



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

Ανακάλυψη Φαρμάκων με χρήση
Βαθιάς Μάθησης

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ

ΙΩΑΝΝΗΣ ΜΕΤΖΑΚΗΣ

Επιβλέπων: Ανδρέας-Γεώργιος Σταφυλοπάτης
Καθηγητής Ε.Μ.Π.

Αθήνα, Νοέμβριος 2024



Εθνικό Μετσόβιο Πολυτεχνείο
Σχολή Ηλεκτρολόγων Μηχανικών και Μηχανικών Υπολογιστών
Τομέας Τεχνολογίας Πληροφορικής και Υπολογιστών

Ανακάλυψη Φαρμάκων με χρήση Βαθιάς Μάθησης

ΔΙΠΛΩΜΑΤΙΚΗ ΕΡΓΑΣΙΑ
ΙΩΑΝΝΗΣ ΜΕΤΖΑΚΗΣ

Επιβλέπων: Ανδρέας-Γεώργιος Σταφυλοπάτης
Καθηγητής Ε.Μ.Π.

Εγκρίθηκε από την τριμελή εξεταστική επιτροπή την 7η Νοεμβρίου 2024.

(Υπογραφή)

(Υπογραφή)

(Υπογραφή)

.....
Ανδρέας-Γεώργιος
Σταφυλοπάτης
Καθηγητής ΕΜΠ

.....
Γιώργος Στάμου
Καθηγητής ΕΜΠ

.....
Στέφανος Κόλλιας
Ομότιμος Καθηγητής

Αθήνα, Νοέμβριος 2024



Εθνικό Μετσόβιο Πολυτεχνείο
Σχολή Ηλεκτρολόγων Μηχανικών και Μηχανικών Υπολογιστών
Τομέας Τεχνολογίας Πληροφορικής και Υπολογιστών

(Υπογραφή)

.....

Ιωάννης Μετζάκης

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

Copyright © Ιωάννης Μετζάκης, 2024.

Με την επιφύλαξη παντός δικαιώματος. All rights reserved

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. Ερωτήματα που αφορούν τη χρήση της εργασίας για κερδοσκοπικό σκοπό πρέπει να απευθύνονται προς το συγγραφέα.

Οι απόψεις και τα συμπεράσματα που περιέχονται σε αυτό το έγγραφο εκφράζουν τον συγγραφέα και δεν πρέπει να ερμηνευθεί ότι αντιπροσωπεύουν τις επίσημες θέσεις του Εθνικού Μετσόβιου Πολυτεχνείου.

Περίληψη

Η παρούσα διπλωματική εργασία παρουσιάζει την εφαρμογή τεχνικών Βαθιάς Μάθησης και συγκεκριμένα παραγωγικών επαναλαμβανόμενων νευρωνικών δικτύων (RNNs) με χρήση Δικτύων Μακράς Βραχυπρόθεσμης Μνήμης (LSTM) για εκ νέου σχεδιασμό φαρμάκων (de novo drug design), με στόχο τον περιορισμό της αναζήτησης στον τεράστιο χημικό χώρο και τη βελτιστοποίηση της συγκεκριμένης διαδικασίας. Για τα δεδομένα εκπαίδευσης και επαλήθευσης του μοντέλου έχει χρησιμοποιηθεί η βάση ChEMBL22.

Παραδοσιακά, οι περισσότερες μέθοδοι για δημιουργία νέων φαρμάκων βασίζονταν στον έλεγχο των υφιστάμενων ενώσεων. Αναπαραστήνοντας όμως τις μοριακές δομές ως SMILES συμβολοσειρές, το LSTM μοντέλο συλλαμβάνει τη σύνταξη αυτών των αναπαραστάσεων με μεγάλη ακρίβεια, δίνοντας τη δυνατότητα για απευθείας δημιουργία νέων ενώσεων. Πέρα από την εκ νέου δημιουργία ενώσεων, το οποίο είναι και το βασικό αντικείμενο της εργασίας, αναλύεται η διαδικασία βελτιστοποίησης του μοντέλου με εφαρμογή της τεχνικής μεταφοράς μάθησης σε συγκεκριμένα σύνολα δεδομένων-στόχων, επιτρέποντας έτσι τη δημιουργία μορίων που είναι δομικά παρόμοια με γνωστές βιοδραστικές ενώσεις, και ελαχιστοποιώντας την εξάρτηση από εκτεταμένα δεδομένα εκπαίδευσης. Τέλος, εξετάζεται και η τεχνική δημιουργίας φαρμάκων με βάση κάποιο υπάρχον τμήμα του τελικού στόχου (fragment-based drug discovery / FBDD), κατά την οποία η διαδικασία παραγωγής ξεκινάει από ήδη γνωστά ενεργά μόρια, επεκτείνοντάς τα σε μεγαλύτερες, πιο σύνθετες μοριακές δομές.

Συνεπώς, η διπλωματική αυτή εστιάζει στη δημιουργία νέων φαρμάκων με βάση συγκεκριμένες προδιαγραφές-στόχους, στη βελτιστοποίηση μέσω μεταφοράς μάθησης σε διαφορετικά δεδομένα-στόχους, καθώς και στον σχεδιασμό μορίων με προυπάρχουσα βάση, ακόμα και σε σενάρια χαμηλών δεδομένων.

Λέξεις Κλειδιά:

Βαθιά Μάθηση, Επαναλαμβανόμενα Νευρωνικά Δίκτυα, Μεταφορά Μάθησης, Δίκτυα Μακράς Βραχυπρόθεσμης Μνήμης, Εκ νέου σχεδιασμός Φαρμάκων, Ανακάλυψη Φαρμάκων, RNN, LSTM, FBDD

Abstract

The following thesis presents the application of Deep Learning techniques, and specifically generative recurrent neural networks (RNNs) using Long Short-Term Memory networks (LSTMs) for de novo drug design, with the goal of narrowing down the vast chemical space and optimizing the drug discovery process. ChEMBL22 dataset has been used for model training and validation.

Traditionally, most methods for discovering new drugs relied on screening existing compounds. However, by representing molecular structures as SMILES strings, the LSTM model captures the syntax of these representations with great accuracy, enabling the direct generation of new compounds. Beyond the generation of novel compounds, which is the main scope of this thesis, the process of model optimization through the application of transfer learning techniques on specific target datasets, allowing the generation of molecules that are structurally similar to known bioactive compounds and minimizing the dependency on extensive training data, is also analyzed. Finally, the technique of fragment-based drug discovery is examined, in which the drug generation process initiates from already known active molecules and extends them into larger, more complex molecular structures.

Thus, this thesis focuses on the generation of new drugs based on specific target specifications, optimization through transfer learning across different target datasets, as well as the design of molecules from pre-existing scaffolds, even in low-data scenarios.

Keywords:

Deep Learning, Recurrent Neural Networks, Transfer Learning, Long Short-Term Memory Networks, De Novo Drug Design, Drug Discovery, Fragment Based Drug Discovery, RNN, LSTM, FBDD

Ευχαριστίες

Για αρχή, θα ήθελα να ευχαριστήσω τον καθηγητή κ. Ανδρέα - Γεώργιο Σταφυλοπάτη για την ευκαιρία που μου έδωσε να εκπονήσω το συγκεκριμένο θέμα, αναλαμβάνοντας την επίβλεψη της παρούσας εργασίας, καθώς και για την σημαντική βοήθειά του. Επίσης, πολύτιμη υπήρξε και η βοήθεια του κ. Γεώργιου Σιόλα, ο οποίος με αρκετή κατανόηση και προθυμότητα καθ' όλη τη διάρκεια της εκπόνησης της διπλωματικής, με καθοδηγούσε με καθοριστικής σημασίας συμβουλές και κατευθύνσεις. Είμαι βαθιά ευγνώμων για τη στήριξή τους και τη δυνατότητα που μου δώθηκε να εκπονήσω τη διπλωματική μου στο Εργαστήριο Συστημάτων Τεχνητής Νοημοσύνης και Μάθησης.

Ακόμα, θα ήθελα να ευχαριστήσω τους φίλους μου για την συμπαράσταση που μου έδειχναν όλα αυτά τα χρόνια, ενώ τέλος, θα ήθελα να αφιερώσω αυτή τη διπλωματική στην οικογένειά μου, η οποία με συνεχή φροντίδα και αγάπη με στηρίζει σε κάθε μου βήμα.

Περιεχόμενα

Περίληψη	6
Abstract	7
Περιεχόμενα.....	11
Κεφάλαιο 1: Εισαγωγή	13
Κεφάλαιο 2: Θεωρητικό Υπόβαθρο	14
2.1 Βαθιά Μάθηση.....	14
2.1.1 Εισαγωγή.....	14
2.1.2 Νευρωνικά Δίκτυα.....	15
2.1.3 Συναρτήσεις Ενεργοποίησης	16
2.1.4 Συναρτήσεις Κόστους	19
2.1.5 Αλγόριθμοι Βελτιστοποίησης	20
2.1.6 Αλγόριθμος Οπισθοδιάδοσης	21
2.1.7 Τεχνικές Κανονικοποίησης (Regularization)	22
2.2 Επαναλαμβανόμενα Νευρωνικά Δίκτυα (RNNs)	24
2.2.1 Εισαγωγή.....	24
2.2.2 Αρχιτεκτονική	24
2.2.3 Επίπεδα.....	25
2.2.4 Εκπαίδευση	26
2.3 Δίκτυα Μακράς Βραχυπρόθεσμης Μνήμης (LSTM).....	27
2.3.1 Εισαγωγή.....	27
2.3.2 Αρχιτεκτονική	27
2.3.3 Πύλες.....	28
Κεφάλαιο 3: Ανακάλυψη Φαρμάκων (Drug Discovery).....	31
3.1 Εισαγωγή στην Ανακάλυψη Φαρμάκων	31
3.1.1 Βήματα της Ανακάλυψης Φαρμάκων	31
3.1.2 Παραδοσιακές Προσεγγίσεις	33
3.1.3 Πολυπλοκότητα	34
3.1.4 Χαρακτηριστικά και Ιδιότητες των Μορίων	35
3.1.5 Κανόνας των 5 του Lipinski (Lipinski's Rule of Five).....	38
3.2 Βαθιά Μάθηση και Drug Discovery	40
Κεφάλαιο 4: Πειραματικό Μέρος	42
4.1 Εισαγωγή.....	42

4.2 Σύντομη Περιγραφή	42
4.3 Βιβλιοθήκες	43
4.3.1 TensorFlow	43
4.3.2 Keras	43
4.3.3 RDKit	43
4.4 Δεδομένα	44
4.4.1 SMILES	44
4.4.2 Προεπεξεργασία Δεδομένων	46
4.4.3 Προετοιμασία Δεδομένων για Εισαγωγή στο Μοντέλο	47
4.5 Μοντέλο Βαθιάς Μάθησης	49
4.5.1 Δομή του Μοντέλου	49
4.5.2 Εκπαίδευση του Μοντέλου	52
4.6 Γενετική Παραγωγή Μορίων	54
4.6.1 Μέθοδος	54
4.6.2 Δημιουργία και Σχολιασμός Αποτελεσμάτων	56
4.7 Παραγωγή Μορίων από Αρχικό Τμήμα	64
4.8 Μεταφορά Μάθησης για Επικεντρωμένη Παραγωγή Μορίων	66
4.8.1 Βήματα της Διαδικασίας	66
4.8.2 Αποτελέσματα	67
Κεφάλαιο 5: Συμπεράσματα και Μελλοντική Έρευνα	71
5.1 Συμπεράσματα	71
5.2 Μελλοντική Έρευνα	72
References	73

Κεφάλαιο 1: Εισαγωγή

Η ανακάλυψη νέων φαρμάκων αποτελεί μία από τις πιο απαιτητικές και χρονοβόρες διαδικασίες στον τομέα της φαρμακευτικής βιομηχανίας, καθώς η εξερεύνηση του χημικού χώρου περιλαμβάνει την ανάλυση και αξιολόγηση δισεκατομμυρίων πιθανοτήτων μορίων. Ωστόσο, ο παραδοσιακός τρόπος ανακάλυψης φαρμάκων βασίζεται κυρίως σε ήδη υπάρχουσες βιβλιοθήκες ενώσεων, οι οποίες καλύπτουν μόνο ένα μικρό ποσοστό του συνολικού χημικού χώρου. Σε αυτό το πλαίσιο, οι σύγχρονες μέθοδοι τεχνητής νοημοσύνης και, ειδικότερα, τα γενετικά μοντέλα με επαναλαμβανόμενα νευρωνικά δίκτυα (Generative Recurrent Neural Networks) εισάγουν μια νέα προσέγγιση για τον εκ νέου σχεδιασμό μορίων.

Η συγκεκριμένη έρευνα επικεντρώνεται στη χρησιμοποίηση επαναλαμβανόμενων νευρωνικών δικτύων με χρήση Δικτύων Μακράς Βραχυπρόθεσμης Μνήμης (LSTM), για την ανάπτυξη μοντέλων που μπορούν να δημιουργήσουν νέες ενώσεις με τη μορφή SMILES ακολουθιών, οι οποίες αντιπροσωπεύουν τις δομές των μορίων. Η καινοτομία αυτής της προσέγγισης έγκειται στη δυνατότητα του μοντέλου να «μαθαίνει» τη σύνταξη των μοριακών αναπαραστάσεων και στη συνέχεια να δημιουργεί νέες δομές χωρίς να απαιτείται η ύπαρξη μιας εκτεταμένης βιβλιοθήκης ενώσεων.

Στην παρούσα εργασία, παρουσιάζεται η δυνατότητα του μοντέλου να προσαρμόζεται με χρήση μεταφοράς μάθησης (transfer learning), επιτρέποντας τη δημιουργία μορίων που είναι εστιασμένα σε συγκεκριμένους βιολογικούς στόχους. Αυτή η διαδικασία είναι ιδιαίτερα χρήσιμη στις πρώτες φάσεις ανακάλυψης φαρμάκων, όταν τα δεδομένα είναι περιορισμένα και απαιτείται αποτελεσματική διαχείριση των διαθέσιμων πόρων. Επίσης, εξετάζεται η δυνατότητα των γενετικών RNNs να υποστηρίξουν τον σχεδιασμό φαρμάκων βασισμένο σε τμήματα μορίων (fragment-based drug discovery), όπου γνωστά ενεργά μόρια επεκτείνονται για να δημιουργηθούν πιο σύνθετες ενώσεις.

Σκοπός της εργασίας είναι να παρουσιάσει την υλοποίηση και τη βελτιστοποίηση των γενετικών μοντέλων RNN στον τομέα του σχεδιασμού φαρμάκων, εστιάζοντας στις μεθόδους που καθιστούν αυτή την τεχνολογία πιο αποδοτική για τη δημιουργία νέων και πρωτοποριακών ενώσεων.

Κεφάλαιο 2: Θεωρητικό Υπόβαθρο

2.1 Βαθιά Μάθηση

2.1.1 Εισαγωγή

Η βαθιά μάθηση αποτελεί έναν από τους πιο σημαντικούς πυλώνες της τεχνητής νοημοσύνης, εστιάζοντας σε προηγμένα μοντέλα μάθησης που βασίζονται σε εκτενή σύνολα δεδομένων. Αυτά τα μοντέλα, γνωστά ως νευρωνικά δίκτυα, αποτελούνται από πολυάριθμα επίπεδα, τα οποία συνεργάζονται αρμονικά για να επεξεργάζονται και να αναλύουν σύνθετα μοτίβα δεδομένων με εντυπωσιακή αποτελεσματικότητα.

Ο όρος “βαθιά” αναφέρεται ακριβώς στον αριθμό των επιπέδων που περιλαμβάνονται στο δίκτυο, επιτρέποντας την εκτέλεση υψηλού επιπέδου υπολογισμών μέσω μη γραμμικών μετασχηματισμών. Κάθε επίπεδο του νευρωνικού δικτύου έχει τη δυνατότητα να επεξεργάζεται διαφορετικούς τύπους χαρακτηριστικών, ξεκινώντας από απλές δομές και φτάνοντας σε πιο περίπλοκες αναπαραστάσεις. Αυτή η ιεραρχία επιτρέπει στα δίκτυα να «μαθαίνουν» και να προσαρμόζονται σε ποικιλία δεδομένων, καθιστώντας τη βαθιά μάθηση εξαιρετικά ισχυρή με πολλές εφαρμογές.

Μερικοί από τους τομείς με τη μεγαλύτερη επίδραση, είναι η αναγνώριση φωνής, όπου έχει βοηθήσει σημαντικά στην αλληλεπίδραση ανθρώπου-μηχανής, ο χώρος των αυτόματων οχημάτων, όπου μέσω της κατανόησης του περιβάλλοντος χώρου επιτρέπει την ανίχνευση εμποδίων και βοηθάει στη λήψη άμεσων αποφάσεων, η ιατρική, αναλύοντας εικόνες και βοηθώντας στη διάγνωση και αναγνώριση ασθενειών, αλλά και η οικονομία, όπου αλγόριθμοι βαθιάς μάθησης χρησιμοποιούνται για την πρόβλεψη χρηματοοικονομικών τάσεων και τη διαχείριση κινδύνων.

2.1.2 Νευρωνικά Δίκτυα

Τα νευρωνικά δίκτυα αποτελούν τη θεμελιώδη μονάδα της βαθιάς μάθησης. Αποτελούνται από στρώματα τεχνητών νευρώνων, οι οποίοι μιμούνται τα βιολογικά νευρωνικά δίκτυα του ανθρώπινου εγκεφάλου [1]. Κάθε νευρώνας αυτού του δικτύου επεξεργάζεται εισερχόμενα δεδομένα και παρέχει έξοδο που μεταδίδεται στα επόμενα στρώματα. Συγκεκριμένα, λαμβάνει μία ή περισσότερες εισόδους, οι οποίες συνήθως έχουν διαφορετικά βάρη μεταξύ τους, και παράγει την έξοδο ως το άθροισμα του πολλαπλασιασμού της κάθε εισόδου με το βάρος της. Στη συνέχεια, η τιμή αυτή “περνάει” μέσα από τη συνάρτηση ενεργοποίησης και στην περίπτωση όπου είναι μεγαλύτερη από το προκαθορισμένο όριο, τότε η πληροφορία μεταδίδεται στον επόμενο νευρώνα.

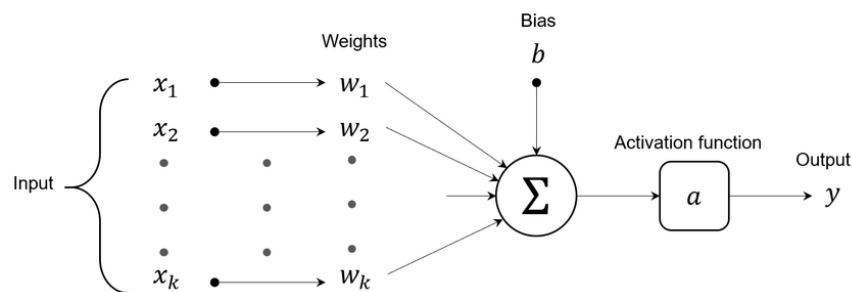


Figure 1: *Artificial Neuron*

Τα πολυστρωματικά νευρωνικά δίκτυα (Deep Neural Networks), περιλαμβάνουν πολλαπλά κρυφά στρώματα που επιτρέπουν τη μοντελοποίηση ιδιαίτερα σύνθετων σχέσεων μεταξύ των δεδομένων [2],[3]. Μέσω της χρήσης μη γραμμικών συναρτήσεων ενεργοποίησης, τα δίκτυα αυτά καταφέρνουν να μάθουν και να αποδώσουν με ακρίβεια ακόμη και σε περιπτώσεις όπου οι σχέσεις των δεδομένων δεν είναι άμεσα εμφανείς.

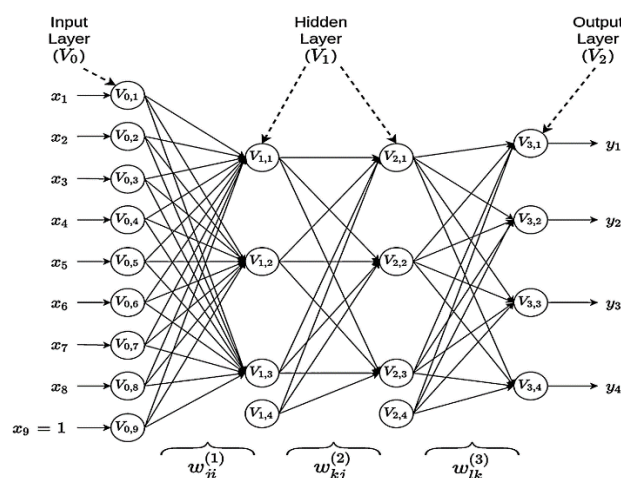


Figure 2: *Artificial Neural Network*

2.1.3 Συναρτήσεις Ενεργοποίησης

Ως συνάρτηση ενεργοποίησης μπορεί να οριστεί μία μαθηματική συνάρτηση που αποφασίζει αν ένας νευρώνας θα "ενεργοποιηθεί" ή όχι, με βάση την είσοδο που λαμβάνει [4]. Οι συναρτήσεις ενεργοποίησης εισάγουν μη γραμμικότητα στο δίκτυο, επιτρέποντας στο μοντέλο να μαθαίνει πιο σύνθετα πρότυπα, καθώς χωρίς αυτές, το δίκτυο θα λειτουργούσε σαν ένα απλό γραμμικό μοντέλο, περιορίζοντας την ικανότητά του να αποτυπώνει πολύπλοκες σχέσεις [5]. Ακόμα, καθορίζουν το εύρος των τιμών εξόδου, βοηθώντας στην κανονικοποίηση των δεδομένων καθώς αυτά περνούν μέσα από το δίκτυο. Παρακάτω παρουσιάζονται οι πιο κοινές συναρτήσεις ενεργοποίησης στη βαθιά μάθηση.

2.1.3.1 Σιγμοειδής Συνάρτηση (Sigmoid)

Η σιγμοειδής συνάρτηση είναι μία από τις πρώτες και πιο απλές συναρτήσεις ενεργοποίησης που χρησιμοποιήθηκαν. Είναι μία μαθηματική συνάρτηση όπου μετατρέπει τις εισόδους σε τιμές μεταξύ 0 και 1. Χρησιμοποιείται κυρίως σε περιπτώσεις δυαδικής ταξινόμησης, καθώς και σε περιπτώσεις όπου η τιμή εξόδου αναμένεται να έχει μορφή πιθανότητας. Ωστόσο, μπορεί να δημιουργήσει προβλήματα "vanishing gradient" σε περίπτωση χρησιμοποίησης αλγορίθμων οπισθοδιάδοσης (backpropagation), ελαχιστοποιώντας τις βαθμίδες, με αποτέλεσμα να μην ενημερώνονται τα βάρη των πρώτων στρωμάτων.

Ορίζεται στο σύνολο των πραγματικών αριθμών και έχει την παρακάτω μορφή:

$$\sigma(x) = \frac{1}{1 + e^{-x}} \quad (1)$$

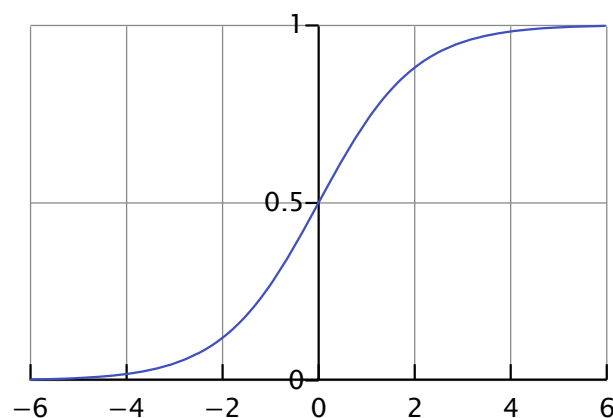


Figure 3: *Sigmoid Function*

2.1.3.2 Συνάρτηση Υπερβολικής Εφαπτομένης (Tanh)

Η συγκεκριμένη συνάρτηση μοιάζει με τη σιγμοειδή, καθώς έχει το ίδιο σχήμα “S”, ωστόσο παρουσιάζει κάποιες σημαντικές διαφορές. Το πεδίο τιμών κυμένεται μεταξύ -1 και 1, δίνοντας έτσι τη δυνατότητα στη συνάρτηση να κεντράρει τα δεδομένα γύρω από το μηδέν, ωστόσο για μεγάλες θετικές ή αρνητικές εισόδους, η βαθμίδα της πλησιάζει το μηδέν, προκαλώντας αντίστοιχα προβλήματα “vanishing gradient”.

