers to Improve Answer Retrieval for Legal Questions*. DOI:10.1145/3462757, 2021.
- [23] Nadia Mushtaq Gardazi, Ali Daud, Muhammad Kamran Malik, Amal Bukhari, Tariq Alsahfi και Bader Alshemaimri. *BERT applications in natural language processing: a review*. DOI:10.1007/s10462-025-11162-5, 2025.
- [24] Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer και Veselin Stoyanov. *RoBERTa: A Robustly Optimized BERT Pretraining Approach*. DOI:http://arxiv.org/abs/1907.11692, 2019.
- [25] Victor Sanh, Lysandre Debut, Julien Chaumond και Thomas Wolf. *DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter*. DOI:http://arxiv.org/abs/1910.01108, 2020.
- [26] Tom B. Brown, 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 και Dario Amodei. *Language Models are Few-Shot Learners*. URL:http://arxiv.org/abs/2005.14165, 2020.
- [27] Noah M Kenney. *A Brief Analysis of the Architecture, Limitations, and Impacts of ChatGPT*. DOI:10.5281/zenodo.7762245, 2023.
- [28] Aditya S Shethiya. *Academia Nexus Journal LLM-Powered Architectures: Designing the Next Generation of Intelligent Software Systems*. URL:https://academianexusjournal.com/index.php/anj/article/view/21, 2, 2023.
- [29] Loïc Kwate Dassi και Loic Kwate Dassi. *Legal-BigBird: An Adapted Long-Range Transformer for Legal Documents*. DOI:10.13140/RG.2.2.14172.92802, 2021.
- [30] John Armour, Richard Parnham και Mari Sako. *Augmented Lawyering*. DOI:10.2139/ssrn.3688896, 2020.
- [31] Junyun Cui, Xiaoyu Shen, Feiping Nie, Zheng Wang, Jinglong Wang και Yulong Chen. *A Survey on Legal Judgment Prediction: Datasets, Metrics, Models and Challenges*. DOI:http://arxiv.org/abs/2204.04859, 2022.
- [32] Daniel Schwarcz, Sam Manning, Patrick Barry, David R Cleveland, JJ Prescott, Beverly Rich, Pablo Arredondo, Victor Bennett, Jonathan Choi, Miryam Gorelashvili, Dan Ho, Bryan Mechell, Amy Monahan, Daniel Rock, James Snelson και Tim Sullivan. *AI-POWERED Lawyering: AI Reasoning Models, Retrieval Augmented Generation and the Future of Legal Practice*. DOI:10.2139/ssrn.5162111, 2023.

- [33] Katie Atkinson, Trevor Bench-Capon και Danushka Bollegala. *Explanation in AI and law: Past, present and future*. DOI:10.1016/j.artint.2020.103387, 2020.
- [34] Zico Junius Fernando και Ariesta Wibisono Anditya. *AI on The Bench: The Future of Judicial Systems in The Age of Artificial Intelligence*. DOI:10.25216/jhp.13.3.2024.523-550, 2024.
- [35] Ilias Chalkidis, Manos Fergadiotis, Prodromos Malakasiotis και Nikolaos Aletras. *Findings of the Association for Computational Linguistics LEGAL-BERT: The Muppets straight out of Law School*. DOI:10.18653/v1/2020.findings-emnlp.261, 2020.
- [36] Catalina Goanta, Nikolaos Aletras, Ilias Chalkidis, Sofia Ranchordas και Gerasimos Spanakis. *Regulation and NLP (RegNLP): Taming Large Language Models*. DOI:http://arxiv.org/abs/2310.05553, 2023.
- [37] Chaojun Xiao, Xueyu Hu, Zhiyuan Liu, Cunchao Tu και Maosong Sun. *Lawformer: A pre-trained language model for Chinese legal long documents*. DOI:10.1016/j.aiopen.2021.06.003, 2, 2021.