Ορίζεται επίσης στο σύνολο των πραγματικών αριθμών και έχει την παρακάτω μορφή:

$$\tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}} \quad (2)$$

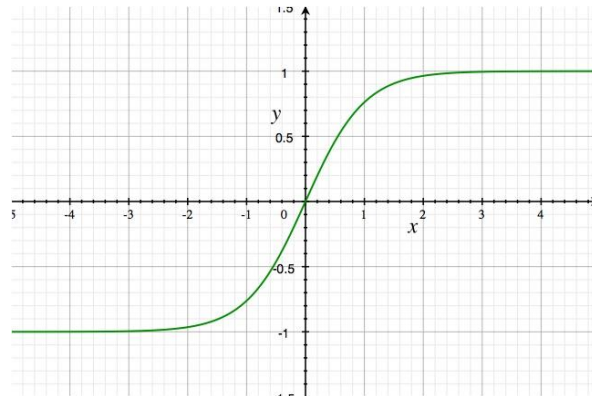


Figure 4: *tanh function*

2.1.3.3 Συνάρτηση Softmax

Η softmax συνάρτηση είναι μία μαθηματική συνάρτηση όπου μετατρέπει ένα διάνυσμα αριθμών σε πιθανότητες, το σύνολο των οποίων ισούται με 1. Χρησιμοποιείται στο τελευταίο επίπεδο των νευρωνικών δικτύων, για να παράγει προβλέψεις σε προβλήματα ταξινόμησης, και ιδιαίτερα σε ταξινομητές πολλαπλών κατηγοριών.

Η μορφή της, για ένα διάνυσμα $z = [z_1, z_2, \dots, z_n]$, είναι η εξής:

$$\text{softmax}(z_i) = \frac{e^{z_i}}{\sum_{j=1}^n e^{z_j}} \quad (3)$$

2.1.3.4 Συνάρτηση ReLU (Rectified Linear Unit)

Η συνάρτηση ReLU θεωρείται η πιο απλή από τις παραπάνω, καθώς επιστρέφει την τιμή της εισόδου, εάν αυτή είναι μεγαλύτερη του μηδενός, αλλιώς επιστρέφει 0. Η απλότητα αυτή, σε συνδυασμό με την υψηλή αποδοτικότητα στην επίλυση προβλημάτων όπως το vanishing gradient, την καθιστούν ιδανική επιλογή για τα περισσότερα πολυεπίπεδα μοντέλα. Παρά τα πλεονεκτήματά της, σε περιπτώσεις όπου οι τιμές εισόδου ενός νευρώνα είναι αρνητικές για μεγάλο χρονικό διάστημα, ο νευρώνας αυτός μπορεί να “νεκρωθεί”, σταματώντας έτσι τη μάθησή του. Το πρόβλημα αυτό προσπαθεί να λύσει μία παραλλαγή της ReLU συνάρτησης, η Leaky ReLU.

Περιγράφεται από την παρακάτω εξίσωση:

$$\text{ReLU}(x) = \max(0, x) \quad (4)$$

2.1.3.5 Συνάρτηση Leaky ReLU

Όπως αναφέρθηκε προηγουμένως, η συνάρτηση αυτή προσφέρει τη λύση στο πρόβλημα των νεκρών νευρώνων που μπορεί να δημιουργηθεί σε ορισμένες περιπτώσεις κατά τη χρήση της απλής ReLU συνάρτησης. Η διαφορά μεταξύ των δύο συναρτήσεων, είναι ότι στην παραλλαγή Leaky ReLU, υπάρχει μία μικρή κλίση για αρνητικές τιμές εισόδου, με αποτέλεσμα να μην σταματούν τελείως οι ενημερώσεις στα βάρη.

Ορίζεται ως εξής:

$$\text{LeakyReLU}(x) = \max(ax, x) \quad (5)$$

όπου a μικρή θετική σταθερά.

2.1.4 Συναρτήσεις Κόστους

Οι συναρτήσεις κόστους είναι ένα κρίσιμο μέρος της εκπαίδευσης των αλγορίθμων μηχανικής μάθησης και νευρωνικών δικτύων. Χρησιμοποιούνται για να μετρήσουν πόσο καλά αποδίδει ένα μοντέλο σε σχέση με την αναμενόμενη έξοδο (ground truth) και να κατευθύνουν τις ενημερώσεις των βαρών ώστε να μειωθεί η απόκλιση μεταξύ της πρόβλεψης και των πραγματικών τιμών. Κατά τη διάρκεια αυτής της διαδικασίας, η συνάρτηση συγκρίνει την πρόβλεψη του μοντέλου με την αληθινή τιμή και επιστρέφει ένα αριθμητικό σφάλμα. Ο στόχος της εκπαίδευσης του μοντέλου είναι να ελαχιστοποιήσει την τιμή αυτής της συνάρτησης, ενώ σε κάθε βήμα της εκπαίδευσης, ο αλγόριθμος χρησιμοποιεί τη συγκεκριμένη τιμή για να ενημερώσει τα βάρη του δικτύου.

Παρακάτω παρουσιάζονται μερικές από τις πιο διαδεδομένες συναρτήσεις κόστους.

2.1.4.1 Μέση Τετραγωνική Απόκλιση (Mean Squared Error)

Χρησιμοποιείται σε περιπτώσεις όπου η έξοδος είναι συνεχής τιμή, και υπολογίζει το μέσο τετραγωνικό σφάλμα μεταξύ της προβλεπόμενης τιμής $\hat{y}_i$ και της πραγματικής τιμής y_i [6]. Περιγράφεται από την εξής εξίσωση:

$$\text{MSE} = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (6)$$

2.1.4.2 Categorical Cross-Entropy

Η συνάρτηση αυτή χρησιμοποιείται για ταξινόμηση σε προβλήματα όπου η έξοδος ανήκει σε πολλές κατηγορίες, υπολογίζοντας τη διαφορά μεταξύ της προβλεπόμενης κατανομής πιθανοτήτων και της πραγματικής κατηγορίας [7]. Περιγράφεται από την παρακάτω συνάρτηση:

$$\text{CE} = - \sum_{i=1}^n y_i \log(\hat{y}_i) \quad (7)$$

2.1.5 Αλγόριθμοι Βελτιστοποίησης

Οι αλγόριθμοι βελτιστοποίησης καθορίζουν πώς τα βάρη του δικτύου ενημερώνονται προκειμένου να ελαχιστοποιηθεί η τιμή της συνάρτησης απώλειας [8],[9]. Αυτό συμβαίνει αναζητώντας τις τιμές των παραμέτρων, όπως τα βάρη του νευρωνικού δικτύου, οι οποίες ελαχιστοποιούν τη συνάρτηση κόστους. Δύο από τους πιο βασικούς αλγορίθμους βελτιστοποίησης, περιγράφονται παρακάτω.

2.1.5.1 Αλγόριθμος Gradient Descent

Είναι ο πιο απλός εκ των δύο, καθώς υπολογίζει την παράγωγο της συνάρτησης απώλειας σε σχέση με τα βάρη, και εν συνεχεία ενημερώνει τα βάρη προς την κατεύθυνση που μειώνει το σχετικό σφάλμα. Η εξίσωση που τον περιγράφει είναι η εξής:

$$w_{new} = w_{old} - \eta \cdot \frac{\partial L}{\partial w} \quad (8)$$

Όπου:

- w_{new} είναι τα ενημερωμένα βάρη.
- w_{old} είναι τα αρχικά βάρη.
- η είναι ο ρυθμός μάθησης (learning rate)
- $\frac{\partial L}{\partial w}$ είναι η παράγωγος της συνάρτησης κόστους L σε σχέση με τα βάρη.

2.1.5.2 Αλγόριθμος Adam (Adaptive Moment Estimation)

Είναι αρκετά πιο περίπλοκος από τον προηγούμενο, καθώς συνδυάζει ιδέες πολλαπλών αλγορίθμων βελτιστοποίησης. Συγκεκριμένα, χρησιμοποιεί το Momentum, όπου προσθέτει έναν όρο “ορμής” στην κατιούσα κλίση, βοηθώντας το μοντέλο να επιταχύνει τη σύγκλιση στις βέλτιστες τιμές, καθώς και το RMSprop (Root Mean Square Propagation), διατηρώντας έτσι έναν κινούμενο μέσο όρο των τετραγώνων κάθε βαθμίδας και προσαρμόζοντας τον ρυθμό μάθησης δυναμικά. Περιγράφεται ως εξής:

$$m_t = \beta_1 m_{t-1} + (1 - \beta_1) \cdot \frac{\partial L}{\partial w} \quad (9)$$

$$v_t = \beta_2 v_{t-1} + (1 - \beta_2) \cdot \left(\frac{\partial L}{\partial w} \right)^2 \quad (10)$$

$$\widehat{m}_t = \frac{m_t}{1 - \beta_1^t}, \quad \widehat{v}_t = \frac{v_t}{1 - \beta_2^t} \quad (11,12)$$

$$w_{new} = w_{old} - \frac{\eta}{\sqrt{\widehat{v}_t} + \epsilon} \cdot \widehat{m}_t \quad (13)$$

2.1.6 Αλγόριθμος Οπισθοδιάδοσης

Ο αλγόριθμος οπισθοδιάδοσης χρησιμοποιείται για την ενημέρωση των βαρών των συνδέσεων στα νευρωνικά δίκτυα, με βάση τα σφάλματα που προκύπτουν κατά τη διαδικασία σύγκρισης της εξόδου με την επιθυμητή έξοδο [2],[10].

Αρχικά, το μοντέλο λαμβάνει τα δεδομένα εισόδου, καθώς και τα αντίστοιχα βάρη τους, υπολογίζοντας έτσι μία έξοδο. Στη συνέχεια, χρησιμοποιώντας μία συνάρτηση κόστους, όπως η μέση τετραγωνική απόκλιση (Mean Squared Error), ή η categorical cross-entropy, υπολογίζεται το σχετικό σφάλμα. Έπειτα, το σφάλμα αυτό οπισθοδιαδίδεται από το επίπεδο εξόδου προς το επίπεδο εισόδου, υπολογίζοντας αρχικά τις παραγώγους του σφάλματος, και εν συνεχεία ενημερώνοντας τα βάρη του κάθε επιπέδου με βάση την επιλεγμένη συνάρτηση βελτιστοποίησης.

Για παράδειγμα, χρησιμοποιώντας έναν απλό αλγόριθμο βελτιστοποίησης, όπως αυτός της κατιούσας κλίσης (gradient descent), η συνάρτηση ενημέρωσης των βαρών είναι η εξής:

$$w_{new} = w_{old} - \eta \cdot \frac{\partial L}{\partial w} \quad (14)$$

Όπου:

- w_{new} είναι τα ενημερωμένα βάρη.
- w_{old} είναι τα αρχικά βάρη.
- η είναι ο ρυθμός μάθησης (learning rate), που καθορίζει το μέγεθος της βήματος.
- $\frac{\partial L}{\partial w}$ είναι η παράγωγος της συνάρτησης κόστους L σε σχέση με τα βάρη.

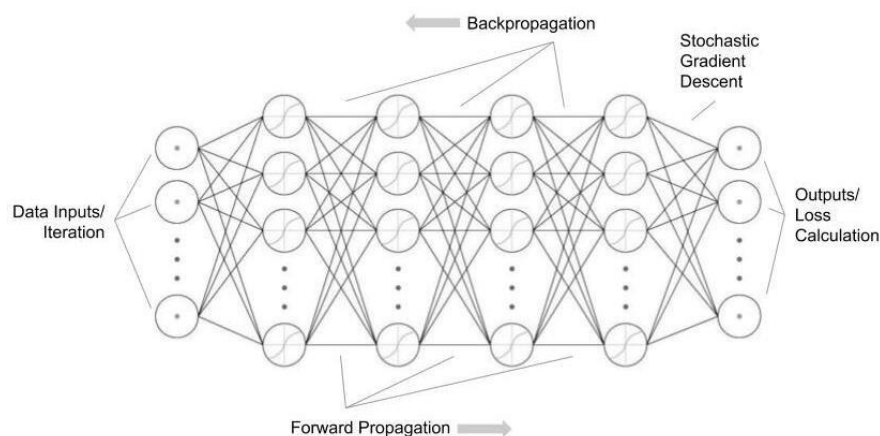


Figure 5: *Backpropagation*

2.1.7 Τεχνικές Κανονικοποίησης (Regularization)

Κατά τη διάρκεια της εκπαίδευσης του μοντέλου, μπορεί να δημιουργηθεί υπερμοντελοποίηση (overfitting), δηλαδή να “μάθει” υπερβολικά καλά τα δεδομένα εκπαίδευσης, αδυνατώντας να κάνει γενικεύσεις σε νέα δεδομένα. Για να αποφευχθεί το παραπάνω, υπάρχουν οι λεγόμενες τεχνικές κανονικοποίησης [11], μερικές από τις οποίες παρουσιάζονται στη συνέχεια.

2.1.7.1 Dropout

Η τεχνική αυτή λειτουργεί απενεργοποιώντας νευρώνες με τυχαίο τρόπο, κατά τη διάρκεια της εκπαίδευσης, με αποτέλεσμα το δίκτυο να μη βασίζεται υπερβολικά σε συγκεκριμένους νευρώνες [12]. Συγκεκριμένα, σε κάθε βήμα της εκπαίδευσης, ορισμένοι νευρώνες εξαιρούνται τυχαία από τους υπολογισμούς, θέτονται δηλαδή ίσοι με 0, με πιθανότητα p . Η συγκεκριμένη τεχνική ωστόσο, μπορεί να οδηγήσει σε υπομοντελοποίηση (underfitting), καθώς σε κάθε βήμα εκπαίδευσης υπάρχει απώλεια δεδομένων.

2.1.7.2 Κανονικοποίηση L2

Στη συγκεκριμένη τεχνική, ένας νέος όρος προστίθεται στην αρχική συνάρτηση κόστους, με σκοπό την αποφυγή της υπερβολικής αύξησης των βαρών. Με αυτόν τον τρόπο, το μοντέλο δε βασίζεται υπερβολικά σε συγκεκριμένα χαρακτηριστικά, με συνέπεια να είναι πιο γενικεύσιμο σε νέα δεδομένα [13]. Ο όρος αυτός αποτελείται από έναν συντελεστή κανονικοποίησης λ και το τετράγωνο των βαρών, με τη συνολική εξίσωση να περιγράφεται παρακάτω:

$$L(w) = L_{original}(w) + \lambda \cdot \sum_{i=1}^n w_i^2 \quad (15)$$

2.1.7.3 Πρώιμη Διακοπή (Early Stopping)

Κατά τη διάρκεια της εκπαίδευσης, παρακολουθείται η απόδοση του μοντέλου, πέρα από το σύνολο εκπαίδευσης, και σε ένα σύνολο επικύρωσης. Εάν η απόδοση σε αυτό το σύνολο αρχίσει να χειροτερεύει, ή έστω σταματήσει να βελτιώνεται, η εκπαίδευση διακόπτεται [14].

2.1.7.4 Εμπλουτισμός Δεδομένων (Data Augmentation)

Ο στόχος της συγκεκριμένης τεχνικής είναι να ενισχύσει το σύνολο των δεδομένων με νέα δεδομένα, τα οποία παραλλαγές των αρχικών [15]. Με αυτόν τον τρόπο, το μοντέλο εκτίθεται σε μεγαλύτερη ποικιλία δεδομένων και μαθαίνει να γενικεύει καλύτερα.

2.2 Επαναλαμβανόμενα Νευρωνικά Δίκτυα (RNNs)

2.2.1 Εισαγωγή

Τα Επαναλαμβανόμενα Νευρωνικά Δίκτυα (Recurrent Neural Networks - RNNs) είναι μια εξελιγμένη κατηγορία νευρωνικών δικτύων, ειδικά σχεδιασμένη για την επεξεργασία ακολουθιακών δεδομένων [17]. Σε αντίθεση με τα παραδοσιακά νευρωνικά δίκτυα, τα RNNs ενσωματώνουν μηχανισμούς μνήμης, επιτρέποντάς τους να διατηρούν και να αξιοποιούν πληροφορίες από προηγούμενες εισόδους [16]. Αυτή τους η ικανότητα τα καθιστά ιδιαίτερα χρήσιμα σε εφαρμογές όπου η αλληλουχία των δεδομένων είναι κρίσιμη. Η “εσωτερική μνήμη” που διαθέτουν τα RNNs, τους επιτρέπει να διατηρούν πληροφορίες από προηγούμενα βήματα της ακολουθίας, επιτρέποντας την πρόβλεψη ακολουθιών [16], ενώ η συνεχής ανατροφοδότηση δίνει τη δυνατότητα χρησιμοποίησης της εξόδου ενός νευρώνα μία χρονική στιγμή, ως είσοδο, είτε στον εαυτό του, είτε σε διαφορετικούς νευρώνες, την επόμενη χρονική στιγμή. Αυτό ορίζεται ως διάνυσμα “κρυφής” κατάστασης (hidden layer) και περιέχει όλες τις προηγούμενες εισόδους στο δίκτυο.

2.2.2 Αρχιτεκτονική

Η βασική δομή ενός RNN μοντέλου περιλαμβάνει έναν βρόχο όπου κάθε κατάσταση εξαρτάται άμεσα από την προηγούμενη κατάσταση [16],[17]. Η εξίσωση ενημέρωσης για ένα RNN είναι η εξής:

$$h_t = f(W_x x_t + W_h h_{t-1} + b) \quad (16)$$

Όπου:

- h_t είναι η κρυφή κατάσταση (hidden state) στο χρονικό βήμα t .
- x_t είναι το δεδομένο εισόδου στο χρονικό βήμα t .
- W_x είναι τα βάρη που συνδέουν την είσοδο με την κρυφή κατάσταση.
- W_h είναι τα βάρη που συνδέουν την προηγούμενη κρυφή κατάσταση με την τρέχουσα.
- b είναι το διάνυσμα μεροληψίας (bias).
- f είναι μια μη γραμμική συνάρτηση ενεργοποίησης.

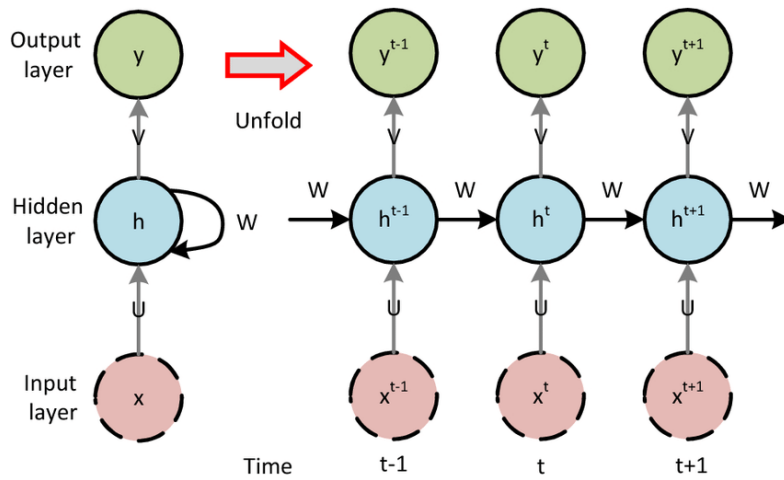


Figure 6: *Recurrent Neural Network*

2.2.3 Επίπεδα

2.2.3.1 Επίπεδο εισόδου (Input Layer)

Το επίπεδο εισόδου είναι υπεύθυνο για τη λήψη της αρχικής εισόδου και τη μετατροπή της σε μορφή που μπορεί να επεξεργαστεί το δίκτυο. Η είσοδος μπορεί να είναι οποιοσδήποτε τύπος ακολουθιακών δεδομένων, όπως λέξεις, χρονοσειρές ή δεδομένα ομιλίας. Ανάλογα με το πρόβλημα, η είσοδος μπορεί να είναι μονοδιάστατη ή πολυδιάστατη, ενώ συχνά τα δεδομένα εισόδου χρειάζεται να επεξεργαστούν ώστε να μετατραπούν σε διανύσματα αναπαράστασης.

Στα RNNs, η είσοδος συνήθως παρουσιάζεται ως μια ακολουθία $\{x_1, x_2, \dots, x_T\}$, όπου κάθε x_t αντιπροσωπεύει την είσοδο στο χρονικό βήμα t , ενώ T είναι το συνολικό μήκος της ακολουθίας.

2.2.3.2 Κρυφό Επίπεδο (Hidden Layer)

Είναι το κεντρικό στοιχείο των RNNs, και είναι το σημείο όπου λαμβάνει χώρα η διατήρηση της "μνήμης" και η εξαρτησιακή μάθηση μεταξύ των χρονικών βημάτων. Αυτό το επίπεδο επιτρέπει την ανατροφοδότηση πληροφοριών από προηγούμενα χρονικά βήματα σε επόμενα χρονικά βήματα, κάτι που είναι το κλειδί για την κατανόηση ακολουθιακών δεδομένων.

Σε κάθε χρονικό βήμα t , το επαναλαμβανόμενο επίπεδο διατηρεί μια κρυφή κατάσταση ht , η οποία περιέχει πληροφορίες για την είσοδο μέχρι εκείνο το χρονικό σημείο. Αυτή η κρυφή κατάσταση ενημερώνεται σε κάθε βήμα και εξαρτάται από την τρέχουσα είσοδο και την κρυφή κατάσταση του προηγούμενου χρονικού βήματος h_{t-1} .

Κάθε νευρώνας στο επαναλαμβανόμενο επίπεδο συνδέεται με την τρέχουσα είσοδο μέσω ενός συνόλου βαρών Wx και με την προηγούμενη κρυφή κατάσταση μέσω άλλου συνόλου βαρών Wh , όπως περιγράφεται και στην εξίσωση (16).

2.2.3.3 Επίπεδο Εξόδου (Output Layer)

Είναι υπεύθυνο για τη μετατροπή της κρυφής κατάστασης σε μια έξοδο κατάλληλη για το πρόβλημα που επιλύουμε. Αυτό το επίπεδο λαμβάνει την κρυφή κατάσταση από το τελευταίο επαναλαμβανόμενο επίπεδο και παράγει μια πρόβλεψη, είτε για κάθε χρονικό βήμα, σε προβλήματα ακολουθιακής εξόδου, είτε για ολόκληρη την ακολουθία.

2.2.4 Εκπαίδευση

Η εκπαίδευση των RNNs γίνεται μέσω της τεχνικής Οπισθοδιάδοσης Μέσα στον Χρόνο (Backpropagation Through Time) [18]. Αυτή είναι μια επέκταση του κλασικού αλγορίθμου οπισθοδιάδοσης, προσαρμοσμένη ώστε να λαμβάνει υπόψη τη χρονική αλληλουχία των δεδομένων. Έτσι, ενώ όπως στα κλασικά νευρωνικά δίκτυα οι παράγωγοι της συνάρτησης κόστους υπολογίζονται σε σχέση με κάθε παράμετρο του δικτύου, στα RNNs αυτό γίνεται σε κάθε χρονικό βήμα ξεχωριστά, και οι παράγωγοι αρθροποιούνται για το σύνολο των βημάτων. Ωστόσο, σε μεγάλα RNNs, ή σε δεδομένα μεγάλων ακολουθιών, η εκπαίδευση μπορεί να παρουσιάσει προβλήματα όπως το vanishing gradient, καθώς σε περιπτώσεις όπου οι παράγωγοι ελαχιστοποιούνται σε μεγάλο βαθμό, τα αρχικά βήματα της εκπαίδευσης καταλήγουν να μην επηρεάζουν επαρκώς την εκπαίδευση.

2.3 Δίκτυα Μακράς Βραχυπρόθεσμης Μνήμης (LSTM)

2.3.1 Εισαγωγή

Το πρόβλημα των vanishing gradients μπορεί να λυθεί με χρήση LSTM δικτύων [16],[19], καθώς είναι σχεδιασμένα με βάση αυτό το πρόβλημα. Όπως αναφέρθηκε προηγουμένως, τα παραδοσιακά RNNs μπορούν να αποθηκεύσουν πληροφορίες από προηγούμενα χρονικά βήματα, ωστόσο για μεγάλα χρονικά διαστήματα η ικανότητα αυτή μειώνεται σημαντικά. Αυτό οδηγεί στο παραπάνω πρόβλημα, όπου οι παράγωγοι που χρησιμοποιούνται για την ενημέρωση των βαρών ελαχιστοποιούνται, με αποτέλεσμα οι πρώιμες χρονικές στιγμές να μην επιδρούν σημαντικά στην εκπαίδευση του δικτύου.

Τα LSTM δίκτυα σχεδιάστηκαν με σκοπό να λύσουν το συγκεκριμένο πρόβλημα, δίνοντας τη δυνατότητα στο δίκτυο να διατηρεί την πληροφορία για μεγαλύτερα χρονικά διαστήματα. Αυτό επιτυγχάνεται με χρήση πυλών και κελιών μνήμης (memory cells), τα οποία ρυθμίζουν τη ροή της πληροφορίας.

2.3.2 Αρχιτεκτονική

Τα δίκτυα LSTM αποτελούνται από μια σειρά επαναλαμβανόμενων μονάδων, καθεμία από τις οποίες αποτελείται από μια κατάσταση κελιού (cell state) και ένα σύνολο πυλών που ρυθμίζουν τη ροή των πληροφοριών. Συγκεκριμένα, ένα LSTM δίκτυο έχει τρεις πύλες, υπεύθυνες για τη διαχείριση του περιεχομένου της μνήμης. Οι πύλες αυτές είναι η πύλη εισόδου (input gate), η πύλη εξόδου (output gate) και η forget gate [19],[20]. Από την πύλη εξόδου, παράγεται το διάνυσμα hidden state, το οποίο περιέχει την πληροφορία για την τρέχουσα κατάσταση του μοντέλου, με βάση τα δεδομένα της ακολουθίας που έχουν επεξεργαστεί μέχρι τη συγκεκριμένη χρονική στιγμή, ενώ οι άλλες δύο πύλες διαμορφώνουν το κελί κατάστασης (cell state), το οποίο αντιπροσωπεύει τη μακροπρόθεσμη μνήμη. Με αυτόν τον τρόπο, το μοντέλο μπορεί να διατηρεί τη μνήμη για μεγάλα χρονικά διαστήματα, λύνοντας το πρόβλημα που δημιουργείται σε αντίστοιχα μεγέθη στα τυπικά RNNs δίκτυα [16],[19].

2.3.3 Πύλες

Παρακάτω παρουσιάζονται αναλυτικά οι 3 πύλες, από τις οποίες αποτελούνται ένα δίκτυο μακράς βραχυπρόθεσμης μνήμης.

2.3.3.1 Forget Gate

Η συγκεκριμένη πύλη αποφασίζει ποια πληροφορία από το προηγούμενη κελί μνήμης πρέπει να απορριφθεί και ποια πρέπει να διατηρηθεί. Εισάγει την ιδέα της επιλεκτικής μνήμης, επιτρέποντας στο δίκτυο να παρακάμπτει περιττές πληροφορίες που μπορεί να μην είναι πλέον χρήσιμες για τη μοντελοποίηση της ακολουθίας. Περιγράφεται από την παρακάτω εξίσωση:

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (17)$$

Όπου:

- f_t περιέχει την “απόφαση” για το ποιες πληροφορίες θα αποθηκευτούν.
- σ είναι η σιγμοειδής συνάρτηση ενεργοποίησης.
- W_f είναι τα βάρη που σχετίζονται με τη συγκεκριμένη πύλη.
- $[h_{t-1}, x_t]$ είναι η συνένωση της προηγούμενης κρυφής κατάστασης h_{t-1} και της τρέχουσας εισόδου x_t .
- b_f είναι η μεροληψία (bias).

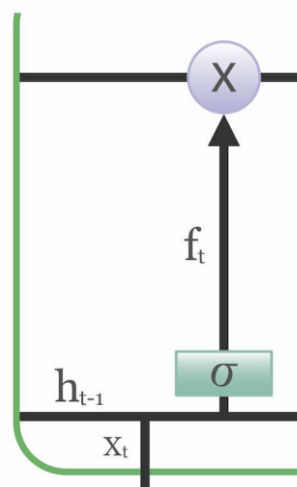


Figure 7: Forget Gate

2.3.3.2 Πύλη Εισόδου (Input Gate)

Η πύλη εισόδου καθορίζει ποιες νέες πληροφορίες θα αποθηκευτούν στο κελί μνήμης. Συνδυάζεται με την κρυφή κατάσταση και την τρέχουσα είσοδο για να παράγει ένα ενημερωμένο κελί μνήμης. Η λειτουργία της περιγράφεται από τις παρακάτω εξισώσεις:

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (18)$$

$$\tilde{C}_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) \quad (19)$$

Όπου:

- i_t είναι η έξοδος της πύλης εισόδου, η οποία καθορίζει ποιες πληροφορίες θα αποθηκευτούν στη μνήμη.
- $\tilde{C}_t$ είναι οι νέες πληροφορίες που μπορούν να προστεθούν στην κυψέλη μνήμης.
- W_i, W_c είναι τα βάρη που μαθαίνονται.
- $\tanh$ χρησιμοποιείται για να περιορίσει τις τιμές μεταξύ -1 και 1.

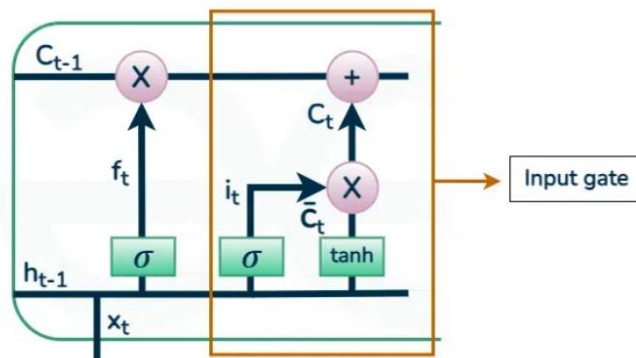


Figure 8: LSTM Input Gate

Στη συνέχεια, ακολουθεί η ανανέωση της κατάστασης του κελιού μνήμης από την προηγούμενη κατάσταση C_t . Αυτή η ενημέρωση γίνεται μέσω της εξίσωσης:

$$C_t = f_t \cdot C_{t-1} + i_t \cdot \tilde{C}_t \quad (20)$$

Όπου:

- C_t είναι το ενημερωμένο κελί μνήμης.
- C_{t-1} είναι το κελί μνήμης από το προηγούμενο χρονικό βήμα.

2.3.3.3 Πύλη Εξόδου (Output Gate)

Η πύλη εξόδου ελέγχει ποια πληροφορία από την κυψέλη μνήμης θα χρησιμοποιηθεί για να παράγει την έξοδο στο τρέχον χρονικό βήμα. Αυτή η πύλη επιτρέπει τη μεταφορά της πληροφορίας από την κυψέλη μνήμης στην κρυφή κατάσταση. Περιγράφεται στις παρακάτω εξισώσεις:

$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (21)$$

$$h_t = o_t \cdot \tanh(C_t) \quad (22)$$

Όπου:

- o_t είναι η έξοδος της πύλης εξόδου, η οποία καθορίζει ποιες πληροφορίες θα περάσουν στην κρυφή κατάσταση.
- h_t είναι η νέα κρυφή κατάσταση.
- C_t είναι το ενημερωμένο κελί μνήμης.

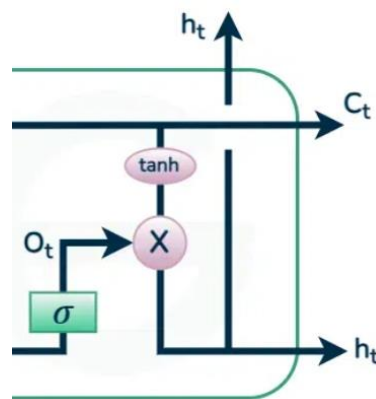


Figure 9: Output Gate

Στο σύνολό του, ένα δίκτυο μακράς βραχυπρόθεσμης μνήμης έχει την παρακάτω μορφή:

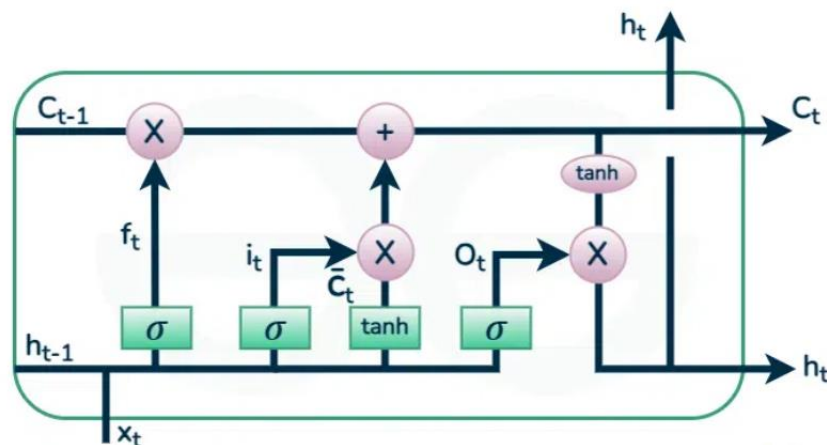


Figure 10: LSTM neural network

Κεφάλαιο 3: Ανακάλυψη Φαρμάκων (Drug Discovery)

3.1 Εισαγωγή στην Ανακάλυψη Φαρμάκων

Η ανακάλυψη φαρμάκων είναι μια πολυδιάστατη διαδικασία, κατά την οποία νέες φαρμακευτικές ουσίες ανακαλύπτονται και αναπτύσσονται για την καταπολέμηση ασθενειών [21]. Περιλαμβάνει τον εντοπισμό πιθανών νέων φαρμάκων, την ανάπτυξή τους και, τελικά, την εισαγωγή τους στην αγορά, μια πορεία γεμάτη πολυπλοκότητες, μακροχρόνιες διαδικασίες και σημαντικά χρηματοοικονομικά κόστη.



Figure 11: *Drug Discovery Pipeline*

3.1.1 Βήματα της Ανακάλυψης Φαρμάκων

Η διαδικασία για την ανακάλυψη ενός νέου φαρμάκου περιλαμβάνει μία σειρά επιστημονικών βημάτων [21], όπου μετατρέπουν μία ιδέα σε λειτουργικό φάρμακο. Παρακάτω, περιγράφονται τα συγκεκριμένα βήματα.

1. Αναγνώριση Στόχων (Target Identification):

Είναι το πρώτο βήμα και σχετίζεται με τον εντοπισμό του βιολογικού στόχου για την εκάστοτε ασθένεια. Οι στόχοι αυτοί είναι είτε πρωτεΐνες, είτε ένζυμα, είτε γονίδια που επηρεάζουν την πορεία της ασθένειας.

2. Αναγνώριση Υποψήφιων Μορίων (Lead identification):

Στο συγκεκριμένο βήμα, γίνεται αναζήτηση για πιθανά μόρια, τα οποία μπορούν να επηρεάσουν τον προκαθορισμένο στόχο. Αυτό γίνεται με τις παρακάτω 2 κύριες προσεγγίσεις:

- a. Εμπειρικός Έλεγχος (Empirical Screening): Ο πιο “απλός”, αλλά ταυρόχρονα και χρονοβόρος τρόπος, όπου χρησιμοποιούνται χημικές βιβλιοθήκες με εκατομμύρια ενώσεις, και μετά από συνεχείς δοκιμές σε βιολογικά συστήματα, επιλέγονται οι ενώσεις που αλληλεπιδρούν με τον στόχο [22].
- b. Εικονική Προβολή (Virtual Screening): Ο πιο καινοτόμος και αποδοτικός τρόπος, ο οποίος είναι και το αντικείμενο μελέτης της παρούσας διπλωματικής, κατά τον οποίο χρησιμοποιούνται υπολογιστικές τεχνικές για να μειωθεί ο χημικός χώρος έρευνας, και να προβλεφθεί πολύ πιο αποτελεσματικά ποιές ενώσεις μπορούν να αλληλεπιδράσουν με τον στόχο [23].

3. Βελτιστοποίηση Υποψηφίων Μορίων (Lead Optimization):

Έχοντας καταλήξει σε έναν αριθμό υποψήφιων μορίων, τα οποία αλληλεπιδρούν θετικά με τον στόχο, ακολουθεί η βελτιστοποίησή τους σε χημικές ιδιότητες όπως η δραστικότητα, η σταθερότητα και η διαλυτότητα.

4. Προκλινική Ανάπτυξη (Preclinical Development):

Σε αυτό το στάδιο, τα υποψήφια φάρμακα που έχουν περάσει τη φάση της βελτιστοποίησης, υποβάλλονται σε προκλινικές δοκιμές για να αξιολογηθεί η ασφάλεια, η τοξικότητα, και η φαρμακοκινητική, το πώς δηλαδή το φάρμακο μετακινείται μέσα στο σώμα στην πορεία του χρόνου.

5. Κλινικές Δοκιμές (Clinical Trials):

Εάν το προηγούμενο βήμα είναι επιτυχές, ακολουθούν οι κλινικές δοκιμές, οι οποίες γίνονται σε ανθρώπους και αποτελούνται από τρεις φάσεις:

- a. Αξιολόγηση της ασφάλειας σε μικρό αριθμό υγιών εθελοντών.
- b. Δοκιμή της αποτελεσματικότητας σε μικρό αριθμό ασθενών.
- c. Δοκιμή της αποτελεσματικότητας σε μεγαλύτερο αριθμό ασθενών.

6. Έγκριση και Κυκλοφορία στην Αγορά:

Τελευταίο βήμα σε αυτή τη διαδικασία αποτελεί η έγκριση του φαρμάκου από τις κατάλληλες ρυθμιστικές αρχές.

3.1.2 Παραδοσιακές Προσεγγίσεις

Ιστορικά, η ανακάλυψη φαρμάκων βασίστηκε σε παραδοσιακές προσεγγίσεις που διαμόρφωσαν το τοπίο της φαρμακευτικής ανάπτυξης [24],[25],[26]. Αυτές οι μέθοδοι περιλαμβάνουν:

3.1.2.1 Εμπειρικός Έλεγχος (Empirical Screening)

Όπως αναφέρθηκε και στην προηγούμενη ενότητα, η συγκεκριμένη προσέγγιση περιλαμβάνει τη δοκιμή χιλιάδων ενώσεων για την ικανότητά τους να παράγουν το επιθυμητό βιολογικό αποτέλεσμα. Απαιτεί συχνά μεγάλες βιβλιοθήκες χημικών ενώσεων και εκτεταμένη εργαστηριακή εργασία για τον εντοπισμό πιθανών υποψηφίων φαρμάκων, γεγονός που την καθιστά χρονοβόρα και δαπανηρή [23]. Ωστόσο, θεωρείται μία αρκετά αποτελεσματική μέθοδος.

3.1.2.2 Σχεδιασμός Φαρμάκων Βασισμένος στη Δομή (Structure-Based Drug Design)

Ο SBDD χρησιμοποιεί τις τρισδιάστατες δομές βιολογικών μορίων για τον σχεδιασμό φαρμάκων [24]. Κατανοώντας το σχήμα και τη λειτουργία των στόχων, οι ερευνητές μπορούν να σχεδιάσουν ενώσεις που να ταιριάζουν ακριβώς, όπως ένα κλειδί σε μια κλειδαριά. Αυτή η μέθοδος έχει οδηγήσει στην ανάπτυξη πολλών επιτυχημένων φαρμάκων, αλλά απαιτεί προηγμένες τεχνολογίες και εξειδίκευση.

3.1.2.3 Ανακάλυψη Φαρμάκων με Βάση Στόχους (Target Based Drug Discovery)

Σε αυτή την προσέγγιση, η προσοχή επικεντρώνεται σε συγκεκριμένους βιολογικούς στόχους που σχετίζονται με ασθένειες [25],[26]. Οι ερευνητές προσπαθούν να κατανοήσουν τους μηχανισμούς δράσης και πώς οι ενώσεις μπορούν να τροποποιήσουν αυτούς τους στόχους. Αυτή η προσέγγιση είναι ιδιαίτερα επιστημονική, αλλά μπορεί να περιορίζεται από την ελλιπή κατανόηση της βιολογίας της ασθένειας.

3.1.3 Πολυπλοκότητα

Η πολυπλοκότητα της ανακάλυψης νέων φαρμάκων είναι μία από τις μεγαλύτερες προκλήσεις στη φαρμακευτική βιομηχανία. Εμπλέκει διάφορους επιστημονικούς τομείς, όπως η χημεία, η βιολογία, η φαρμακολογία και η κλινική έρευνα [21]. Κάθε στάδιο της διαδικασίας ανακάλυψης φαρμάκων παρουσιάζει μοναδικά εμπόδια.

3.1.3.1 Εντοπισμός Κατάλληλων Στόχων

Ο εντοπισμός του σωστού βιολογικού στόχου είναι αρκετά κρίσιμος, καθώς πολλές ασθένειες περιλαμβάνουν περίπλοκες βιολογικές διεργασίες και πολλαπλούς στόχους, καθιστώντας δύσκολο τον εντοπισμό του πιο αποτελεσματικού. Επιπλέον, οι στόχοι πρέπει να επικυρωθούν για να διασφαλιστεί ότι παίζουν κρίσιμο ρόλο στην πορεία της ασθένειας.

3.1.3.2 Έλεγχος και Βελτιστοποίηση Ενώσεων

Αφού εντοπιστούν πιθανές ενώσεις, υποβάλλονται σε εκτεταμένο έλεγχο για να αξιολογηθεί η αποτελεσματικότητά και η ασφάλειά τους. Αυτή η φάση μπορεί να περιλαμβάνει χιλιάδες ενώσεις, και η βελτιστοποίησή τους για καλύτερη απόδοση μπορεί να απαιτεί πολλαπλούς γύρους δοκιμών και τροποποιήσεων.

3.1.3.3 Προκλινικές και Κλινικές Δοκιμές

Μετά τον εντοπισμό μιας κύριας ένωσης, οι ερευνητές πρέπει να διεξάγουν προκλινικές μελέτες για να αξιολογήσουν την ασφάλεια και την αποτελεσματικότητά της σε ανθρώπους και ζώα. Εάν το βήμα αυτό κριθεί επιτυχημένο, η ένωση προχωρά στις κλινικές δοκιμές, οι οποίες μπορεί να διαρκέσουν αρκετά χρόνια για να ολοκληρωθούν. Οι κλινικές δοκιμές, όπως συζητήθηκε προηγουμένως, χωρίζονται σε φάσεις, καθεμία με συγκεκριμένους στόχους, και μπορεί να εμπλέκουν χιλιάδες συμμετέχοντες.

3.1.4 Χαρακτηριστικά και Ιδιότητες των Μορίων

Σημαντικό βήμα στη διαδικασία ανακάλυψης νέων φαρμάκων αποτελεί η συστηματική αξιολόγηση συγκεκριμένων μοριακών ιδιοτήτων [31]-[33], οι οποίες σχετίζονται με τη βιολογική δραστηριότητα των μορίων και καθορίζουν την πιθανότητα ενός υποψήφιου μορίου να επιτύχει τους στόχους του. Παρακάτω, περιγράφονται μερικά από τα πιο σημαντικά φυσικοχημικά χαρακτηριστικά των μορίων, τα οποία αξιολογούνται κατά την διάρκεια της προαναφερθείσας διαδικασίας.

3.1.4.1 Μοριακό Βάρος (Molecular Weight)

Ως μοριακό βάρος ενός μορίου, ορίζεται το σύνολο της μάζας όλων των ατόμων που το απαρτίζουν και ως μονάδα μέτρησης έχει το Dalton (Da). Θεωρείται από τις πλέον σημαντικές ιδιότητες που αξιολογούνται για την διαπίστωση της βιοδιαθεσιμότητας και της απορροφησιμότητας ενός φαρμάκου, ενώ επίσης αποτελεί και μία από τις συνθήκες του κανόνα των 5 του Lipinski (Rule of Five) [31], ο οποίος θα αναλυθεί στη συνέχεια.

Σύμφωνα λοιπόν με τον παραπάνω κανόνα, το ιδανικό μοριακό βάρος ενός μορίου πρέπει είναι μικρότερο των 500 Da. Συγκεκριμένα, σε περιπτώσεις πολύ υψηλού μοριακού βάρους, τα μόρια δυσκολεύονται να διαπεράσουν τις κυτταρικές μεμβράνες, γεγονός που οδηγεί σε μειωμένη απορρόφησή τους από το σώμα. Ωστόσο, ούτε οι αρκετά χαμηλές τιμές μοριακού βάρους είναι επιθυμητές, καθώς σε αυτές τις περιπτώσεις ενώ μπορεί να επιτυγχάνεται υψηλή απορροφησιμότητα, η σταθερότητα του μορίου μειώνεται δραματικά.

3.1.4.2 Λιποφιλικότητα (LogP)

Εκφράζει τη συγγένεια ενός μορίου ως προς ένα λιπόφιλο περιβάλλον, την ικανότητα δηλαδή, να διαλύεται σε λιπαρά μέσα και να διεισδύει σε κυτταρικές μεμβράνες. Είναι επίσης μέρος του κανόνα των 5 του Lipinski, καθώς έχει σημαντικό ρόλο στην κατανόηση της ισορροπίας μεταξύ διαλυτότητας και διαπερατότητας.

Κατάλληλο θεωρείται ένα μόριο, του οποίου ο λογαριθμικός συντελεστής κατανομής (logP) έχει τιμές μεταξύ 0 και 3, δίνοντάς του τη δυνατότητα να διεισδύει επιτυχώς στις κυτταρικές μεμβράνες, χωρίς όμως να χάνει την υδροδιαλυτότητά του. Μεγαλύτερες τιμές θα δημιουργούσαν συσσώρευση του μορίου σε λιπώδεις ιστούς, μειώνοντας έτσι τη διαλυτότητα του φαρμάκου. Από την άλλη πλευρά, γίνεται αντιληπτό πως χαμηλότερες τιμές θα δημιουργούσαν πρόβλημα σχετικά με την απορροφησιμότητα, καθώς το μόριο θα δυσκολευόταν να διεισδύσει στις κυτταρικές μεμβράνες.

3.1.4.3 Αριθμός Δοτών Υδρογόνου (Hydrogen Bond Donors)

Αναφέρεται στον αριθμό των ομάδων σε ένα μόριο, οι οποίες μπορούν να λειτουργήσουν ως δότες υδρογονικών δεσμών (όπως -OH, -NH). Είναι εξίσου σημαντική ιδιότητα με τις προηγούμενες, καθώς συμβάλλει στον καθορισμό της διαλυτότητας και της απορροφησιμότητας των μορίων.

Συγκεκριμένα, οι υδρογονικοί δεσμοί διαμορφώνουν τη συγγένεια των μορίων με τους στόχους τους, επηρεάζοντας άμεσα τις αλληλεπιδράσεις με τους υποδοχείς, τα ένζυμα και άλλες πρωτεΐνες. Έτσι, ένας υψηλός αριθμός δοτών υδρογόνου μπορεί να οδηγήσει στην αύξηση της διαλυτότητας του μορίου, μειώνοντας ωστόσο την ικανότητά του να διεισδύει στις κυτταρικές μεμβράνες. Είναι μέρος του κανόνα του Lipinski, βάσει του οποίου ο μέγιστος αριθμός δοτών υδρογόνου που πρέπει να έχει ένα μόριο, είναι 5.

3.1.4.4 Αριθμός Αποδεκτών Υδρογόνου (Hydrogen Bond Acceptors)

Αντίστοιχα, αναφέρεται στον αριθμό των ομάδων σε ένα μόριο, οι οποίες μπορούν να λειτουργήσουν ως αποδέκτες υδρογονικών δεσμών (όπως -N, -C=O). Έχουν παρόμοια συμβολή με τους δότες, καθώς επίσης διαμορφώνουν τη συγγένεια των φαρμάκων με τους στόχους τους και επηρεάζουν τη διαλυτότητα των μορίων.