- [38] Aldo Gangemi. *Design patterns for legal ontology construction*. 2007.
- [39] Marco Tulio Ribeiro, Sameer Singh και Carlos Guestrin. "Why should i trust you?" *Explaining the predictions of any classifier*. *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, σελίδες 1135–1144. DOI:10.1145/2939672.2939778, 2016.
- [40] Scott Lundberg και Su In Lee. *A Unified Approach to Interpreting Model Predictions*. DOI:http://arxiv.org/abs/1705.07874, 2017.
- [41] Michael Feldman, Sorelle Friedler, John Moeller, Carlos Scheidegger και Suresh Venkatasubramanian. *Certifying and removing disparate impact*. DOI:http://arxiv.org/abs/1412.3756, 2014.
- [42] Ilias Chalkidis, Abhik Jana, Dirk Hartung, Michael Bommarito, Daniel Martin Katz και Nikolaos Aletras. *LexGLUE: A Benchmark Dataset for Legal Language Understanding in English*. DOI:10.18653/v1/2022.acl-long.297, σελίδες 4310–4330.
- [43] Nicholas Pipitone και Ghita Houir Alami. *LegalBench-RAG: A Benchmark for Retrieval-Augmented Generation in the Legal Domain*. DOI:http://arxiv.org/abs/2408.10343, 2024.
- [44] Raquel Silveira, Caio Ponte, Vitor Almeida, Vlória Pinheiro και Vasco Furtado. *LegalBert-pt: A Pretrained Language Model for the Brazilian Portuguese Legal Domain*. DOI:10.1007/978-3-031-45392-218, σελίδες 268–282, 2023.
- [45] Stella Douka, Hadi Abdine, Michalis Vazirgiannis, Rajaa El Hamdani και David Restrepo Amariles. *JuriBERT: A Masked-Language Model Adaptation for French Legal Text*. DOI:http://arxiv.org/abs/2110.01485, 2022.

- [46] Sarat Kiran. *Hybrid Retrieval-Augmented Generation (RAG) Systems with Embedding Vector Databases*. DOI:10.32628/CSEIT25112702, σελίδες 2694–2702, 2025.
- [47] Lee F Peoples. *Artificial Intelligence and Legal Analysis: Implications for Legal Education and the Profession*. DOI:10.2139/ssrn.5123122, 2025.
- [48] Tonya Custis, Frank Schilder, Thomas Vacek, Thomson Reuters, Gayle McElvain και Hector Martinez Alonso. *Westlaw Edge AI features demo: KeyCite Overruling Risk, Litigation Analytics, and WestSearch plus*. *Proceedings of the 17th International Conference on Artificial Intelligence and Law, ICAIL 2019*, σελίδες 256–257. DOI:10.1145/3322640.3326739, 2019.
- [49] João Albertode Oliveira Lima. *Unlocking Legal Knowledge with Multi-Layered Embedding-Based Retrieval*. DOI:http://arxiv.org/abs/2411.07739, 2024.
- [50] Xiaoxian Yang, Zhifeng Wang, Qi Wang, Ke Wei, Kaiqi Zhang και Jiangang Shi. *Large language models for automated Q&A involving legal documents: a survey on algorithms, frameworks and applications*. DOI:10.1108/IJWIS-12-2023-0256, 20:413–435, 2024.
- [51] Fábio Souza, Rodrigo Nogueira και Roberto Lotufo. *BERTimbau: Pretrained BERT Models for Brazilian Portuguese*. DOI:10.1007/978-3-030-61377-828, σελίδες 403–417, 2020.
- [52] Vuong Pham, Hoang Huy Le, Thinh Phu Ngo, Binh Nguyen, Diem Nguyen και Hien D. Nguyen. *Enhancing legal research through knowledge-infused information retrieval for Vietnamese labor law*. DOI:10.11591/tjai.v13.i4.pp3962-3973, σελίδες 3962–3973, 2024.