Ένας υπερβολικά μεγάλος αριθμός αποδεκτών υδρογόνου οδηγεί στην αύξηση των μη ομοιοπολικών δεσμών του μορίου με τους στόχους, καθιστώντας το αρκετά πολικό και μειώνοντας έτσι την ικανότητά του να διαπερνά τις κυτταρικές μεμβράνες. Για αυτόν το λόγο, σύμφωνα και με τον κανόνα των 5 του Lipinski, ένα μόριο θα πρέπει να έχει το πολύ 10 αποδέκτες υδρογονικών δεσμών, έτσι ώστε να έχει υψηλή βιοδιαθεσιμότητα.

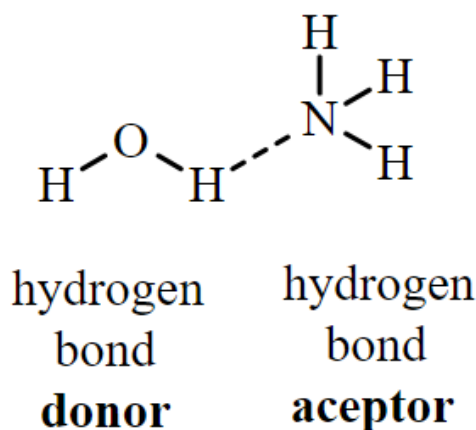


Figure 12: *Hydrogen Bond Donors and Acceptors*

3.1.4.5 QED Score (Quantitative Estimate of Drug-likeness Score)

Ο δείκτης QED περιγράφει μία συνολική εικόνα της “φαρμακευτικότητας” του φαρμάκου, λαμβάνοντας υπόψιν πολλαπλές φυσικοχημικές ιδιότητες, μεταξύ των οποίων βρίσκονται το μοριακό βάρος, η λιποφιλικότητα, καθώς και οι δότες/αποδέκτες υδρογονικών δεσμών [33]. Οι τιμές του κυμαίνονται στο διάστημα μεταξύ 0 και 1, με την “φαρμακευτικότητα” του φαρμάκου να αυξάνεται όσο η τιμή πλησιάζει στο 1, το οποίο θεωρείται και η καλύτερη τιμή αυτού του δείκτη. Έτσι, υψηλές τιμές QED σε ένα μόριο υποδηλώνουν χαμηλή τοξικότητα, ιδανική διαπερατότητα και κυρίως κατάλληλη βιοδιαθεσιμότητα, αυξάνοντας τις πιθανότητες επιτυχίας του ως φάρμακο.

Βασίζεται στις συναρτήσεις επιθυμίας (desirability functions), που πρωτοαναφέρθηκαν από τον Harrington E.C. στο άρθρο του “*The Desirability Function*” το 1965 και στη συνέχεια εξελίχθηκε από τους Derringer και Suich το 1983 [34]. Οι συναρτήσεις αυτές χρησιμοποιούνται για τη μετατροπή πολυάριθμων παραμέτρων, οι οποίες μετρούνται σε διαφορετικές κλίμακες, σε μία ενιαία κλίμακα μεταξύ 0 και 1, με σκοπό τον υπολογισμό συνολικών δεικτών, όπως είναι και ο QED.

Περιγράφεται από την παρακάτω εξίσωση:

$$QED = \exp\left(\frac{1}{n} \sum_{i=1}^n \ln(d_i)\right) \quad (1)$$

Όπου:

- d_i είναι η συνάρτηση επιθυμίας για κάθε μοριακό δείκτη (όπως το μοριακό βάρος και το LogP).
- n είναι ο αριθμός των μοριακών δεικτών.

3.1.5 Κανόνας των 5 του Lipinski (Lipinski's Rule of Five)

Όπως φάνηκε και από τις συνεχείς αναφορές στον συγκεκριμένο κανόνα στις προηγούμενες παραγράφους, ο κανόνας των 5 του Lipinski [31] αποτελεί έναν από τους πιο διαδεδομένους εμπειρικούς κανόνες στη φαρμακευτική για την αξιολόγηση της “φαρμακευτικότητας” ενός μορίου, καθορίζοντας τη βιοδιαθεσιμότητα του μορίου κατά τη διά στόματος χορήγησή του. Δημιουργήθηκε από τον Christopher Lipinski και την ομάδα του το 1997, και ονομάστηκε έτσι καθώς οι αριθμοί που χρησιμοποιούνται στις συνθήκες του είναι πολλαπλάσια του 5.

Όπως αναφέρθηκε προηγουμένως, περιέχει τέσσερις συνθήκες, από τις οποίες θα πρέπει να παραβιάζεται το πολύ μία, έτσι ώστε το μόριο να έχει υψηλή πιθανότητα επιτυχίας ως φάρμακο. Αυτό συμβαίνει διότι βασίστηκε στην ανάλυση των ιδιοτήτων ήδη γνωστών επιτυχών φαρμάκων. Έτσι, ενώ υπήρχε η αντίληψη πως κατά βάση ένα τέτοιο φάρμακο θα έπρεπε να πληροί και τις τέσσερις αυτές συνθήκες, διαπιστώθηκε πως εάν παραβιαζόταν μόλις μία από αυτές, το φάρμακο θα μπορούσε ακόμα να έχει υψηλή “φαρμακευτικότητα” και απορροφιστικότητα.

Οι τέσσερις αυτές συνθήκες, οι οποίες αναφέρθηκαν περιληπτικά και προηγουμένως, είναι οι εξής:

- Μοριακό βάρος: Το μόριο θα πρέπει να έχει μοριακό βάρος μικρότερο από 500 Dalton (Da). Μόρια με υψηλότερο μοριακό βάρος δυσκολεύονται να απορροφηθούν και να διαπεράσουν τις κυτταρικές μεμβράνες.
- Λιποφιλικότητα: Το LogP θα πρέπει να είναι μικρότερο από 5, εξασφαλίζοντας έτσι την ισορροπία μεταξύ λιποδιαλυτότητας και υδατοδιαλυτότητας.
- Αριθμός Δοτών Υδρογόνου: Το μόριο θα πρέπει να έχει 5 ή λιγότερους δότες υδρογονικών δεσμών, έτσι ώστε να διατηρείται η διαπερατότητα της κυτταρικής μεμβράνης.
- Αριθμός Αποδεκτών Υδρογόνου: Το μόριο θα πρέπει να έχει 10 ή λιγότερους αποδέκτες υδρογονικών δεσμών, συμβάλλοντας έτσι στην υψηλή βιοδιαθεσιμότητα του μορίου.

Μπορεί να δίνει μία γενική εικόνα για τη βιοδιαθεσιμότητα των δια στόματος χορηγούμενων φαρμάκων, καθώς ένα μόριο που πληροί τον κανόνα των 5 είναι πολύ πιο πιθανό να απορροφηθεί από τον άνθρωπο, να φτάσει επιτυχώς στον στόχο του και να έχει φαρμακολογική δράση, ωστόσο δεν είναι απόλυτος. Υπάρχουν και περιπτώσεις φαρμάκων, όπως για παράδειγμα τα αντιβιοτικά που έχουν υψηλότερο μοριακό βάρος, τα οποία παραμένουν εξαιρετικά αποτελεσματικά.

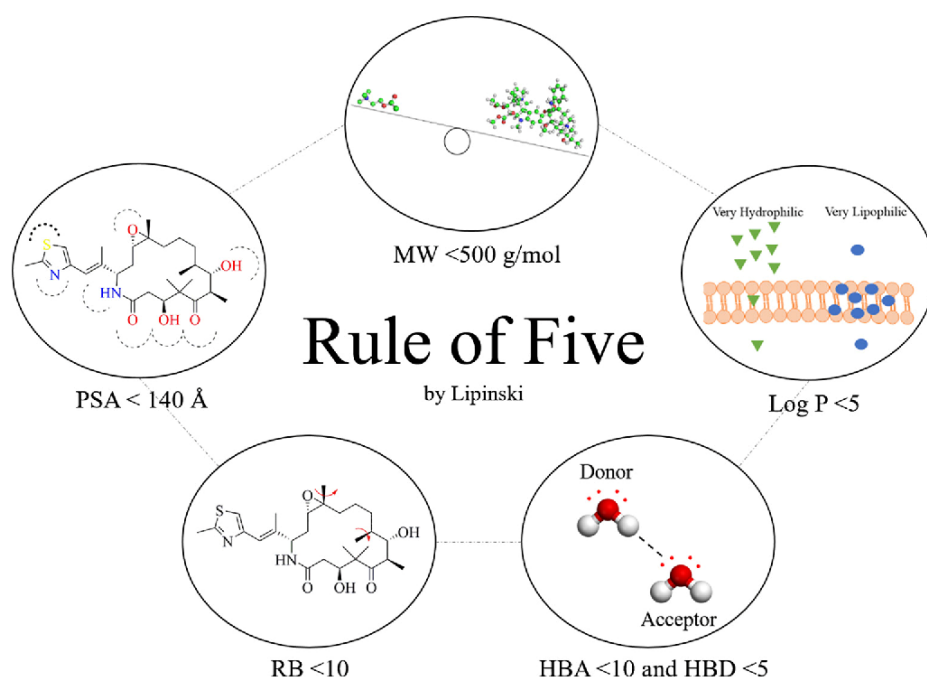


Figure 13: *Lipinski's Rule of Five*

3.2 Βαθιά Μάθηση και Drug Discovery

Η ανακάλυψη νέων φαρμάκων είναι μια από τις πιο πολύπλοκες και πολυδιάστατες διαδικασίες στη βιοϊατρική έρευνα, που απαιτεί τεράστια ποσότητα δεδομένων, επαναληπτικά πειράματα, και συχνά, πολλά χρόνια για να επιτευχθεί. Η παραδοσιακή προσέγγιση περιλαμβάνει χρονοβόρες διαδικασίες, όπως τον εμπειρικό έλεγχο χιλιάδων χημικών ενώσεων και τη συστηματική βελτιστοποίησή τους, μέχρι να βρεθεί μια ένωση που να είναι κατάλληλη για κλινικές δοκιμές [21]. Η βαθιά μάθηση έχει εισέλθει στον τομέα αυτό [27]-[30], προσφέροντας νέες δυνατότητες και εργαλεία που μπορούν να μειώσουν τα κόστη, να αυξήσουν την αποδοτικότητα και να επιταχύνουν τη διαδικασία ανακάλυψης φαρμάκων.

Έχοντας αρχικά εφαρμοστεί σε τομείς όπως η επεξεργασία εικόνας και η αναγνώριση φωνής, έχει βρει νέα εφαρμογή στην ανακάλυψη φαρμάκων, όπου η ικανότητά της να διαχειρίζεται τεράστιους όγκους δεδομένων και να αναγνωρίζει πολύπλοκα μοτίβα έχει τεράστια αξία. Τα δεδομένα που χρησιμοποιούνται σε αυτή τη διαδικασία περιλαμβάνουν χημικές ιδιότητες, βιολογικές πληροφορίες, προκλινικές και κλινικές μελέτες, καθώς και γενετικά δεδομένα, καθιστώντας τη απαραίτητο εργαλείο για να αντλήσει χρήσιμες πληροφορίες από δεδομένα με τέτοια πολυπλοκότητα.

Η είσοδος της βαθιάς μάθησης στην ανακάλυψη φαρμάκων συνδέεται στενά με την αύξηση της υπολογιστικής ισχύος και τη διαθεσιμότητα μεγάλων βιοϊατρικών συνόλων δεδομένων. Στο παρελθόν, τα δεδομένα ήταν αποσπασματικά, δύσκολα στη συλλογή και επεξεργασία, και η ανάλυσή τους απαιτούσε παραδοσιακές μεθόδους που συχνά ήταν αναποτελεσματικές σε μεγάλα σύνολα δεδομένων. Η βαθιά μάθηση, χάρη στην ικανότητά της να μαθαίνει από δεδομένα χωρίς να απαιτείται να προκαθοριστούν συγκεκριμένοι κανόνες, εισήλθε ως μια τεχνολογία που μπορούσε να εξάγει μοτίβα και προβλέψεις από μεγάλα και περίπλοκα δεδομένα. Αυτή η δυνατότητα της βαθιάς μάθησης να εξάγει γρήγορα και αξιόπιστα προβλέψεις έκανε την τεχνολογία αυτή αναπόσπαστο εργαλείο στην ανακάλυψη φαρμάκων.

Ένα από τα κύρια προβλήματα που αντιμετωπίζει η φαρμακευτική βιομηχανία είναι η ανάγκη να εξεταστούν χιλιάδες, αν όχι εκατομμύρια, χημικές ενώσεις, ώστε να εντοπιστούν οι πιο υποσχόμενες για περαιτέρω ανάπτυξη. Αυτή η διαδικασία, η οποία αναλύθηκε στην προηγούμενη ενότητα, απαιτεί παραδοσιακά πολύ χρόνο και κόπο. Με τη βοήθεια της βαθιάς μάθησης όμως, μπορεί να επιταχυνθεί σημαντικά. Τα νευρωνικά δίκτυα μπορούν να εκπαιδευτούν με τεράστια σύνολα δεδομένων που περιλαμβάνουν τις χημικές ιδιότητες διαφόρων ενώσεων και τις βιολογικές τους επιδράσεις, και στη συνέχεια, να προβλέψουν ποιες ενώσεις είναι πιο πιθανό να έχουν τα επιθυμητά θεραπευτικά αποτελέσματα, χωρίς να χρειάζεται να περάσουν όλες από δοκιμές σε εργαστήριο.

Πέρα όμως από την αντιμετώπιση του παραπάνω προβλήματος, η χρήση βαθιάς μάθησης συμβάλει θετικά και στο πρόβλημα της βελτιστοποίησης φαρμάκων. Όπως αναφέρθηκε προηγουμένως, αφού εντοπιστεί μια υποψήφια χημική ένωση, χρειάζεται να βελτιστοποιηθεί για να αυξηθεί η αποτελεσματικότητά της και να μειωθούν οι πιθανές παρενέργειες. Αυτό είναι ένα πολύπλοκο πρόβλημα που περιλαμβάνει την τροποποίηση της δομής των μορίων και την πρόβλεψη του τρόπου, με τον οποίο οι αλλαγές αυτές θα επηρεάσουν την αποτελεσματικότητα και την ασφάλεια του φαρμάκου. Με τη χρησιμοποίηση μοντέλων βαθιάς μάθησης, μπορούν να δημιουργηθούν μοντέλα που προβλέπουν την αποτελεσματικότητα των μοριακών αλλαγών [27], επιτρέποντας την πιο γρήγορη και ακριβή βελτιστοποίηση των φαρμάκων.

Ένα άλλο σημαντικό πεδίο όπου η βαθιά μάθηση παίζει κρίσιμο ρόλο είναι η ανακάλυψη νέων στόχων για φάρμακα (drug targets) [27],[30]. Η επιλογή του σωστού βιολογικού στόχου, όπως μια πρωτεΐνη ή ένα ένζυμο που σχετίζεται με μια ασθένεια, είναι ζωτικής σημασίας για την επιτυχία ενός φαρμάκου. Στο παρελθόν, η επιλογή των στόχων αυτών βασιζόταν σε χρόνια έρευνας και γνώσεων σχετικά με τις βιολογικές διεργασίες. Με αυτή την προσέγγιση, μπορούν να αναλυθούν τεράστια βιοϊατρικά δεδομένα, όπως γονιδιακές εκφράσεις και πρωτεϊνικές αλληλεπιδράσεις, και να εντοπιστούν νέοι στόχοι που θα μπορούσαν να οδηγήσουν σε καινοτόμες θεραπείες. Αυτή η δυνατότητα όχι μόνο επιταχύνει την ανακάλυψη νέων στόχων, αλλά και ανοίγει τον δρόμο για πιο εξατομικευμένες θεραπείες που προσαρμόζονται στα συγκεκριμένα βιολογικά χαρακτηριστικά κάθε ασθενούς.

Ένας ακόμα τομέας που επωφελείται από την εισαγωγή της βαθιάς μάθησης στον τομέα της φαρμακευτικής είναι η πρόβλεψη της τοξικότητας και των παρενεργειών νέων φαρμάκων. Ένα από τα μεγαλύτερα προβλήματα στη διαδικασία ανάπτυξης φαρμάκων είναι ότι πολλά υποψήφια φάρμακα αποτυγχάνουν στις κλινικές δοκιμές λόγω τοξικότητας ή ανεπιθύμητων παρενεργειών. Η βαθιά μάθηση μπορεί να χρησιμοποιηθεί για την ανάλυση βιολογικών και κλινικών δεδομένων, προβλέποντας ποια φάρμακα είναι πιο πιθανό να προκαλέσουν ανεπιθύμητες αντιδράσεις σε συγκεκριμένα βιολογικά συστήματα. Αυτό δίνει τη δυνατότητα για σημαντική μείωση στα ποσοστά αποτυχίας στις κλινικές δοκιμές, κάνοντας τη διαδικασία ανάπτυξης φαρμάκων πιο αποδοτική και λιγότερο δαπανηρή.

Παρά όμως την πρόοδο που έχει σημειώσει στην ανακάλυψη φαρμάκων, υπάρχουν ακόμα προκλήσεις που πρέπει να ξεπεραστούν. Ένα από τα μεγαλύτερα προβλήματα είναι η ποιότητα και η διαθεσιμότητα των δεδομένων. Αν και τα νευρωνικά δίκτυα μπορούν να μάθουν από τεράστιες ποσότητες δεδομένων, η ακρίβεια των προβλέψεών τους εξαρτάται από την ποιότητα των δεδομένων που χρησιμοποιούνται για την εκπαίδευσή τους. Τα δεδομένα της φαρμακευτικής έρευνας είναι συχνά ελλιπή ή ανακριβή, γεγονός που μπορεί να επηρεάσει αρνητικά τις προβλέψεις της βαθιάς μάθησης.

Κεφάλαιο 4: Πειραματικό Μέρος

4.1 Εισαγωγή

Η υλοποίηση του συγκεκριμένου πειράματος βασίζεται στην έρευνα και την καινοτόμο προσέγγιση που αναπτύσσουν σχετικά με τις νέες μεθόδους ανακάλυψης φαρμάκων, οι A. Gupta, A. Müller, B. Huisman, J. Fuchs, P. και G. Schneider, στο άρθρο τους με τίτλο “Generative Recurrent Networks for De Novo Drug Design” [27]. Στόχος του πειράματος είναι η δημιουργία νέων, φαρμακολογικά ενεργών μορίων, με χρήση ενός RNN μοντέλου με LSTM κελιά, καθώς και η εξέταση πιο στοχευμένων σεναρίων, όπως η εκπαίδευση του μοντέλου με εφαρμογή τεχνικής μάθησης σε συγκεκριμένα σύνολα δεδομένων-στόχων με σκοπό τη δημιουργία μορίων, δομικά παρόμοιων με ήδη υπάρχουσες βιοδραστικές ενώσεις, αλλά και η διαδικασία δημιουργίας φαρμάκων με βάση κάποιο υπάρχον τμήμα του τελικού στόχου.

4.2 Σύντομη Περιγραφή

Τα βήματα που ακολουθήθηκαν στη διάρκεια αυτού του πειράματος περιγράφονται περιληπτικά παρακάτω:

1. Αρχικά, ως σύνολο δεδομένων επιλέχθηκαν 500.000 χημικές ενώσεις σε μορφή SMILES συμβολοσειρών από τη βάση δεδομένων ChEMBL (<https://www.ebi.ac.uk/chembl>), οι οποίες μετά από την απαραίτητη προεπεξεργασία αποθηκεύτηκαν σε νέο αρχείο.
2. Μετά την ολοκλήρωση του πρώτου σταδίου επεξεργασίας, το οποίο βασίστηκε σε καθαρά χημικές αρχές, τα δεδομένα του νέου ανανεωμένου αρχείου προετοιμάστηκαν καταλλήλως, έτσι ώστε να ταιριάζουν στις απαιτήσεις του μοντέλου.
3. Στη συνέχεια, ορίστηκε η αρχιτεκτονική του μοντέλου, με βάση τις βέλτιστες τεχνικές υλοποίησης και παραμετροποίησης LSTM μοντέλων, καθώς και ένα σύνολο συναρτήσεων που χρησιμοποιήθηκαν στη συνέχεια.
4. Μετά την εκπαίδευση του μοντέλου, αξιολογήθηκε η απόδοσή του με κριτήριο τις ιδιότητες των παραγόμενων ενώσεων, όπως αυτές περιγράφονται στην ενότητα 3.1.4.
5. Έπειτα, εξετάστηκε η ικανότητα του μοντέλου να δημιουργεί ενώσεις, έχοντας ως βάση ένα συγκεκριμένο τμήμα, συγκρίνοντας τις ιδιότητές τους.
6. Τέλος, μέσω μεταφοράς μάθησης, το μοντέλο εκπαιδεύτηκε σε συγκεκριμένα δεδομένα-στόχους, και εξετάστηκε η δυνατότητά του να προσαρμόζεται σε αυτά και να παράγει δομικά παρόμοια μόρια.

4.3 Βιβλιοθήκες

Για την υλοποίηση του πειράματος χρησιμοποιήθηκε η γλώσσα προγραμματισμού python, ενώ ως περιβάλλον εκτέλεσης επιλέχθηκε το google colab. Ο λόγος για τον οποίο προτιμήθηκε αυτή η γλώσσα είναι οι πολύ ισχυρές και πλήρεις βιβλιοθήκες που διαθέτει, ειδικά σε σχέση με την μηχανική μάθηση, αλλά και σχετικά με την επεξεργασία και οπτικοποίηση δεδομένων χημικής πληροφορίας. Έτσι, για τον ορισμό και την εκπαίδευση του LSTM μοντέλου χρησιμοποιήθηκαν η TensorFlow και το Keras, ενώ για την ανάλυση, επεξεργασία και αναπαράσταση των μορίων χρησιμοποιήθηκε η RDKit. Πέρα από τις παραπάνω, χρησιμοποιήθηκαν και βιβλιοθήκες γενικής χρήσης, όπως numpy, pandas, sklearn, matplotlib και άλλες.

4.3.1 TensorFlow

Η TensorFlow θεωρείται μία από τις πιο διαδεδομένες βιβλιοθήκες για κατασκευή, ορισμό και εκπαίδευση μοντέλων μηχανικής μάθησης [35]. Είναι πλατφόρμα ανοιχτού κώδικα και επιτρέπει τον χειρισμό μεγάλων ποσοτήτων δεδομένων, παρέχοντας δυνατότητες όπως GPU/TPU υποστήριξη για ταχύτερη εκπαίδευση, καθώς και εργαλεία παρακολούθησης και βελτιστοποίησης της εκπαίδευσης, όπως callbacks και checkpoints.

4.3.2 Keras

Το Keras είναι ένα υψηλού επιπέδου API, το οποίο χρησιμοποιείται για την κατασκευή και ανάπτυξη μοντέλων βαθιάς μάθησης. Αποτελεί μέρος της TensorFlow βιβλιοθήκης και ξεχωρίζει για την απλότητα και ευελιξία που προσδίδει στη δημιουργία νευρωνικών δικτύων. Προτερήματα όπως αυτά, το καθιστούν ως ένα από τα πιο εύχρηστα εργαλεία στον χώρο της μηχανικής μάθησης, γεγονός που αποδεικνύεται από την ευρεία χρήση του από την επιστημονική κοινότητα, καθώς και από μερικά από τα πιο γνωστά ερευνητικά κέντρα και ινστιτούτα, όπως το CERN, η NASA και το NIH.

4.3.3 RDKit

Η συγκεκριμένη βιβλιοθήκη διαφέρει από τις παραπάνω, καθώς σχετίζεται με την ανάλυση, επεξεργασία και αναπαράσταση χημικής πληροφορίας. Είναι επίσης λογισμικό ανοιχτού κώδικα και χρησιμοποιείται ευρέως από την επιστημονική κοινότητα για την ανακάλυψη φαρμάκων. Μέσω της δυνατότητας που παρέχει για μετατροπή μεταξύ διαφόρων μορφών χημικών δεδομένων, επιτρέπει την ανάλυση μορίων μέσω μηχανικής μάθησης.

4.4 Δεδομένα

Για τις ανάγκες του πειράματος, χρησιμοποιήθηκε η ανοιχτής πρόσβασης βάση δεδομένων βιοδραστικών μορίων ChEMBL (<https://www.ebi.ac.uk/chembl>). Αρχικά, επιλέχθηκαν 500.000 μόρια σε SMILES μορφή, έχοντας ως κριτήριο επιλογής τον κανόνα των 5 του Lipinski [31]. Στη συνέχεια, τα μόρια αυτά πέρασαν από μία διαδικασία “καθαρισμού” τους, σκοπός της οποίας είναι η βελτίωση της αξιοπιστίας και συνέπειας των μορίων. Τέλος, τα από χημικής πλευράς επεξεργασμένα μόρια, έπειτα από περεταίρω επεξεργασία, μετατράπηκαν σε αναγνώσιμες για το μοντέλο ακολουθίες.

Παρακάτω, αναλύεται η μορφή SMILES που έχει επιλεγεί για την αναπαράσταση των μορίων, καθώς και τα 2 στάδια επεξεργασίας τους.

4.4.1 SMILES

Η SMILES (Simplified Molecular Input Line Entry System) συμβολοσειρά είναι μία σειριακή αναπαράσταση μοριακών δομών ως ακολουθία χαρακτήρων, που παρέχει έναν τυποποιημένο τρόπο γραφής τους, φιλικό προς επεξεργασία από υπολογιστικά συστήματα [36]. Είναι ευρέως διαδεδομένη στον χώρο της ανακάλυψης φαρμάκων μέσω υπολογιστικών συστημάτων, καθώς είναι εύκολα ερμηνεύσιμη από μοντέλα μηχανικής μάθησης. Θεωρείται πολύ πιο εύχρηστη και συμπαγής σε σχέση με τρισδιάστατες μορφές αναπαράστασης, μιας και ο μονοδιάστατος και γραμμικός τρόπος γραφής της την καθιστά σχετικά απλή και εύκολη προς χρήση. Τέλος, αποτελεί ιδανική επιλογή για μοντέλα που λειτουργούν κυρίως με αλληλουχίες δεδομένων, όπως το LSTM δίκτυο που έχει χρησιμοποιηθεί και στη συγκεκριμένη υλοποίηση.

Αποτελείται από μία ακολουθία συμβόλων, τα οποία εκφράζουν όλες τις χημικές πληροφορίες του μορίου. Συγκεκριμένα:

- Άτομα: Τα στοιχεία που απαρτίζουν το μόριο αντιπροσωπεύονται από το σύμβολό τους στον Περιοδικό Πίνακα.
- Δεσμοί: Οι απλοί δεσμοί παραλείπονται, οι διπλοί συμβολίζονται με “=”, ενώ οι τριπλοί με “#”.
- Διακλαδώσεις: Για την αναπαράσταση διακλαδώσεων μορίων χρησιμοποιούνται παρενθέσεις.
- Κύκλοι: Κατά την διαδικασία μετατροπής του μορίου σε μορφή SMILES, στα άκρα κάθε κύκλου του μορίου προστίθεται ένας αριθμός, που υποδηλώνει ότι αυτά τα άκρα ενώνονται.

CN1C=CC2=C(C=C1)C(=O)C(=O)C2

The diagram shows a central benzene ring with several substituents. The ring carbons are labeled C1, C2, C3, C4, C5, and C6. Substituents include an oxygen atom (O) at C1, a carbonyl group (C=O) at C2, a nitrogen atom (N) at C3, and a methyl group (CH3) at C4. Dashed lines with arrows extend from the left and right sides of the structure.

O=C2C(O)C=CC3C2(C4)C5C1C(O)CCC5CC3N(C)C4

Figure 14: *Evolution of SMILES representation of Morphine*

4.4.2 Προεπεξεργασία Δεδομένων

Η διαδικασία αυτή, αποτελείται από τα παρακάτω τέσσερα βήματα, τα οποία πραγματοποιούνται για κάθε μόριο ξεχωριστά:

1. Χημική τυποποίηση του μορίου: Σε αυτό το βήμα, μέσω της μεθόδου `MolStandardize.normalize()` της RDKit βιβλιοθήκης, το μόριο τυποποιείται διορθώνοντας τυχόν συντακτικά λάθη ή ασυνήθιστες δομές, διασφαλίζοντας έτσι μία ομοιόμορφη δομή για την περαιτέρω επεξεργασία.
2. Αφαίρεση μη βιολογικά ενεργών τμημάτων: Στη συνέχεια, μέσω της μεθόδου `MolStandardize.LargestFragmentChooser()`, αφαιρούνται ουσίες που δεν συμβάλλουν στη βιολογική δραστηριότητα του μορίου, όπως άλατα ή άλλες ουσίες σταθεροποίησης ή αύξησης της διαλυτότητας του μορίου. Με αυτόν τον τρόπο, δίνεται η δυνατότητα στο μοντέλο να εστιάσει στο κύριο βιολογικά ενεργό τμήμα του μορίου, μειώνοντας έτσι την πολυπλοκότητά του και βελτιώνοντας την ακρίβεια των αποτελεσμάτων του.
3. Εξουδετέρωση: Όπως το αρχικό βήμα, και αυτό αποσκοπεί στην τυποποίηση του μορίου, αφαιρώντας τυχόν φορτία που μπορεί να έχει, μέσω της μεθόδου `MolStandardize.Uncharger()`. Αυτό είναι απαραίτητο, καθώς πιθανά φορτία μπορούν να προσθέσουν αστάθεια στο μοντέλο, ενώ η απουσία αυτών δημιουργεί μία ουδέτερη και ομοιόμορφη μορφή στο σύνολο των δεδομένων.
4. Κανονικοποίηση SMILES: Στο τελευταίο αυτό βήμα, πραγματοποιείται η κανονικοποίηση του μορίου, κατά την οποία η πληροφορία σχετικά με την τρισδιάστατη διάταξη των ατόμων αφαιρείται, δημιουργώντας έτσι μία μοναδική αναπαράσταση κάθε μορίου. Με αυτόν τον τρόπο, τα δεδομένα αποκτούν συνέπεια και διευκολύνεται η σύγκριση μεταξύ των μορίων.



Figure 15: *Preprocessing Steps*

4.4.3 Προετοιμασία Δεδομένων για Εισαγωγή στο Μοντέλο

Έπειτα από την κανονικοποίηση και τον “καθαρισμό” των μορίων, λαμβάνοντας υπόψιν χημικές και φαρμακευτικές αρχές, τα δεδομένα χρειάζεται να επεξεργαστούν έτσι ώστε να μετατραπούν σε αναγνώσιμη μορφή, συμβατή με το μοντέλο βαθιάς μάθησης. Η διαδικασία αυτή αποτελείται από τα τρία παρακάτω βήματα:

1. Κατακερματισμός της SMILES συμβολοσειράς: Στην πρώτη φάση της κωδικοποίησης, η κάθε συμβολοσειρά SMILES διαχωρίζεται σε ξεχωριστούς χαρακτήρες. Οι χαρακτήρες αυτοί μπορεί να είναι είτε άτομα, είτε ειδικά σύμβολα και αριθμοί που αναπαράσταν δεσμούς, κύκλους και διακλαδώσεις. Έτσι, μετά την ολοκλήρωση της διαδικασίας διαχωρισμού, τα μεμονωμένα πλέον σύμβολα αποθηκεύονται σε μία λίστα.

Για παράδειγμα, στην περίπτωση όπου η SMILES συμβολοσειρά εισόδου ήταν η “C(=O)”, μετά τη διαδικασία κατακερματισμού της το τελικό αποτέλεσμα θα ήταν η λίστα “[‘C’, ‘(’, ‘=’, ‘O’, ‘)’]”.

2. Συμπλήρωση Ακολουθίας (Padding): Στη συνέχεια, η ακολουθία που έχει παραχθεί από το πρώτο βήμα συμπληρώνεται με ειδικούς χαρακτήρες, με σκοπό την ομοιομορφία των δεδομένων. Συγκεκριμένα, στην αρχή της ακολουθίας προστίθεται ο χαρακτήρας “G”, όπου υποδηλώνει την έναρξή της, ενώ στο τέλος της ακολουθίας προστίθεται ο χαρακτήρας “E” για να δηλώσει την ολοκλήρωσή της. Ωστόσο, τα μοντέλα LSTM, όπως τα περισσότερα νευρωνικά δίκτυα, απαιτούν οι είσοδοι να έχουν ίσο μήκος σε κάθε batch. Συνεπώς, για να επιτευχθεί αυτό, μετά τον χαρακτήρα “E” οι ακολουθίες συμπληρώνονται με τον χαρακτήρα “A” μέχρι να φτάσουν το προκαθορισμένο κοινό μήκος που θα πρέπει να έχουν όλες οι ακολουθίες. Ο αριθμός αυτός προκύπτει από το μήκος της μεγαλύτερης συμβολοσειράς SMILES. Έτσι, συμπληρώνοντας παντού τον συγκεκριμένο χαρακτήρα, εξασφαλίζεται ότι όλες οι ακολουθίες έχουν το ίδιο μήκος, ανεξαρτήτως μεγέθους της αρχικής ακολουθίας, καθιστώντας τις διαχειρίσιμες από το LSTM μοντέλο.

Για παράδειγμα, στην περίπτωση όπου η αρχική ακολουθία ήταν η “[‘C’, ‘(’, ‘=’, ‘O’, ‘)’]” και το προκαθορισμένο μήκος κάθε ακολουθίας μετά την ολοκλήρωση του βήματος αυτού ήταν 10, το τελικό διάλυμα θα ήταν το “[‘G’, ‘C’, ‘(’, ‘=’, ‘O’, ‘)’’, ‘E’, ‘A’, ‘A’, ‘A’, ‘A’]”.

3. Κωδικοποίηση Ακολουθίας: Στο τελευταίο βήμα, η κάθε ακολουθία μήκους n , μετατρέπεται σε ένα διάνυσμα $n-1$ μηδενικών και ενός άσσου. Συγκεκριμένα, για τον κάθε διαθέσιμο χαρακτήρα που μπορεί να υπάρχει στην ακολουθία, έχει δημιουργηθεί ένα μοναδικό διάνυσμα, το οποίο έχει 1 στη θέση που αντιστοιχεί στο συγκεκριμένο σύμβολο και 0 σε όλες τις άλλες. Έτσι, δημιουργείται κάτι σαν λεξικό, μέσω του οποίου γίνεται η μετέπειτα αντιστοίχιση των ακολουθιών.

Για παράδειγμα, για το διάνυσμα του προηγούμενου βήματος, θα δημιουργούνταν το παρακάτω λεξικό:

'G': [1, 0, 0, 0, 0, 0, 0, 0, 0],
 'C': [0, 1, 0, 0, 0, 0, 0, 0, 0],
 '(': [0, 0, 1, 0, 0, 0, 0, 0, 0],
 '=: [0, 0, 0, 1, 0, 0, 0, 0, 0],
 'O': [0, 0, 0, 0, 1, 0, 0, 0, 0],
 ')': [0, 0, 0, 0, 0, 1, 0, 0, 0],
 'E': [0, 0, 0, 0, 0, 0, 1, 0, 0],
 'A': [0, 0, 0, 0, 0, 0, 0, 1, 0]

Μετά την αντιστοίχιση, ο τελικός πίνακας, σε μορφή κατάλληλη για είσοδο στο LSTM μοντέλο, θα ήταν ο εξής:

```
[ [1, 0, 0, 0, 0, 0, 0, 0, 0],
  [0, 1, 0, 0, 0, 0, 0, 0, 0],
  [0, 0, 1, 0, 0, 0, 0, 0, 0],
  [0, 0, 0, 1, 0, 0, 0, 0, 0],
  [0, 0, 0, 0, 1, 0, 0, 0, 0],
  [0, 0, 0, 0, 0, 1, 0, 0, 0],
  [0, 0, 0, 0, 0, 0, 1, 0, 0],
  [0, 0, 0, 0, 0, 0, 0, 1, 0],
  [0, 0, 0, 0, 0, 0, 0, 0, 1],
  [0, 0, 0, 0, 0, 0, 0, 0, 1],
  [0, 0, 0, 0, 0, 0, 0, 0, 1] ]
```

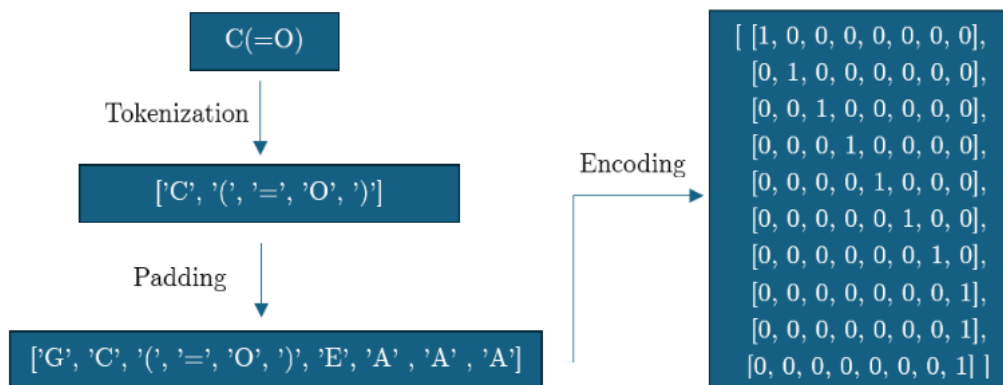


Figure 16: Encoding Steps

4.5 Μοντέλο Βαθιάς Μάθησης

Το μοντέλο που χρησιμοποιήθηκε για τις ανάγκες της εργασίας είναι ένα LSTM δίκτυο, το οποίο έχει σχεδιαστεί για την επεξεργασία και εκπαίδευση ακολουθιακών δεδομένων. Η μορφή των δεδομένων εισόδου, δηλαδή των SMILES συμβολοσειρών, όπου στην ουσία αποτελούν ακολουθίες χαρακτήρων που αντιστοιχούν σε άτομα, δεσμούς, διακλαδώσεις και κύκλους, καθιστά το LSTM μοντέλο ιδανική επιλογή για την εκπαίδευσή τους. Ο κύριος λόγος για το παραπάνω είναι πως τα LSTM μοντέλα είναι ικανά να μαθαίνουν μακροπρόθεσμες εξαρτήσεις, διατηρώντας πληροφορίες σε βάθος χρόνου, με αποτέλεσμα να μπορούν να προβλέψουν με μεγάλη ακρίβεια την επόμενη ακολουθία ενός μορίου.

Για τη δημιουργία και εκπαίδευση του μοντέλου σε python χρησιμοποιήθηκαν οι βιβλιοθήκες TensorFlow και Keras.

4.5.1 Δομή του Μοντέλου

Το μοντέλο αποτελείται από τρία επίπεδα, δύο LSTM και ένα Dense, τα οποία αναλύονται παρακάτω:

1. LSTM επίπεδο (1^ο):

- Αποτελείται από 256 νευρώνες, επιτρέποντάς του να διατηρεί και να επεξεργάζεται μακροπρόθεσμες εξαρτήσεις για τα δεδομένα έχοντας επαρκή χωρητικότητα μνήμης, χωρίς όμως να γίνεται υπερβολικά βαρύ υπολογιστικά ή να οδηγείται σε καταστάσεις υπερπροσαρμογής.
- Ως είσοδο δέχεται τους κωδικοποιημένους πίνακες των SMILES συμβολοσειρών.
- Για την αρχικοποίηση των βαρών χρησιμοποιείται η RandomNormal, διασφαλίζοντας έτσι την σταθεροποίηση των αρχικών βημάτων εκπαίδευσης και την εξασφάλιση μίας ισορροπημένης κατάστασης, καθώς τα βάρη αρχικοποιούνται ακολουθώντας την κανονική κατανομή.
- Η παράμετρος `return_sequences` έχει οριστεί ως `True`, καθώς δεδομένου ότι έχουμε δύο διαδοχικά LSTM επίπεδα, είναι απαραίτητο να επιστραφεί όλη η ακολουθία εξόδου, έτσι ώστε να χρησιμοποιηθεί ως είσοδος στο επόμενο, χωρίς να χαθεί οποιαδήποτε πληροφορία σε κάποιο προηγούμενο χρονικό βήμα.
- Για ποσοστό απόρριψης Dropout έχει επιλεγεί μία μέτρια τιμή, ίση με 0.3, επιτρέποντας έτσι στις συνδέσεις μεταξύ των δύο LSTM επιπέδων να παραμένουν ενεργές και σταθερές.

2. LSTM επίπεδο (2°):

- Σε αντίθεση με την συνήθη τακτική μείωσης των νευρώνων σε 128 στο δεύτερο επίπεδο, στη συγκεκριμένη περίπτωση διατηρήθηκαν σε 256 με σκοπό την εξασφάλιση της σταθερότητας και της δυνατότητας για επεξεργασία της πληροφορίας για την πλήρη ακολουθία εξαρτήσεων του μοντέλου, καθώς τα SMILES δεδομένα αποτελούν πολύπλοκα ακολουθιακά δεδομένα.
- Ως είσοδο δέχεται τα δεδομένα εξόδου του προηγούμενου επιπέδου.
- Αντίστοιχα, ως συνάρτηση αρχικοποίησης χρησιμοποιείται επίσης η RandomNormal.
- Η παράμετρος return_sequences είναι επίσης ορισμένη ως True, παρόλο που αποτελεί το τελευταίο επίπεδο LSTM και βρίσκεται πριν το Dense επίπεδο. Ο λόγος για τον οποίο έγινε η συγκεκριμένη επιλογή είναι η ανάγκη για πρόβλεψη κάθε βήματος μία ακολουθίας, και όχι μόνο το τελευταίο βήμα.
- Σε αυτό το επίπεδο έχει αυξηθεί η ένταση της απόρριψης σε 0.5, επιτρέποντας στο μοντέλο να “ξεχάσει” λιγότερο σημαντικές πληροφορίες και μειώνοντας ακόμα περισσότερο την πιθανότητα υπερπροσαρμογής στα δεδομένα εκπαίδευσης.

3. Dense επίπεδο (3°):

- Ο αριθμός των μονάδων σε αυτό το επίπεδο είναι ίσος με το μήκος του κωδικοποιημένου διανύσματος εισόδου, το οποίο αντιπροσωπεύει το σύνολο των διαφορετικών χαρακτήρων της SMILES ακολουθίας.
- Ως είσοδο δέχεται αντίστοιχα τον τρισδιάστατο πίνακα εξόδου του δεύτερου LSTM επιπέδου.
- Ως συνάρτηση ενεργοποίησης έχει επιλεγεί η Softmax, μία συνάρτηση ιδανική για προβλήματα πολυκατηγορικής ταξινόμησης όπως το συγκεκριμένο, αφού μετατρέπει τις εξόδους του επιπέδου σε πιθανότητες, οι οποίες αντιστοιχούν σε κάθε πιθανό χαρακτήρα. Έτσι, διαλέγοντας την πιθανότητα με την υψηλότερη τιμή, το μοντέλο μπορεί να προβλέψει τον χαρακτήρα που ακολουθεί σε κάθε περίπτωση.
- Ομοίως με τα προηγούμενα επίπεδα, η συνάρτηση αρχικοποίησης παραμένει η RandomNormal.

Για την επιλογή όλων των παραπάνω παραμέτρων, έγιναν δοκιμές μέσω κάποιων αναλυτικών περιγραμμάτων, σε συνδυασμό με την υπάρχουσα βιβλιογραφία και τις προτιμώμενες μεθόδους υλοποίησης ανάλογα την περίπτωση. Αυτή η απόφαση πάρθηκε λόγω της υψηλής περιπλοκότητας του μοντέλου, γεγονός που θα αύξανε σημαντικά το υπολογιστικό κόστος και θα καθιστούσε αρκετά απαιτητική και χρονοβόρα την παραπάνω διαδικασία.

Μία ακόμα σημαντική απόφαση ήταν η χρησιμοποίηση δύο LSTM επιπέδων αντί για ένα. Αυτή η επιλογή δίνει την ικανότητα στο μοντέλο να “μαθαίνει” πιο σύνθετα μοτίβα και εξαρτήσεις των ακολουθιών, καθώς στο πρώτο επίπεδο καταγράφονται βασικές σχέσεις μεταξύ των συμβόλων, οι οποίες μεταβιβάζονται στο δεύτερο επίπεδο για περαιτέρω επεξεργασία. Επιπλέον, βοηθά το μοντέλο να διατηρήσει τη μνήμη για περισσότερα βήματα στην ακολουθία, μειώνοντας έτσι τα πιθανά vanishing gradient προβλήματα, κάτι το οποίο είναι αρκετά κρίσιμο σε ακολουθιακά δεδομένα μεγάλου μεγέθους όπως τα συγκεκριμένα.

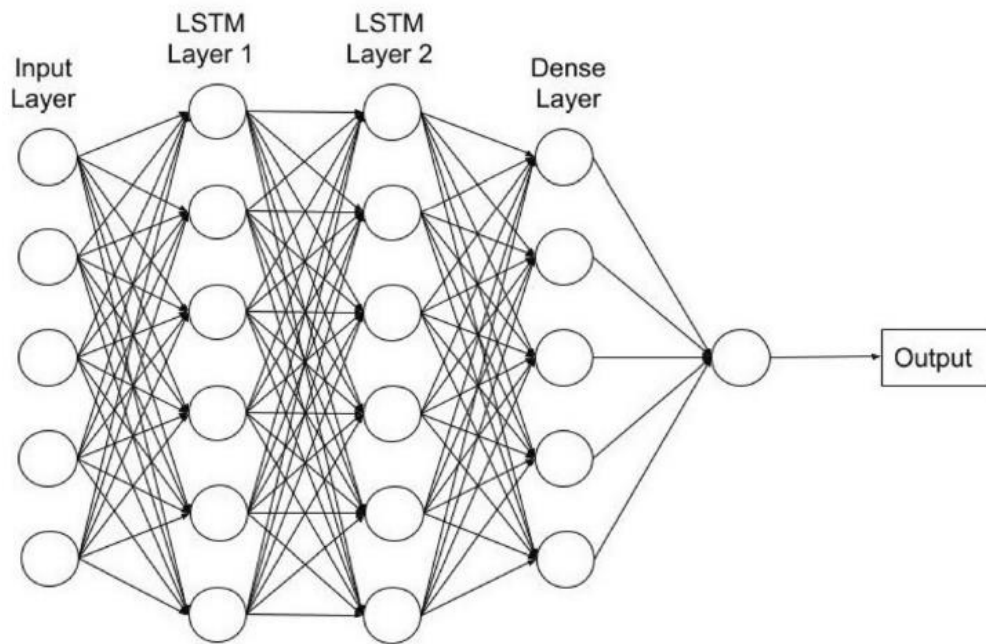


Figure 17: *Model Architecture Diagram*

4.5.2 Εκπαίδευση του Μοντέλου

Η διαδικασία εκπαίδευσης ενός νευρωνικού δικτύου περιλαμβάνει διάφορα στάδια, ενώ επίσης εξαρτάται από αρκετές παραμέτρους. Τα βήματα που ακολουθήθηκαν, καθώς και επιλογές των παραμέτρων που χρησιμοποιήθηκαν, περιγράφονται παρακάτω:

- **Δεδομένα:** Τα πλέον επεξεργασμένα δεδομένα διαχωρίζονται σε δύο τμήματα, το σύνολο εκπαίδευσης και το σύνολο επικύρωσης. Συγκεκριμένα, για σύνολο επικύρωσης επιλέγεται το 10% των δεδομένων και χρησιμοποιείται για την παρακολούθηση της απόδοσης του μοντέλου κατά τη διάρκεια της εκπαίδευσης, εντοπίζοντας σημάδια υπερπροσαρμογής και διασφαλίζοντας την ικανότητα γενίκευσης του μοντέλου.
- **Αλγόριθμος Βελτιστοποίησης:** Η διαδικασία εκπαίδευσης του μοντέλου βασίζεται αρκετά στον αλγόριθμο βελτιστοποίησης, αφού είναι υπεύθυνος για το πώς το μοντέλο “μαθαίνει” και προσαρμόζει τα βάρη του. Για το μοντέλο αυτού του πειράματος έχει επιλεγεί ο Adam, ο οποίος μπορεί να προσαρμόζει δυναμικά τον ρυθμό μάθησης, δίνοντας έτσι τη δυνατότητα στο μοντέλο να συγκλίνει πιο γρήγορα, αποφεύγοντας προβλήματα όπως η αστάθεια στις τιμές των gradients.
- **Συνάρτηση Απώλειας:** Για τον υπολογισμό των σφαλμάτων κατά τη διάρκεια της εκπαίδευσης έχει χρησιμοποιηθεί η κατηγορική cross-entropy, η οποία υπολογίζει την απόκλιση μεταξύ των πραγματικών τιμών και των προβλεπόμενων πιθανοτήτων του μοντέλου. Έτσι, μέσω συνεχούς προσαρμογής των βαρών του, το μοντέλο προσπαθεί να ελαχιστοποιήσει αυτή την απώλεια.
- **Batch Size:** Απαραίτητος είναι ο διαχωρισμός των δεδομένων σε ομάδες δεδομένων (batches), έτσι ώστε να μειωθεί το υπολογιστικό κόστος και να αυξηθεί η ταχύτητα εκπαίδευσης του μοντέλου, εξισορροπώντας έτσι τη σταθερότητά του. Επίσης, οι ομάδες δεδομένων βοηθούν και στην προσαρμογή των βαρών, καθώς για κάθε ένα από αυτά υπολογίζεται το σφάλμα του και μέσω του αλγορίθμου οπισθοδιάδοσης προσαρμόζονται τα βάρη του μοντέλου. Για τις ανάγκες του συγκεκριμένου πειράματος ως αριθμός batch size επιλέχθηκε το 256. Έτσι ώστε να υπάρξει μία ισορροπία μεταξύ της σταθερότητας των gradients και των διαθέσιμων υπολογιστικών πόρων.
- **Εποχές:** Αφότου λοιπόν έχουν επιλεγεί όλες οι παραπάνω παράμετροι, μπορεί να ξεκινήσει η εκπαίδευση του μοντέλου, η οποία πραγματοποιείται σε μικρά βήματα, τις λεγόμενες εποχές, μέσω των οποίων παρακολουθείται η απόδοσή του. Λόγω του μεγάλου όγκου των δεδομένων και της πολυπλοκότητας του μοντέλου, επιλέχθηκε να εκπαιδευτεί για 30 εποχές.