- [53] Jiale Wei, Shuchi Wu, Ruochen Liu, Xiang Ying, Jingbo Shang και Fangbo Tao. *Tuning LLMs by RAG Principles: Towards LLM-native Memory*. DOI:http://arxiv.org/abs/2503.16071, 2025.
- [54] Masha Medvedeva, Martijn Wieling και Michel Vols. *Rethinking the field of automatic prediction of court decisions*. DOI:10.1007/s10506-021-09306-3, 31:195–212, 2023.
- [55] Zhongxiang Sun. *A Short Survey of Viewing Large Language Models in Legal Aspect*. DOI:http://arxiv.org/abs/2303.09136, 2023.
- [56] Jiaqi Wang. *Rainbow RAG: An LLM-Powered RAG System for Contract Review*. Master's thesis, Delft University of Technology, Delft, Netherlands, 2023.
- [57] Emre Mumcuoğlu, Ceyhun E. Öztürk, Haldun M. Ozaktas και Aykut Koç. *Natural language processing in law: Prediction of outcomes in the higher courts of Turkey*. DOI:10.1016/j.ipm.2021.102684, 58, 2021.
- [58] André Lage-Freitas, Héctor Allende-Cid, Orivaldo Santana και Livia Oliveira-Lage. *Predicting Brazilian Court Decisions*. DOI:10.7717/peerj-cs.904, 2022.

- [59] Intisar Almuslim και Diana Inkpen. *Legal Judgment Prediction for Canadian Appeal Cases*. 2022 7th International Conference on Data Science and Machine Learning Applications (CDMA), σελίδες 163–168. DOI:10.1109/CDMA54072.2022.00032, 2022.
- [60] Shubham Kumar Nigam, Balaramamahanthi Deepak Patnaik, Shivam Mishra, Noel Shallum, Kripabandhu Ghosh και Arnab Bhattacharya. *TathyaNyaya and FactLegalLlama: Advancing Factual Judgment Prediction and Explanation in the Indian Legal Context*. DOI:http://arxiv.org/abs/2504.04737, 2025.
- [61] Zlatan Morić, Vedran Dakić και Siniša Urošev. *An AI-Based Decision Support System Utilizing Bayesian Networks for Judicial Decision-Making*. DOI:10.3390/systems13020131, 2025.
- [62] Mohamed Bayan Kmainasi, Ali Ezzat Shahroor και Amani Al-Ghraibah. *Can Large Language Models Predict the Outcome of Judicial Decisions?* DOI:http://arxiv.org/abs/2501.09768, 2025.
- [63] Alex Biedermann, Timothy Lau και Ronald Allen. *Assessing AI Output in Legal Decision-Making with Nearest Neighbors Articles: Assessing AI Output in Legal Decision-Making with Nearest Neighbors*. URL:https://elibrary.law.psu.edu/pslr/vol124/iss3/1, 2020.
- [64] Ali Shariq Imran, Henrik Hodnefeld, Zenun Kastrati, Noureen Fatima, Sher Muhammad Daudpota και Ahmad Wani. *Classifying European Court of Human Rights Cases Using Transformer-Based Techniques*. DOI:10.1109/ACCESS.2017.DOI.
- [65] Nu Wang. *'Black Box Justice': Robot Judges and AI-based Judgment Processes in China's Court System*. DOI:10.1109/ISTAS50296.2020.9462216, 2020-Νοεμβέρ:58–65, 2020.
- [66] Narges Farzaneh Kondori και Saeed Rouhani. *Identifying and ranking critical success factors for digital court*. DOI:10.1108/JOCM-03.
- [67] Changqing Shi, Tania Sourdin και Bin Li. *The Smart Court - A New Pathway to Justice in China?* DOI:10.36745/ijca.367, σελίδες 1–19, 2021.
- [68] Olha Kovalchuk & Vladyslav Teremetskyi. *SMART TECHNOLOGIES IN JUSTICE: PERSPECTIVES FOR UKRAINE*. DOI:10.32849/2663-5313/2023.5.13.
- [69] Yusuf Dwi Prasetyo και Dani Muhtada. *The Effectiveness of Using E-Court Services in Realizing Electronic Religious Courts*. DOI:10.21043/qjjs.v10i2.17281, 2025.
- [70] Alicia L. Bannon και Douglas Keith. *Remote Court: Principles for Virtual Proceedings During the COVID-19 Pandemic and Beyond*. URL:https://scholarlycommons.law.northwestern.edu/nulr/vol115/iss6/7/, 115(6):1875–1904, 2021.
- [71] Rahman S M Wahidur, Sumin Kim, Haeung Choi, David S Bhatti, Heungno Lee και Senior Member. *Legal Query RAG*. DOI:10.1109/ACCESS.2023.DOI.

- [72] Youssra Amazou, Faouzi Tayalati, Houssam Mensouri, Abdellah Azmani και Monir Azmani. *Accurate AI Assistance in Contract Law Using Retrieval-Augmented Generation to Advance Legal Technology*. DOI:10.14569/IJACSA.2025.01602113, 16, 2025.
- [73] Xiaotong Bing. *The Path of Formulating the Basic Law of Artificial Intelligence in China – Analysis of the Desirability of the EU Artificial Intelligence Act*. DOI:10.56397/slj.2023.09.09, 2:68–73, 2023.
- [74] Joel Niklaus, Jakob Merane, Luka Nenadic, Sina Ahmadi, Yingqiang Gao, Cyrill A. H. Chevalley, Claude Humbel, Christophe Gösken, Lorenzo Tanzi, Thomas Lüthi, Stefan Palombo, Spencer Poff, Boling Yang, Nan Wu, Matthew Guilloid, Robin Mamié, Daniel Brunner, Julio Pereyra και Niko Grupen. *SwiLTra-Bench: The Swiss Legal Translation Benchmark*. DOI:http://arxiv.org/abs/2503.01372, 2025.
- [75] Adalmirde Oliveira Gomes, Tomasde Aquino Guimaraes και Luiz Akutsu. *The relationship between judicial staff and court performance: Evidence from Brazilian state courts*. DOI:10.18352/ijca.214, 8:12–19, 2016.
- [76] Yash Mahajan, Matthew Freestone, Sathyanarayanan Aakur και Santu Karmaker. *Revisiting Word Embeddings in the LLM Era*. DOI:http://arxiv.org/abs/2502.19607, 2025.
- [77] Weikang Yuan, Kaisong Song, Zhuoren Jiang, Junjie Cao, Yujie Zhang, Jun Lin, Kun Kuang, Ji Zhang και Xiaozhong Liu. *LeCoDe: A Benchmark Dataset for Interactive Legal Consultation Dialogue Evaluation*. DOI: https://arxiv.org/abs/2505.19667, 2025.
- [78] Rajaa El Hamdani, Thomas Bonald, Fragkiskos Malliaros, Nils Holzenberger και Fabian Suchanek. *The Factuality of Large Language Models in the Legal Domain*. DOI:http://arxiv.org/abs/2409.11798, 2024.
- [79] Dan Hendrycks, Collin Burns, Anya Chen και Spencer Ball. *CUAD: An Expert-Annotated NLP Dataset for Legal Contract Review*. DOI:http://arxiv.org/abs/2103.06268, 2021.
- [80] Jiuzhou Han, Paul Burgess και Ehsan Shareghi. *Evaluating LLM-based Approaches to Legal Citation Prediction: Domain-specific Pre-training, Fine-tuning, or RAG? A Benchmark and an Australian Law Case Study*. DOI:http://arxiv.org/abs/2412.06272, 2024.