- Callbacks: Στην εκκίνηση της εκπαίδευσης, ορίζονται και κάποιες callback συναρτήσεις, οι οποίες παίζουν σημαντικό ρόλο στη διασφάλιση της βέλτιστης απόδοσης του μοντέλου.
 - ModelCheckpoint: Στο τέλος κάθε εποχής εξετάζεται η validation loss τιμή, και εάν έχει βελτιωθεί σε σχέση με την μέχρι πρότερος καλύτερη τιμή της, τότε η callback συνάρτηση ModelCheckpoint αναλαμβάνει να αποθηκεύσει τη συγκεκριμένη κατάσταση του μοντέλου, δίνοντας την δυνατότητα για επιστροφή στην καλύτερη έκδοσή του.
 - EarlyStopping: Η συνάρτηση αυτή χρησιμοποιείται για τον τερματισμό της εκπαίδευσης σε περίπτωση που η προαναφερθείσα τιμή validation loss δεν έχει βελτιωθεί σε ένα διάστημα 5 εποχών. Κατά την εκπαίδευση ωστόσο του συγκεκριμένου μοντέλου, δεν χρειάστηκε να επέμβει.
 - TensorBoard: Μία ακόμα callback συνάρτηση που έχει επιλεγεί για την εκπαίδευση του μοντέλου είναι η TensorBoard, μέσω της οποίας καταγράφεται η συνολική διαδικασία, επιτρέποντας την οπτικοποίηση της πορείας της εκπαίδευσης μέσω γραφημάτων των σφαλμάτων. Τα γραφήματα αυτά παραθέτονται παρακάτω.

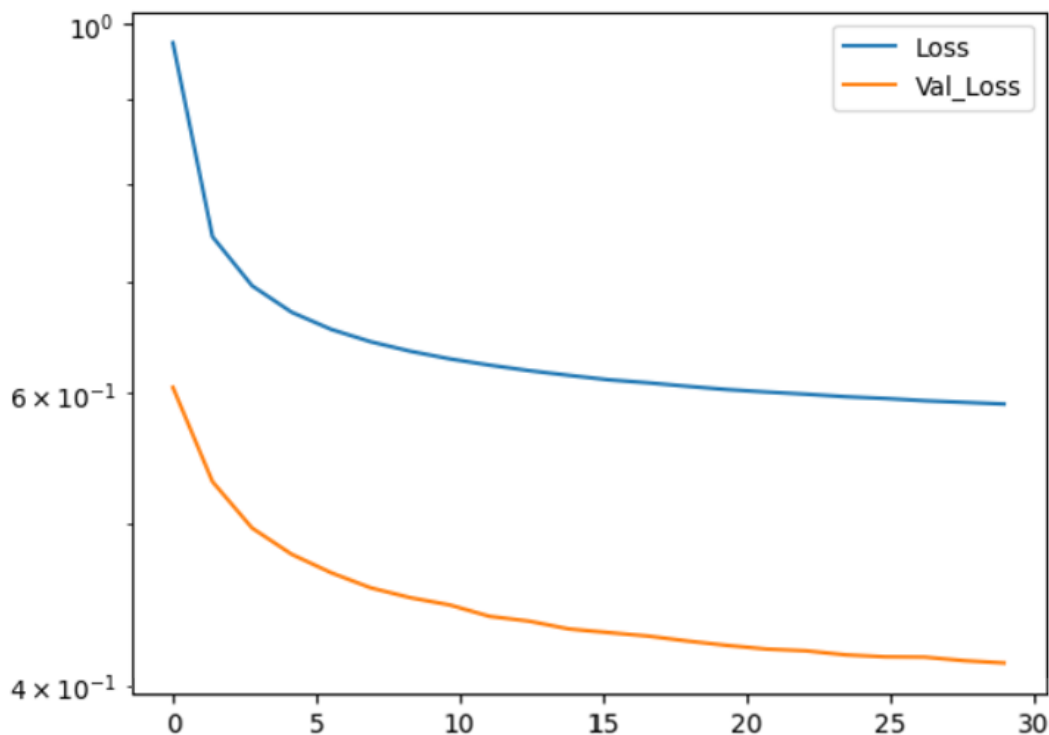


Figure 18: *Training and Validation Loss over Epochs*

4.6 Γενετική Παραγωγή Μορίων

4.6.1 Μέθοδος

Μετά την ολοκλήρωση της εκπαίδευσης, το μοντέλο έχει καταφέρει να μάθει τη δομή και τα χαρακτηριστικά των μοριακών ακολουθιών. Σκοπός λοιπόν αυτής της φάσης είναι η χρησιμοποίηση των γνώσεων που έχει αποκομίσει το μοντέλο, για τη δημιουργία νέων χημικών ακολουθιών.

Η διαδικασία παραγωγής νέων μορίων γίνεται μέσω δειγματοληψίας, όπου σε κάθε βήμα επιλέγεται ο επόμενος χαρακτήρας της ακολουθίας με βάση την κατανομή πιθανοτήτων που προβλέπει το μοντέλο. Σημαντικό ρόλο έχει και η παράμετρος temperature, η οποία επηρεάζει τον βαθμό ποικιλομορφίας των αποτελεσμάτων. Τα βήματα της διαδικασίας παρουσιάζονται αναλυτικά παρακάτω:

1. Η διαδικασία ξεκινάει καλώντας τη συνάρτηση sample, η οποία δέχεται ως παραμέτρους τον αριθμό των επιθυμητών προς παραγωγή SMILES συμβολοσειρών και το εάν είναι απαραίτητο τα μόρια που αντιπροσωπεύουν αυτές οι συμβολοσειρές να είναι έγκυρα. Ακόμα, δέχεται και το αρχικό σύμβολο των παραγόμενων ακολουθιών, το οποίο από προεπιλογή είναι το “G”, όπου όπως έχει αναφερθεί σε προηγούμενο κεφάλαιο, συμβολίζει την έναρξη της ακολουθίας.
2. Στη συνέχεια, και όσο το πιο πρόσφατο σύμβολο δεν είναι το “E”, το οποίο συμβολίζει τη λήξη της ακολουθίας, το μοντέλο προχωράει στην πρόβλεψη των επόμενων συμβόλων. Συγκεκριμένα, για αρχή επαναλαμβάνει τη διαδικασία επεξεργασίας των δεδομένων που περιγράφηκε στην ενότητα 4.4.3. Παίρνει δηλαδή την μέχρι πρότενος ακολουθία, τη διασπά σε ξεχωριστούς χαρακτήρες και στη συνέχεια τους κωδικοποιεί ώστε να είναι αναγνώσιμοι από το μοντέλο.
3. Έπειτα, αυτή η ακολουθία εισάγεται στο μοντέλο, το οποίο επιστρέφει τις πιθανότητες για τον χαρακτήρα που ακολουθεί. Αυτή η κατανομή πιθανοτήτων θα χρησιμοποιηθεί στο επόμενο βήμα για τον καθορισμό του προβλεπόμενου χαρακτήρα.

4. Ο τρόπος που γίνεται αυτό δίνει και την ουσιαστική επιλογή ως προς την τυχαιότητα, καθώς και την “διαφορετικότητα” του επόμενου συμβόλου. Συγκεκριμένα, εισάγεται η παράμετρος της λεγόμενης θερμοκρασίας [37], μία τιμή η οποία καθορίζει το πόσο συντηρητική ή “εξερευνητική” θα είναι η παραγόμενη ακολουθία. Παίρνοντας λοιπόν τις προβλέψεις, η συνάρτηση παραγωγής των μορίων τις τροποποιεί μέσω μίας λογαριθμικής κλίμακας και στη συνέχεια διαιρεί το αποτέλεσμα με την παράμετρο θερμοκρασίας. Έτσι, για υψηλότερες τιμές αυτής της παραμέτρου, η παραγόμενη ακολουθία έχει μεγαλύτερη ποικιλία και διαφορετικότητα, ωστόσο μειώνεται η πιθανότητα να είναι έγκυρη, ενώ για μικρότερες τιμές η πρόβλεψη μπορεί να είναι πιο σταθερή, παραμένει όμως συντηρητική.
5. Το σύμβολο που προκύπτει από τον παραπάνω υπολογισμό προστίθεται στην ήδη υπάρχουσα ακολουθία και στη συνέχεια η διαδικασία επαναλαμβάνεται έως ότου παραχθεί το σύμβολο λήξης “E”.
6. Όταν λοιπόν έρθει ο συγκεκριμένος χαρακτήρας, η συνάρτηση επιστρέφει τη συμβολοσειρά SMILES, έχοντας αφαιρέσει τα σύμβολα αναφοράς “G” και “E”.

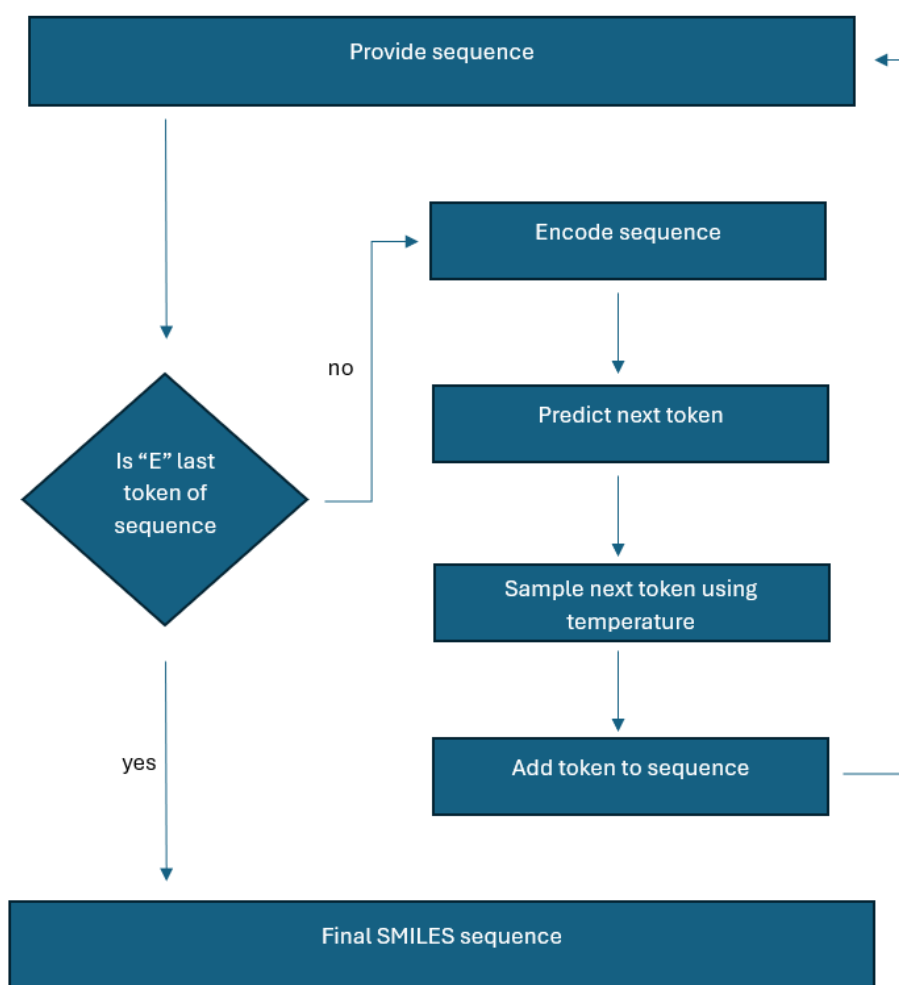


Figure 19: *Sampling Process*

4.6.2 Δημιουργία και Σχολιασμός Αποτελεσμάτων

Δεδομένου ότι η παράμετρος θερμοκρασίας επηρεάζει σε πολύ μεγάλο βαθμό τη διαδικασία δειγματοληψίας χαρακτήρων, άρα και τα νέα μόρια που προκύπτουν, επιλέχθηκαν τρεις διαφορετικές τιμές της για την πραγματοποίηση των πειραμάτων. Σκοπός είναι η σύγκριση των αποτελεσμάτων για κάθε τιμή της παραμέτρου και η επιλογή εκείνης που αποφέρει τον καλύτερο συνδυασμό εγκυρότητας και ποικιλομορφίας των μορίων. Οι τρεις διαφορετικές τιμές που δοκιμάστηκαν για την παράμετρο της θερμοκρασίας είναι οι 0.5, 0.75 και 1. Παρακάτω παρουσιάζονται αρχικά τα βήματα που εκτελέστηκαν για τη δημιουργία των αποτελεσμάτων, και στη συνέχεια τα διαφορετικά αποτελέσματα για τις τρεις τιμές της θερμοκρασίας.

4.6.2.1 Βήματα

1. Σε πρώτη φάση, γίνεται δειγματοληψία 1000 SMILES ακολουθιών. Για να υπολογιστεί η εγκυρότητα των νέων ακολουθιών χρειάζεται να μετατραπούν σε μόρια μέσω της συνάρτησης Chem.MolFromSmiles της βιβλιοθήκης RDKit, η οποία επιστρέφει None εάν η SMILES συμβολοσειρά που δέχεται ως είσοδο δεν αντιστοιχεί σε έγκυρο μόριο.
2. Στη συνέχεια, γίνεται έλεγχος της καινοτομίας των νέων μορίων, υπολογίζοντας πόσα από τα μόρια που προέκυψαν από τη δειγματοληψία υπάρχουν στα δεδομένα που χρησιμοποιήθηκαν για την εκπαίδευση του μοντέλου.
3. Έπειτα, υπολογίζονται τα χημικά χαρακτηριστικά των μορίων και συγκρίνονται με αυτά των αρχικών μορίων. Συγκεκριμένα, η GetMACCSKeysFingerprint συνάρτηση της RDKit βιβλιοθήκης υπολογίζει 166 δυαδικά χαρακτηριστικά. Το κάθε ένα από αυτά παίρνει την τιμή 1 εάν το μόριο έχει το συγκεκριμένο μοτίβο και 0 εάν δεν το έχει. Ωστόσο, λόγω της δυσκολίας αναπαράστασης τόσων διανυσμάτων 166 τιμών, χρειάζεται να γίνει μείωση των διαστάσεων.

Αυτό μπορεί να γίνει με αρκετούς τρόπους. Ένας από αυτούς είναι η χρήση της τεχνικής PCA (Principal Component Analysis), ενός γραμμικού μετασχηματισμού των δεδομένων σε 2 διαστάσεις, έτσι ώστε να είναι πιο εύκολη η οπτικοποίησή τους. Ωστόσο, για τα δεδομένα αυτού του πειράματος, η τεχνική UPAM (Uniform Manifold Approximation and Projection) καθίσταται ιδανική για τη μείωση των διαστάσεων, καθώς μπορεί και αναγνωρίζει και μη γραμμικούς συνδυασμούς, αποδίδοντας έτσι καλύτερα την πολύπλοκη γεωμετρία του χημικού χώρου.

Τα μειωμένων διαστάσεων δεδομένα αναπαρίστανται στη συνέχεια μέσω διαγραμμάτων διασποράς (scatter plots).

4. Στο προηγούμενο βήμα παρουσιάζεται ουσιαστικά η ικανότητα του μοντέλου να παράγει μόρια, των οποίων ο χημικός χώρος συμπίπτει με αυτόν των αρχικών μορίων. Στη συνέχεια, η ανάλυση εστιάζει σε συγκεκριμένες ιδιότητες, οι οποίες έχουν περιγραφεί εκτενώς στην ενότητα 3.1.4. Ειδική μνεία γίνεται στις τιμές μοριακού βάρους (molecular weight) και λιποφιλικότητας (logP), ως οι δύο πιο σημαντικές για τη βιοδραστικότητα του μορίου.
5. Τέλος, για να γίνει ακόμα πιο κατανοητή η επιρροή της θερμοκρασίας στη διαδικασία δειγματοληψίας, παραθέτονται τα διαγράμματα διασποράς των παραπάνω μετρικών, σε αντιπαραβολή με εκείνες των αρχικών δεδομένων.

4.6.2.2 Αποτελέσματα

- Τιμή Θερμοκρασίας 0.5:
 - Αριθμός Μορίων για Δειγματοληψία: 1000
 - Ποσοστό Έγκυρων SMILES: 88.6%
 - Ποσοστό Μοναδικών SMILES: 90.4%

Διάγραμμα UMAP:

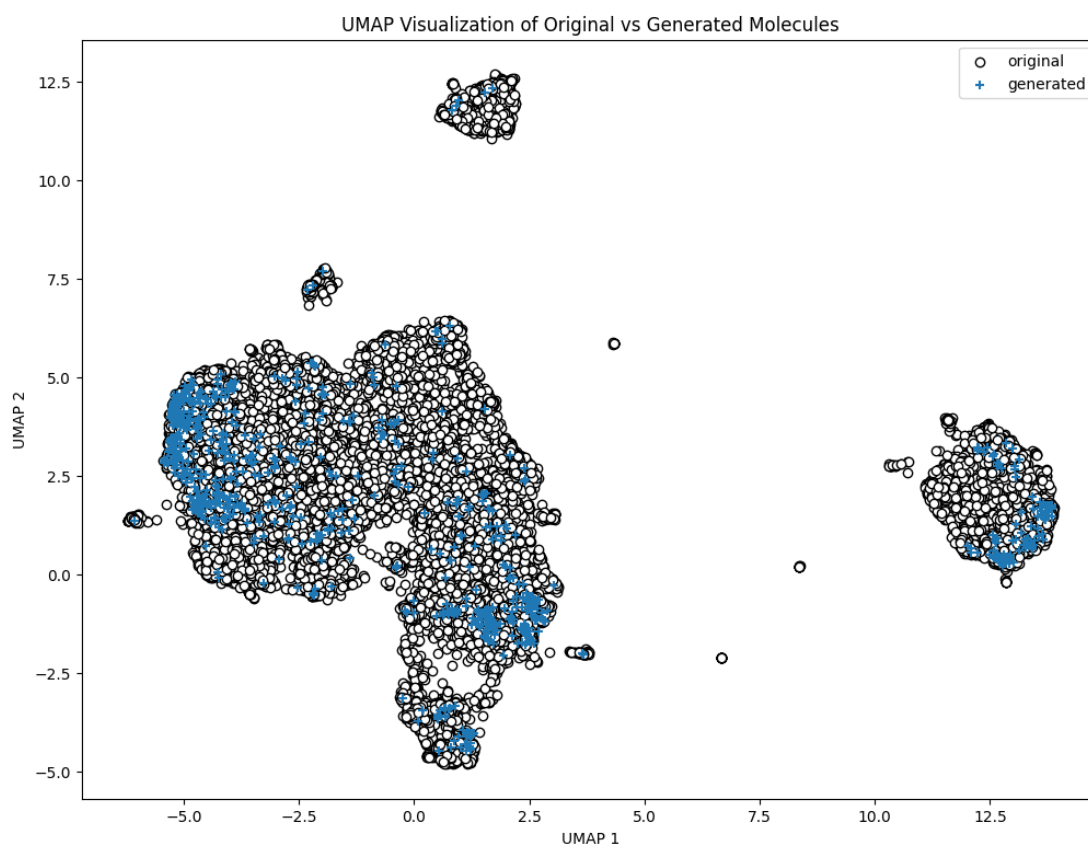


Figure 20: UMAP Analysis with temperature = 0.5

Από το παραπάνω διάγραμμα φαίνεται ξεκάθαρα ο διαχωρισμός σε τρεις συστάδες δεδομένων (clusters), οι οποίες αντιπροσωπεύουν διαφορετικούς χημικούς χώρους και μοριακά μοτίβα. Το σημαντικό ωστόσο που παρατηρείται σε αυτή την αναπαράσταση είναι ότι τα παραγόμενα μόρια βρίσκονται εντός των συστάδων που σχηματίζουν τα αρχικά μόρια. Αυτό σημαίνει ότι το μοντέλο κατάφερε να αφομιώσει τα χαρακτηριστικά των αρχικών δεδομένων και να τα αναπαράξει κατά τη δειγματοληψία νέων μορίων.

Διαγράμματα Μεμονωμένων Ιδιοτήτων:

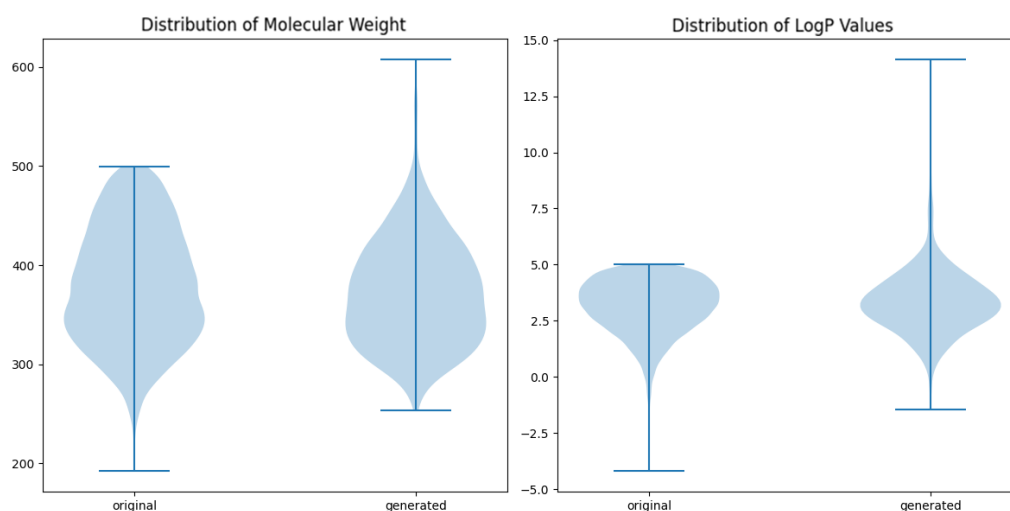


Figure 21: *Molecular Weight & logP distribution with temperature = 0.5*

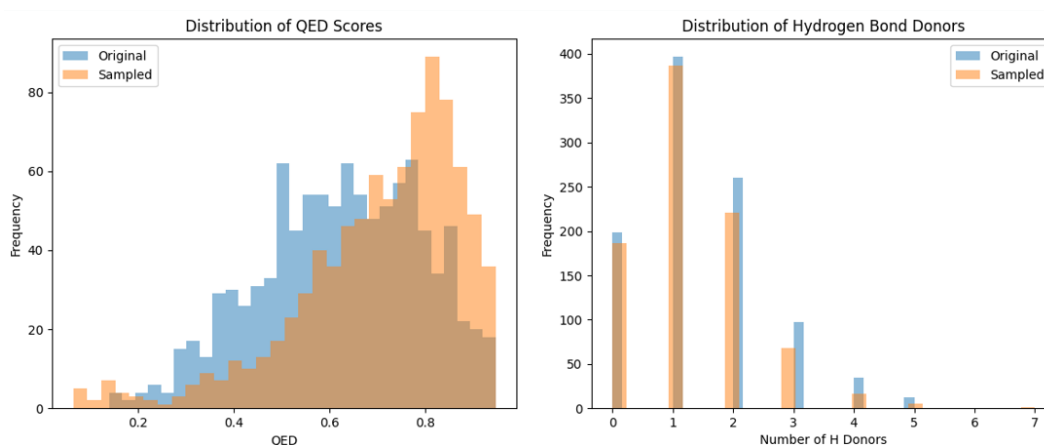


Figure 22: *QED Score & Number of Donors with temperature = 0.5*

Όπως φαίνεται στα παραπάνω διαγράμματα, οι κατανομές του μοριακού βάρους και του logP συγκλίνουν στις ίδιες τιμές, ενώ η κατανομή του QED τείνει να συγκλίνει σε πιο υψηλές τιμές για τα νέα μόρια. Αυτό δείχνει πως τα παραγόμενα μόρια δε διαφέρουν σημαντικά στα κύρια χαρακτηριστικά, ενώ ταυτόχρονα έχουν μία αυξημένη πιθανότητα για καλή βιολογική δραστηριότητα και φαρμακοκινητική συμπεριφορά.