- [81] Yejin Bang, Ziwei Ji, Alan Schelten, Anthony Hartshorn, Tara Fowler, Cheng Zhang, Nicola Cancedda και Pascale Fung. *HalluLens: LLM Hallucination Benchmark*. DOI:http://arxiv.org/abs/2504.17550, 2025.
- [82] Sandra Wachter, Brent Mittelstadt και Luciano Floridi. *Transparent, explainable, and accountable AI for robotics*. DOI:10.1126/scirobotics.aan6080, 2, 2017.

- [83] Hanze Sun. *Bridging the Accountability Gap in AI Decision-Making: An Integrated Analysis of Legal Precedents and Scholarly Perspectives*. DOI:10.54691/24vexn67, σελίδα 2025, 2025.
- [84] Marina Matić Bošković. *Implications of EU AI regulation for criminal justice*. DOI:10.56461/iuprlrc.2024.5.ch8, 2024.
- [85] Jiayi Ye, Yanbo Wang, Yue Huang, Dongping Chen, Qihui Zhang, Nuno Moniz, Tian Gao, Werner Geyer, Chao Huang, Pin Yu Chen, Nitesh V Chawla και Xiangliang Zhang. *Justice or Prejudice? Quantifying Biases in LLM-as-a-Judge*. DOI:http://arxiv.org/abs/2410.02736, 2024.
- [86] Jumai Adedaja Fabuyi. *Leveraging Synthetic Data as a Tool to Combat Bias in Artificial Intelligence (AI) Model Training*. DOI:10.9734/jerr/2024/v26i121337, 26:24–46, 2024.
- [87] Angela Huyue Zhang. *The Promise and Perils of China’s Regulation of Artificial Intelligence*. DOI:10.2139/ssrn.4708676, 2024.
- [88] Elizabeth Edenberg και Alexandra Wood. *Disambiguating Algorithmic Bias: From Neutrality to Justice*. AIES 2023 - Proceedings of the 2023 AAAI/ACM Conference on AI, Ethics, and Society. DOI:10.1145/3600211.3604695, 2023.
- [89] Oakley Parker. *Data Governance and Ethical AI: Developing Legal Frameworks to Address Algorithmic Bias and Discrimination*. DOI:10.13140/RG.2.2.11601.34406, 2024.
- [90] Yufei Guo, Muzhe Guo, Juntao Su, Zhou Yang, Mengqiu Zhu, Hongfei Li, Mengyang Qiu και Shuo Shuo Liu. *Bias in Large Language Models: Origin, Evaluation, and Mitigation*. DOI:http://arxiv.org/abs/2411.10915, 2024.
- [91] Laura Weidinger, John Mellor, Maribeth Rauh, Conor Griffin, Jonathan Uesato, Po Sen Huang, Myra Cheng, Mia Glaese, Borja Balle, Atoosa Kasirzadeh, Zac Kenton, Sasha Brown, Will Hawkins, Tom Stepleton, Courtney Biles, Abeba Birhane, Julia Haas, Laura Rimell, Lisa Anne Hendricks, William Isaac, Sean Legassick, Geoffrey Irving και Iason Gabriel. *Ethical and social risks of harm from Language Models*. DOI:http://arxiv.org/abs/2112.04359, 2021.
- [92] Yiran Hu, Huanghai Liu, Qingjing Chen, Ning Zheng, Chong Wang, Yun Liu, Charles L. A. Clarke και Weixing Shen. *J&H: Evaluating the Robustness of Large Language Models Under Knowledge-Injection Attacks in Legal Domain*. DOI:http://arxiv.org/abs/2503.18360, 2025.
- [93] Nadjia Madaoui. *The Impact of Artificial Intelligence on Legal Systems: Challenges and Opportunities*. DOI:10.21564/2414-990x.164.289266, 1:285–303, 2024.
- [94] Amal Fawzy Ahmed Awad. *Predictive Justice and the Age of Artificial Intelligence*. DOI:10.37136/0515-013-999-064, 2021.