- Τιμή Θερμοκρασίας 0.75:
 - Αριθμός Μορίων για Δειγματοληψία: 1000
 - Ποσοστό Έγκυρων SMILES: 73.9%
 - Ποσοστό Μοναδικών SMILES: 94.6%

Διάγραμμα UMAP:

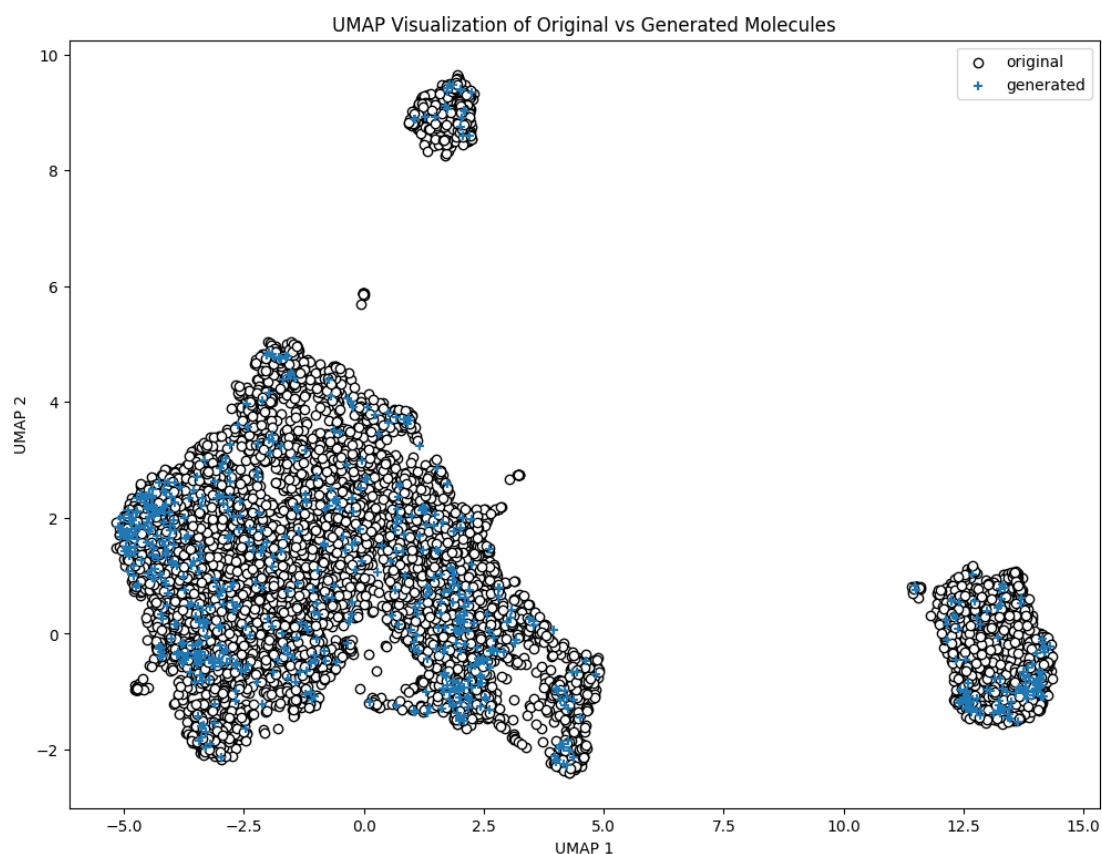


Figure 23: UMAP Analysis with temperature = 0.75

Στην παραπάνω εικόνα φαίνεται πως παρόλο που υπάρχει μία αύξηση στην παράμετρο της θερμοκρασίας, τα παραγόμενα μόρια εξακολουθούν να συμπίπτουν με τον χημικό χώρο των αρχικών δεδομένων. Αυτό αποδεικνύει την ικανότητα του μοντέλου να μπορεί να αφωμιώνει τις ιδιότητες των μορίων και να παράγει νέα μόρια με αντίστοιχες ιδιότητες.

Διαγράμματα Μεμονωμένων Ιδιοτήτων:

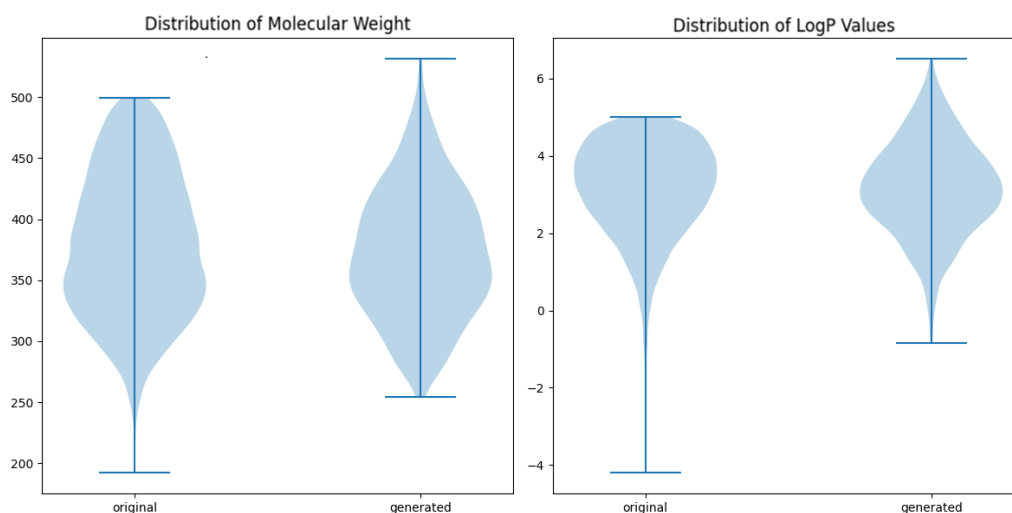


Figure 24: *Molecular Weight & logP distribution with temperature = 0.75*

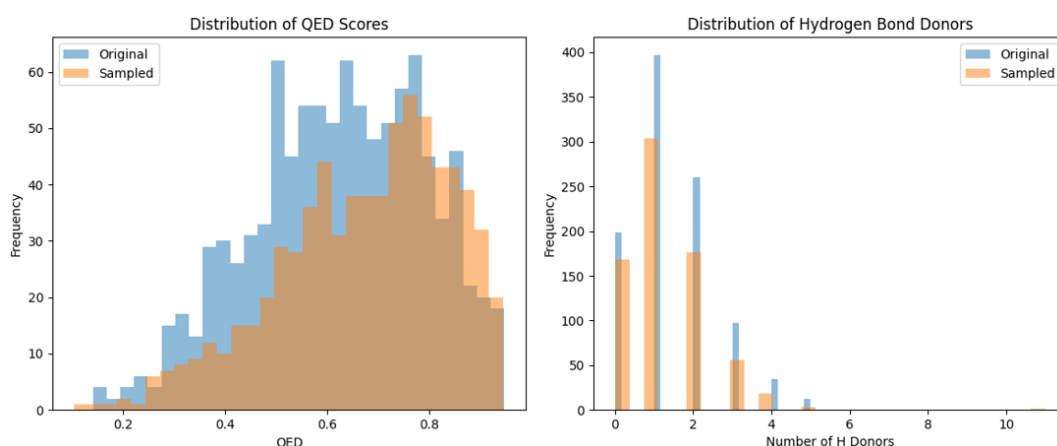


Figure 25: *QED Score & Number of Donors with temperature = 0.75*

Σε αυτήν την περίπτωση υπάρχουν πιο εμφανείς διαφορές σε σχέση με τα αρχικά δεδομένα. Οι κατανομές του μοριακού βάρους και της λιποφιλικότητας (logP) των νέων μορίων παρουσιάζουν μεγαλύτερη ομοιομορφία σε σχέση με αυτές των αρχικών, γεγονός που φανερώνει την ύπαρξη ποικιλομορφίας. Το QED Score παραμένει ικανοποιητικό, δίνοντας υψηλές πιθανότητες για φαρμακευτική αποτελεσματικότητα των ενώσεων.

- Τιμή Θερμοκρασίας 1:
 - Αριθμός Μορίων για Δειγματοληψία: 1000
 - Ποσοστό Έγκυρων SMILES: 47.2%
 - Ποσοστό Μοναδικών SMILES: 99.8%

Διάγραμμα UMAP:

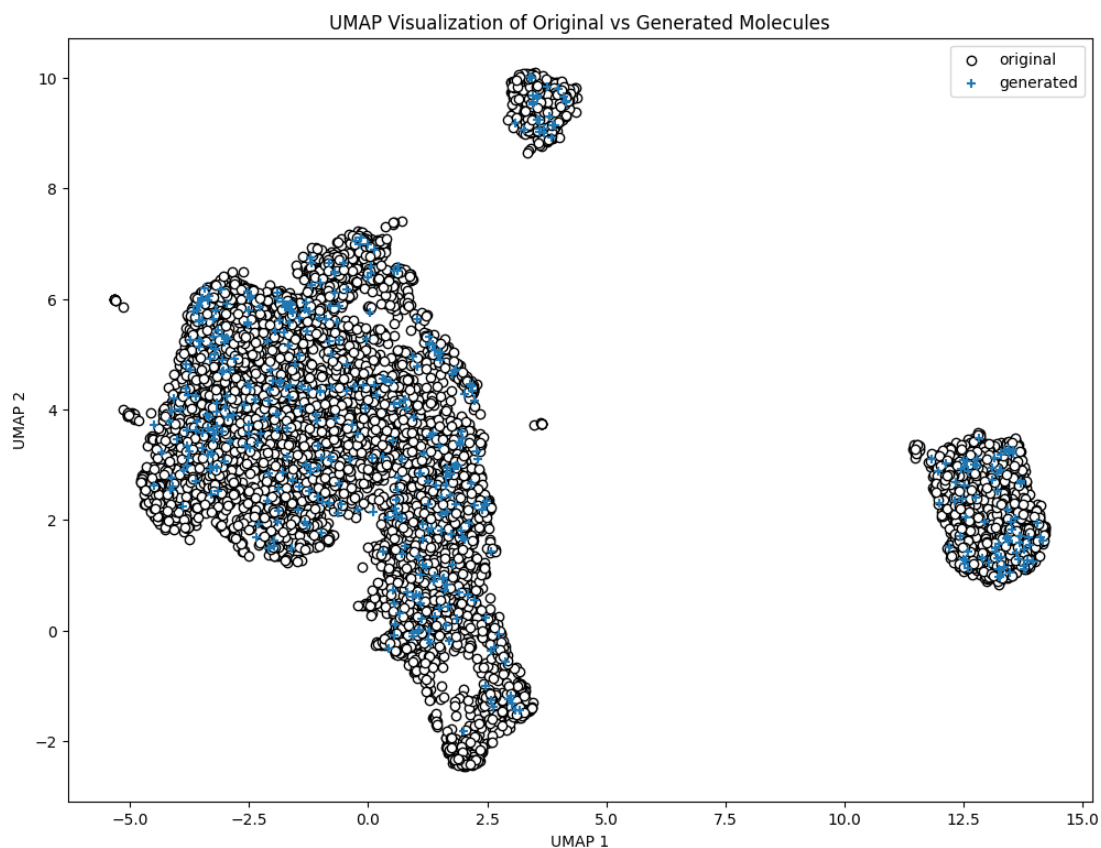


Figure 26: *UMAP Analysis with temperature = 1*

Στην παραπάνω εικόνα φαίνεται πως παρόλο που υπάρχει μία αύξηση στην παράμετρο της θερμοκρασίας, τα παραγόμενα μόρια εξακολουθούν να συμπίπτουν με τον χημικό χώρο των αρχικών δεδομένων. Αυτό αποδεικνύει την ικανότητα του μοντέλου να μπορεί να αφωμιώνει τις ιδιότητες των μορίων και να παράγει νέα μόρια με αντίστοιχες ιδιότητες.

Διαγράμματα Μεμονωμένων Ιδιοτήτων:

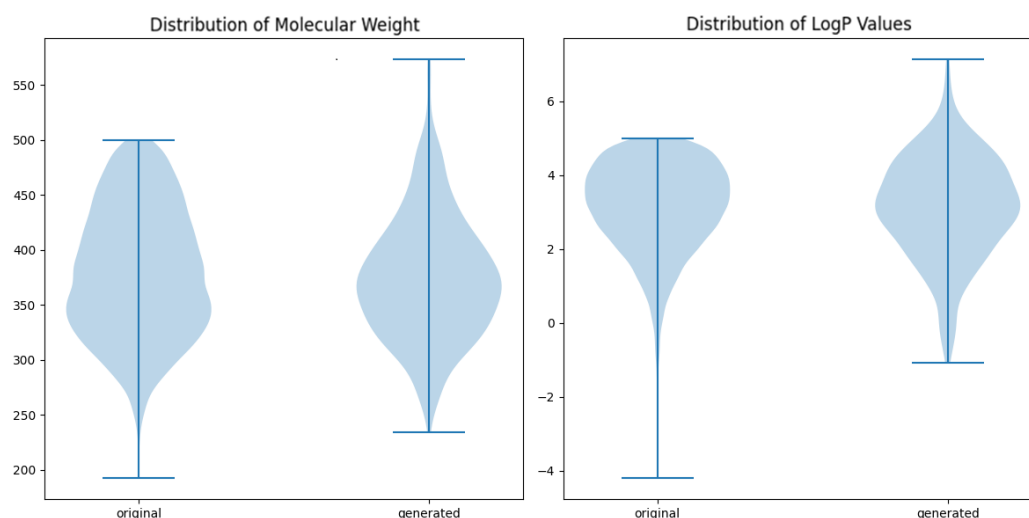


Figure 27: *Molecular Weight & logP distribution with temperature = 1*

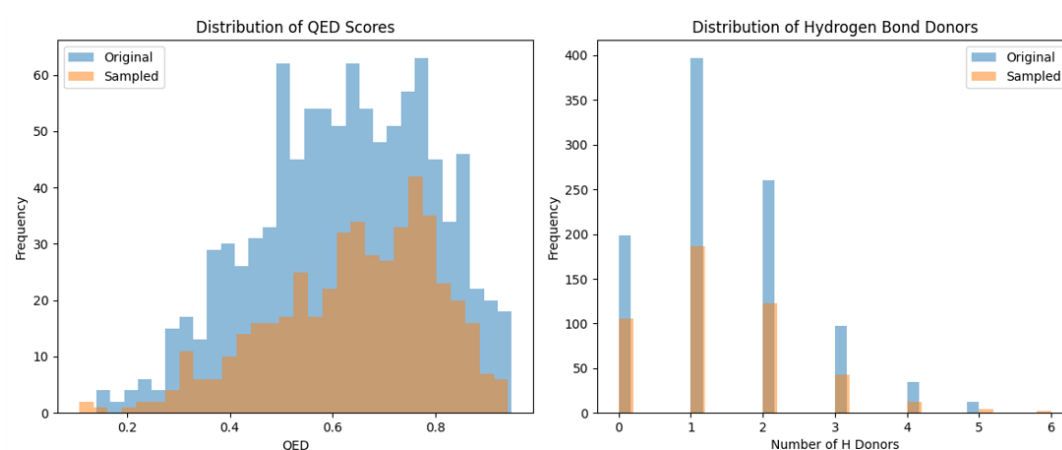


Figure 28: *QED Score & Number of Donors with temperature = 1*

Στην τελευταία περίπτωση, για θερμοκρασία ίση με 1 δηλαδή, η κατανομές του μοριακού βάρους και της λιποφιλικότητας των παραγόμενων μορίων διαφέρουν και πάλι ως προς την διασπορά τους. Μπορεί να μην υπάρχει τόσο μεγάλη αύξηση της ποικιλομορφίας σε σχέση με τις χαμηλότερες τιμές θερμοκρασίας, ωστόσο από τη σχετικά πιο επίπεδη κατανομή των νέων μορίων στο QED Score διάγραμμα, παρατηρείται η επίδραση της υψηλής θερμοκρασίας στην ικανότητα σχηματισμού αποδοτικών, φαρμακευτικά, μορίων.

- Σύγκριση των αποτελεσμάτων:

Τιμές		Αρχικά Δεδομένα	Temp = 0.5	Temp = 0.75	Temp = 1
Μοριακό Βάρος	mean	379.69	372.80	370.78	374.84
	std	55.65	53.73	55.87	58.20
logP	mean	3.09	3.40	3.27	3.13
	std	1.21	1.15	1.38	1.49
QED Score	mean	0.61	0.71	0.63	0.60
	std	0.17	0.17	0.16	0.17
Number of H Donors	mean	1.41	1.28	1.32	1.43
	std	1.11	0.97	1.05	1.10

Ο παραπάνω πίνακας παρέχει μία σαφή εικόνα για το πώς η παράμετρος θερμοκρασίας επηρεάζει τις ιδιότητες των παραγόμενων μορίων. Συγκεκριμένα, για χαμηλότερες τιμές, υπάρχει μία αυξημένη καταλληλότητα των νέων μορίων για φαρμακευτική χρήση, χωρίς όμως μεγάλη ποικιλία, ενώ όσο η τιμή της θερμοκρασίας αυξάνεται, οι ιδιότητες αυτές αλλάζουν με αντιστρόφως ανάλογο τρόπο.

Ωστόσο, σημαντική παραμένει η ικανότητα του μοντέλου να παράγει μόρια κοντά στον χημικό χώρο των αρχικών, ανεξάρτητα από την τιμή της θερμοκρασίας, όπως φαίνεται από τα 3 διαγράμματα UMAP που παρατέθηκαν στις αντίστοιχες ενότητες.

4.7 Παραγωγή Μορίων από Αρχικό Τμήμα

Μία ακόμα δυνατότητα του μοντέλου που εξετάζεται στο παρόν πείραμα, είναι η δημιουργία μορίων με βάση κάποιο προυπάρχον τμήμα, με πιθανώς γνωστές βιοχημικές ιδιότητες.

Η διαδικασία δειγματοληψίας μορίων με βάση κάποιο αρχικό τμήμα δεν αλλάζει δραστηριότητα από την ήδη υπάρχουσα διαδικασία για την γενική παραγωγή μορίων. Η μοναδική διαφοροποίηση είναι πως αντί να δοθεί ως αρχική ακολουθία ο χαρακτήρας “G”, δίνεται η η ακολουθία του τμήματος αναφοράς μαζί με το σύμβολο “G” πριν την έναρξή της.

Η ουσία που επιλέχθηκε να είναι η βάση για τη μετέπειτα δειγματοληψία μορίων είναι η νικοτιναμίδα, η οποία αποτελεί τη φυσική δραστήρια μορφή της βιταμίνης B3. Αποτελείται από έναν δακτύλιο πυριδίνης με μία αμιδική ομάδα, χαρακτηριστικά που την καθιστούν χημικά σταθερή και βιολογικά δραστήρια, και άρα κατάλληλη υποψήφια ουσία για βάση των παραγόμενων μορίων.

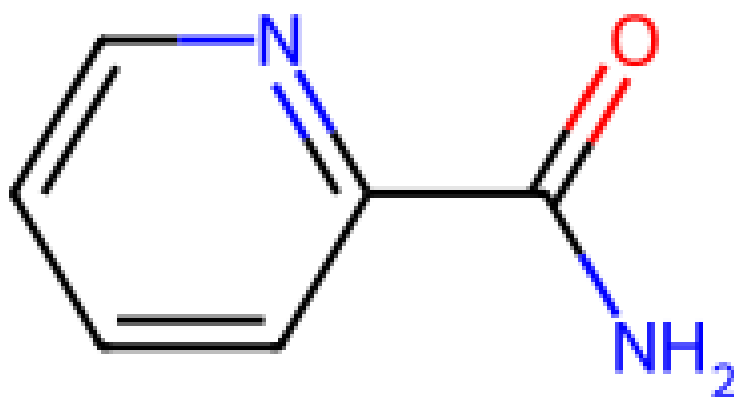


Figure 29: *Nicotinamide Chemical Structure*

Μετά τη διαδικασία δειγματοληψίας, τα 1000 μόρια που έχουν παραχθεί παρουσιάζουν τα παρακάτω χαρακτηριστικά:

- ο Ποσοστό Έγκυρων SMILES: 78.2%
- ο Ποσοστό Μοναδικών SMILES από το σύνολο των έγκυρων: 93.57%
- ο Μέση τιμή Μοριακού Βάρους: 343.3 Da
- ο Μέση τιμή LogP: 2.85
- ο Μέση τιμή QED Score: 0.75
- ο Μέση τιμή Αριθμού Δοτών Υδρογόνου: 0.86

Από τα παραπάνω ποσοστά φαίνεται πως το μοντέλο είναι ικανό να παράγει ένα ικανοποιητικό ποσοστό έγκυρων SMILES ακολουθιών με βάση κάποιο αρχικό τμήμα, χωρίς να δημιουργεί πολλά διπλότυπα, καθώς μετά την κανονικοποίηση των δεδομένων το ποσοστό κοινών ακολουθιών είναι μόλις 6.43%.

Επίσης, σημαντικές είναι και οι τιμές των υπόλοιπων παραμέτρων των νέων μορίων, οι οποίες βρίσκονται μέσα στο αποδεκτό εύρος των αποδοτικών φαρμακευτικά ουσιών, ενώ ειδικά η αρκετά υψηλή τιμή του δείκτη QED score τα καθιστά ικανά για να εξελιχθούν σε φαρμακολογικά δραστικά μόρια.

Παρακάτω παρουσιάζεται η χημική δομή 5 τυχαίων μορίων που παράχθηκαν έχοντας ως αρχικό τμήμα τη νικοτιναμίδα, μαζί με το ίδιο το τμήμα, όπου γίνεται εμφανής η ξεκάθαρη παρουσία του σε αυτά.