- [95] Marta Gamito Cantero, Giulia Gentile και Marta Cantero Gamito. *ALGORITHMS, RULE OF LAW, AND THE FUTURE OF JUSTICE: IMPLICATIONS IN THE ESTONIAN JUSTICE SYSTEM*. DOI:10.2870/640834, 2023.
- [96] Alexandros Tassios, Stergios Tegos, Christos Bouas, Konstantinos Manousaridis, Maria Papoutsoglou, Maria Kaltsa, Eleni Dimopoulou, Thanassis Mavropoulos, Stefanos Vrochidis και Georgios Meditskos. *LLM Performance in Low-Resource Languages: Selecting an Optimal Model for Migrant Integration Support in Greek*. DOI:10.3390/fi17060235, 2025.
- [97] Marios Koniaris, Dimitris Galanis, Eugenia Giannini και Panayiotis Tsanakas. *Evaluation of Automatic Legal Text Summarization Techniques for Greek Case Law*. DOI:10.3390/info14040250, 2023.
- [98] John Koutsikakis, Ilias Chalkidis, Prodromos Malakasiotis και Ion Androutsopoulos. *GREEK-BERT: The greeks visiting sesame street*. *ACM International Conference Proceeding Series*, σελίδες 110–117. DOI:10.1145/3411408.3411440, 2020.
- [99] Antonios Karampatzos και Georgios Malos. *The Role of Case Law and the Prospective Overruling in the Greek Legal System*, σελίδες 163–184. DOI:10.1007/978-3-319-16175-4_7, 2015.
- [100] Shai Farber. *Harmonizing AI and human instruction in legal education: a case study from Israel on training future legal professionals*. DOI:10.1080/09695958.2024.2430018, 2024.
- [101] Danish Anthal και Raj Kumar. *Revolutionizing Clinical Legal Education in India through Corporate Social Responsibility*, 2025.
- [102] Kaan Thilakarathna, Menaka Harankaha και D R Menaka Harankaha. *INTERNATIONAL JOURNAL OF LAW MANAGEMENT & HUMANITIES AI in Legal Education: Balancing Opportunities and Challenges in the Sri Lankan Context*. DOI:https://doi.org/10.1000/IJLMH.118515, 2024.
- [103] Anisa Shaikh και Bharati Vidyapeeth. *Critical Examination of The Role & Impact of Artificial Intelligence (AI) in Indian Legal Education* THE JOURNAL OF ORIENTAL RESEARCH MADRAS. DOI:10.54691/24vexn67.
- [104] Hanze Sun. *Bridging the Accountability Gap in AI Decision-Making: An Integrated Analysis of Legal Precedents and Scholarly Perspectives*. DOI:10.54691/24vexn67, 5.
- [105] Patritzia Giamperi. *Assessing the Quality of AI and MT in Legal Translation*. DOI:10.54103/2035-7680/28886, 2025.
- [106] Omkar Komera. *Explainable AI (XAI): Bridging the Gap Between Complexity and Interpretability*. Preprint, ResearchGate, 2024.

- [107] Philipp Hacker, Ralf Krestel, Stefan Grundmann και Felix Naumann. *Explainable AI under contract and tort law: legal incentives and technical challenges*. DOI:10.1007/s10506-020-09260-6, 28, 2020.
- [108] Paul Burgess, Iwan Williams, Lizhen Qu και Weiqing Wang. *Using Generative AI to Identify Arguments in Judges' Reasons: Accuracy and Benefits for Students*. DOI:10.5204/lthj.3637, 6, 2024.
- [109] M. Mikail Demir, Hakan T. Otal και M. Abdullah Canbaz. *LegalGuardian: A Privacy-Preserving Framework for Secure Integration of Large Language Models in Legal Practice*. DOI:<http://arxiv.org/abs/2501.10915>, 2025.