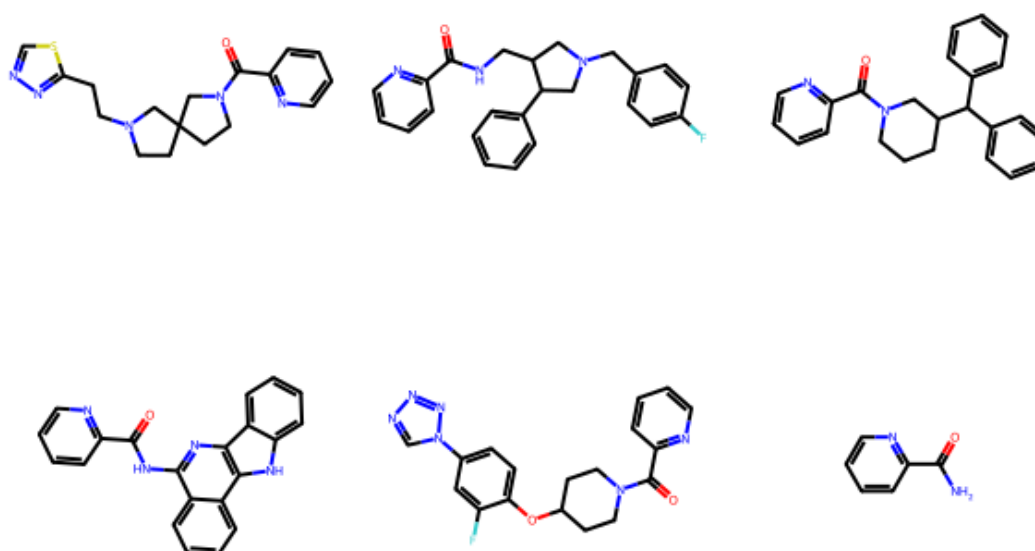


Figure 30: Five FBDD Generated Molecules & Starting Fragment

4.8 Μεταφορά Μάθησης για Επικεντρωμένη Παραγωγή Μορίων

Το τελευταίο κομμάτι που εξετάζεται σε αυτή την εργασία είναι η χρήση της τεχνικής μεταφοράς μάθησης για παραγωγή μορίων, με σκοπό να έχουν χαρακτηριστικά συγκεκριμένων τύπων μορίων. Αυτό γίνεται μέσω διαδικασίας fine-tuning του ήδη εκπαιδευμένου μοντέλου σε ένα συγκεκριμένο υποσύνολο δεδομένων, με αποτέλεσμα την εξειδίκευσή του σε αυτά. Η διαδικασία αυτή αποτελείται από αρκετά βήματα, αλλά παραμένει η ίδια για κάθε υποσύνολο που εξετάζεται σε αυτό το κεφάλαιο.

4.8.1 Βήματα της Διαδικασίας

1. **Φόρτωση Δεδομένων:** Η διαδικασία ξεκινάει με τη φόρτωση του αρχικού μοντέλου, το οποίο έχει ήδη εκπαιδευτεί σε ένα μεγάλο σύνολο δεδομένων, μαθαίνοντας γενικά χαρακτηριστικά και χημικές ιδιότητες των μορίων. Το αρχικό μοντέλο λοιπόν, έχοντας μάθει τα γενικά μοτίβα των δεδομένων, είναι κατάλληλο να χρησιμοποιηθεί ως αφετηρία για την περαιτέρω εκπαίδευση, εξοικονομώντας έτσι χρόνο και υπολογιστικούς πόρους. Στη συνέχεια, φορτώνονται και τα δεδομένα, πάνω στα οποία θα γίνει το fine-tuning, και συνήθως περιλαμβάνουν ομάδες ή τύπους μορίων που παρουσιάζουν συγκεκριμένες ιδιότητες και χαρακτηριστικά. Προφανώς, τα δεδομένα αυτά ακολουθούν όλα τα βήματα προεπεξεργασίας που πραγματοποιήθηκαν και στα αρχικά.
2. **Εκπαίδευση του Μοντέλου:** Κατά τη διαδικασία της εκπαίδευσης, το μοντέλο εξειδικεύεται στα νέα δεδομένα, διατηρώντας ωστόσο μεγάλο μέρος της πληροφορίας που έμαθε από το ευρύτερο σύνολο. Η εκπαίδευση γίνεται με χαμηλό ρυθμό μάθησης, έτσι ώστε να αποφευχθούν μεγάλες αλλαγές στα βάρη και να εξασφαλιστεί πως το μοντέλο δε θα χάσει τα βασικά μοτίβα που απέκτησε από την αρχική εκπαίδευση, αλλά ταυτόχρονα θα μπορεί να προσαρμόζεται στα νέα δεδομένα. Για αυτόν τον λόγο, έχει επιλεγεί και η εκπαίδευση του μοντέλου για μόλις 5 εποχές, καθώς θεωρούνται αρκετές για την προσαρμογή του στις νέες πληροφορίες και λίγες για την υπερπροσαρμογή του σε αυτές.
3. **Παραγωγή Μορίων:** Σε αυτό το σημείο, και αφού έχει ολοκληρωθεί η εκ νέου εκπαίδευση του μοντέλου, ξεκινάει η διαδικασία παραγωγής μορίων, η οποία είναι ακριβώς η ίδια με αυτή που ακολουθήθηκε για την γενικευμένη παραγωγή μορίων και περιγράφεται στην ενότητα 4.6.1.

4.8.2 Αποτελέσματα

Η παραπάνω διαδικασία πραγματοποιήθηκε σε 3 διαφορετικά υποσύνολα δεδομένων βιολογικών στόχων, με σκοπό την ανάδειξη της δυνατότητας του μοντέλου να προσαρμόζεται σε αυτά. Η τιμή της θερμοκρασίας επιλέχθηκε να είναι ίση με 0.5, καθώς στη συγκεκριμένη ενότητα δίνεται περισσότερη βάση στη δημιουργία έγκυρων μορίων, τα οποία βρίσκονται στον ίδιο χημικό χώρο με τα αρχικά τους δεδομένα.

- Thyroid Receptor:
 - Αριθμός Μορίων που χρησιμοποιήθηκαν για εκπαίδευση: 5148
 - Αριθμός Μορίων για Δειγματοληψία: 1000
 - Ποσοστό Έγκυρων SMILES: 71.3%
 - Ποσοστό Μοναδικών SMILES: 87.6%

Διάγραμμα UMAP:

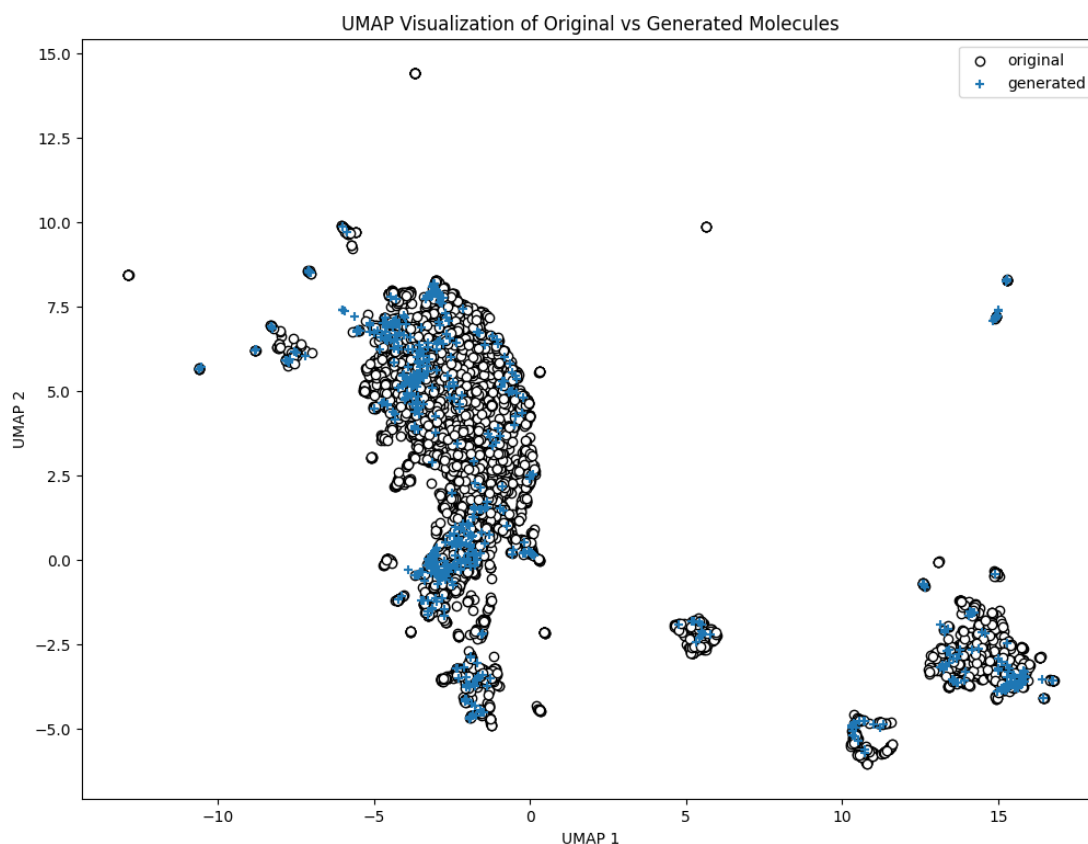


Figure 31: UMAP Analysis of TR ligands

Όπως φαίνεται παραπάνω, παρόλο που υπάρχουν αρκετά clusters που δείχνουν τις διαφορετικές ιδιότητες των δεδομένων, τα παραγόμενα μόρια καταφέρνουν να τις μιμηθούν και να μοιραστούν τον ίδιο χημικό χώρο. Αυτό φαίνεται ακόμα πιο έντονα στον παρακάτω πίνακα, όπου αντιπαραβάλλονται οι μέσες τιμές ορισμένων σημαντικών παραμέτρων, για τα αρχικά δεδομένα, καθώς και αυτά που προέκυψαν από τη δειγματοληψία μετά την περαιτέρω εκπαίδευση. Ακόμα, για να γίνει ακόμα πιο αισθητή η επιρροή των νέων δεδομένων στις ιδιότητες των παραγόμενων μορίων, έχει προστεθεί και μία τρίτη στήλη με τις τιμές που έχουν τα δεδομένα που δημιουργήθηκαν από το μοντέλο μετά την αρχική του εκπαίδευση, πριν δηλαδή χρησιμοποιηθεί οποιαδήποτε τεχνική μεταφοράς μάθησης.

Τιμές (mean)	TR (αρχικά)	TR (παραγόμενα)	Παραγόμενα χωρίς Μεταφορά Μάθησης (T = 0.5)
Μοριακό Βάρος	410.17	392.23	372.80
logP	3.94	4.02	3.40
QED Score	0.54	0.57	0.71
No. of H Donors	1.19	1.09	1.28

Τα αντίστοιχα δεδομένα παρατίθενται και για τους επόμενους δύο στόχους, επιβεβαιώνοντας τις αρχικές εκτιμήσεις.

- PPAR γ :
 - Αριθμός Μορίων που χρησιμοποιήθηκαν για εκπαίδευση: 6326
 - Αριθμός Μορίων για Δειγματοληψία: 1000
 - Ποσοστό Έγκυρων SMILES: 74.3%
 - Ποσοστό Μοναδικών SMILES: 81.64%

Διάγραμμα UMAP:

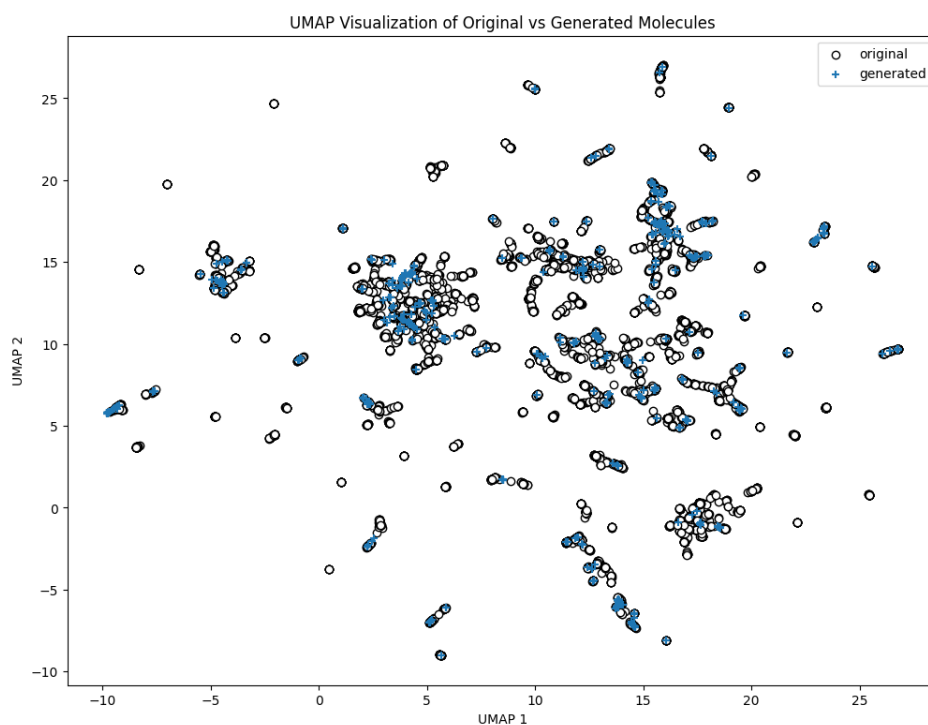


Figure 32: UMAP Analysis of PPAR γ ligands

Σε αυτή την εικόνα ο διαχωρισμός των δεδομένων εισόδου είναι ακόμα πιο έντονος, ωστόσο ακόμα και έτσι, το μοντέλο μπορεί να δημιουργήσει μόρια των οποίων ο χημικός χώρος συμπίπτει με των αρχικών.

Τιμές (mean)	PPAR γ (αρχικά)	PPAR γ (παραγόμενα)	Παραγόμενα χωρίς Μεταφορά Μάθησης (T = 0.5)
Μοριακό Βάρος	451.69	446.13	372.80
logP	5.29	5.57	3.40
QED Score	0.41	0.40	0.71
No. of H Donors	1.32	1.30	1.28

- LXRβ:
 - ο Αριθμός Μορίων που χρησιμοποιήθηκαν για εκπαίδευση: 1359
 - ο Αριθμός Μορίων για Δειγματοληψία: 1000
 - ο Ποσοστό Έγκυρων SMILES: 76.1%
 - ο Ποσοστό Μοναδικών SMILES: 58.13%

Διάγραμμα UMAP:

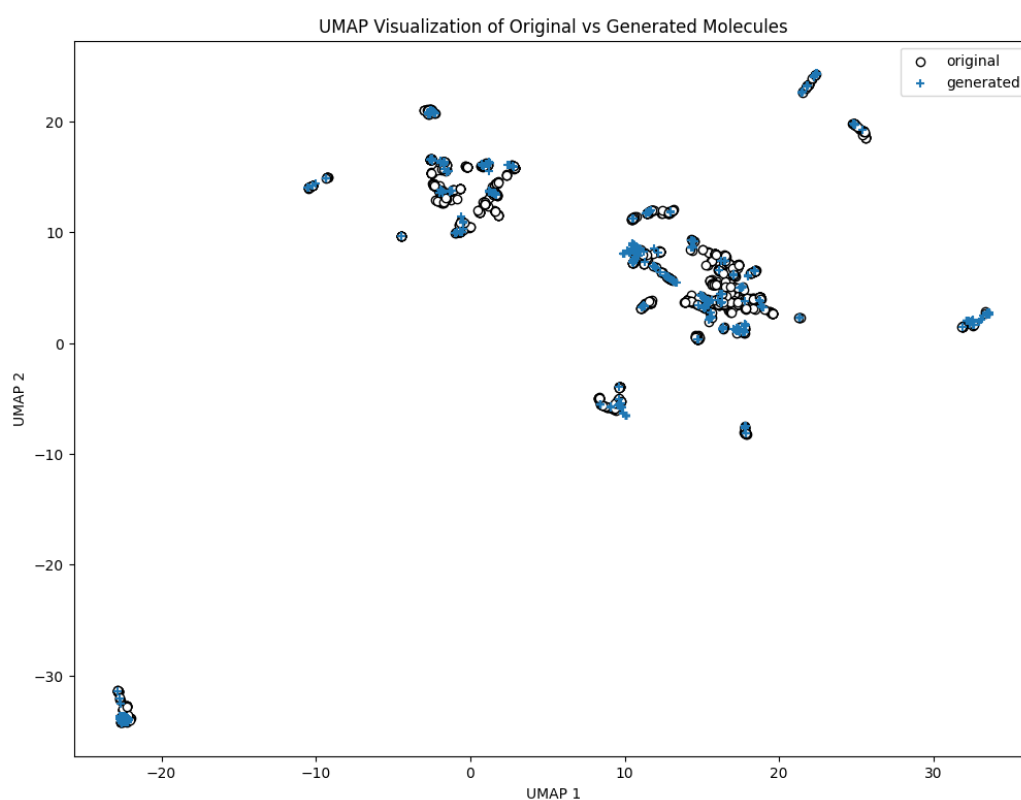


Figure 33: UMAP Analysis of LXRβ ligands

Για τους συγκεκριμένους βιολογικούς στόχους, παρατηρείται ένα αρκετά μειωμένο ποσοστό μοναδικών SMILES. Αυτό οφείλεται αρκετά στον μικρό αριθμό δεδομένων εξειδίκευσης, γεγονός που οδηγεί σε αύξηση της υπερπροσαρμογής πάνω σε αυτά.

Τιμές (mean)	LXRβ (αρχικά)	LXRβ (παραγόμενα)	Παραγόμενα χωρίς Μεταφορά Μάθησης (T = 0.5)
Μοριακό Βάρος	437.60	433.50	372.80
logP	5.14	5.52	3.40
QED Score	0.47	0.47	0.71
No. of H Donors	0.94	0.77	1.28

Κεφάλαιο 5: Συμπεράσματα και Μελλοντική Έρευνα

5.1 Συμπεράσματα

Το βασικότερο συμπέρασμα που βγήκε από την παρούσα διπλωματική, είναι η ικανότητα του LSTM μοντέλου να παράγει, με μεγάλη επιτυχία, χημικές δομές με πιθανές φαρμακολογικές ιδιότητες. Η διαδικασία παραγωγής νέων μορίων μέσω γενετικής προσέγγισης επέτρεψε τη διατήρηση βασικών, επιθυμητών για φαρμακευτικές εφαρμογές, δομικών χαρακτηριστικών, καθώς το μοντέλο κατάφερε να “μάθει” και να “καταλάβει” τις βασικές ιδιότητες και σχέσεις μεταξύ των μορίων που χρησιμοποιήθηκαν στην αρχική εκπαίδευση. Βέβαια, σημαντικό ρόλο έπαιξε η εύρεση και επιλογή των κατάλληλων παραμέτρων, είτε κατά τον ορισμό του μοντέλου, είτε κατά τη διαδικασία εκπαίδευσής του. Συγκεκριμένα, ήταν καταλυτική η χρήση δύο LSTM επιπέδων, επιλογή που έδωσε στο μοντέλο τη δυνατότητα να επεξεργάζεται πιο σύνθετα μοτίβα και εξαρτήσεις των ακολουθιών, ενώ ταυτόχρονα μείωσε τον κίνδυνο του vanishing gradient προβλήματος. Ακόμα, φάνηκε η σημαντική επιρροή της παραμέτρου θερμοκρασίας κατά τη διαδικασία δειγματοληψίας νέων ακολουθιών, ως προς την εγκυρότητα και την μοναδικότητά τους, δύο έννοιες που όπως διαπιστώθηκε, είναι αντιστρόφως ανάλογες.

Ένα ακόμα συμπέρασμα που προέκυψε από τα πειράματα που διεξήχθησαν είναι πως αφού το μοντέλο έχει μάθει να προβλέπει τους επόμενους χαρακτήρες μίας ακολουθίας τέτοιων, είναι επίσης ικανό να προβλέψει μόρια ξεκινώντας τη δειγματοληψία από μία ήδη υπάρχουσα συμβολοσειρά, και όχι από την αρχή. Αυτό φάνηκε με μεγάλη επιτυχία δημιουργώντας νέα μόρια, τα οποία περιείχαν όλα ένα συγκεκριμένο αρχικό τμήμα.

Τέλος, η δυνατότητα του μοντέλου να εστιάζει σε υποσύνολα δεδομένων βιολογικών στόχων μέσω μηχανικής μάθησης, επέτρεψε τη δημιουργία μορίων με συγκεκριμένες ιδιότητες, γεγονός που το καθιστά ένα πολύτιμο εργαλείο για τη στοχευμένη ανακάλυψη νέων ενώσεων.

5.2 Μελλοντική Έρευνα

Βάσει των συμπεραμάτων, αλλά και της συνολικής προσέγγισης της εργασίας, γίνονται σαφείς οι ικανότητες των μοντέλων βαθιάς μάθησης στον τομέα της ανακάλυψης φαρμάκων. Έτσι, δίνεται η δυνατότητα να ακολουθηθούν διάφορες κατευθύνσεις για να επεκταθεί η συγκεκριμένη προσέγγιση δίνοντας ακόμα πιο ενδιαφέροντα αποτελέσματα.

Μία από αυτές θα μπορούσε να είναι η επέκταση του αρχικού συνόλου δεδομένων με περισσότερα και διαφορετικού τύπου μόρια, σε μία προσπάθεια βελτίωσης της γενίκευσης του μοντέλου, και ταυτόχρονης μείωσης της όποιας υπερπροσαρμογής του.

Μία αρκετά ενδιαφέρουσα προσέγγιση θα ήταν και η χρήση ενισχυτικής μάθησης για ακόμα πιο συστηματική εξερεύνηση του χημικού χώρου, προσφέροντας τη δυνατότητα ανακάλυψης νέων μορίων που πληρούν συγκεκριμένα κριτήρια.

Τέλος, ο συνδυασμός των δύο τελευταίων μεθόδων που αναλύθηκαν στο πειραματικό μέρος, της παραγωγής μορίων από αρχικό τμήμα, καθώς και της χρήσης μεταφοράς μάθησης για παραγωγή μορίων, θα μπορούσε να ενισχύσει την ανακάλυψη νέων ουσιών, βασισμένων σε ήδη υπάρχοντα μόρια, με επιβεβαιωμένες φαρμακευτικές ιδιότητες.

References

- [1] Jiao, L., Shang, R., Liu, F., & Zhang, W. (2020). The models and structure of neural networks. , 47-79. <https://doi.org/10.1016/b978-0-12-819795-0.00002-5>.
- [2] Schmidhuber, J. (2014). Deep learning in neural networks: An overview. *Neural networks : the official journal of the International Neural Network Society*, 61, 85-117 . <https://doi.org/10.1016/j.neunet.2014.09.003>.
- [3] Khan, A., Sohail, A., Zahoor, U., & Qureshi, A. (2019). A survey of the recent architectures of deep convolutional neural networks. *Artificial Intelligence Review*, 53, 5455 - 5516. <https://doi.org/10.1007/s10462-020-09825-6>.
- [4] Parhi, R., & Nowak, R. (2019). The Role of Neural Network Activation Functions. *IEEE Signal Processing Letters*, 27, 1779-1783. <https://doi.org/10.1109/LSP.2020.3027517>.
- [5] Alkhoully, A., Mohammed, A., & Hefny, H. (2021). Improving the Performance of Deep Neural Networks Using Two Proposed Activation Functions. *IEEE Access*, 9, 82249-82271. <https://doi.org/10.1109/ACCESS.2021.3085855>.
- [6] Parhi, R., & Nowak, R. (2019). The Role of Neural Network Activation Functions. *IEEE Signal Processing Letters*, 27, 1779-1783. <https://doi.org/10.1109/LSP.2020.3027517>.
- [7] Gordon-Rodríguez, E., Loaiza-Ganem, G., Pleiss, G., & Cunningham, J. (2020). Uses and Abuses of the Cross-Entropy Loss: Case Studies in Modern Deep Learning. *ArXiv*, abs/2011.05231.
- [8] Soydaner, D. (2020). A Comparison of Optimization Algorithms for Deep Learning. *ArXiv*, abs/2007.14166. <https://doi.org/10.1142/S0218001420520138>.
- [9] Shulman, D. (2023). Optimization Methods in Deep Learning: A Comprehensive Overview. *ArXiv*, abs/2302.09566. <https://doi.org/10.48550/arXiv.2302.09566>.
- [10] Lillicrap, T., Santoro, A., Marris, L., Akerman, C., & Hinton, G. (2020). Backpropagation and the brain. *Nature Reviews Neuroscience*, 21, 335 - 346. <https://doi.org/10.1038/s41583-020-0277-3>.
- [11] Moradi, R., Berangi, R., & Minaei, B. (2019). A survey of regularization strategies for deep models. *Artificial Intelligence Review*, 53, 3947 - 3986. <https://doi.org/10.1007/s10462-019-09784-7>.
- [12] Srivastava, N., Hinton, G., Krizhevsky, A., Sutskever, I., & Salakhutdinov, R. (2014). Dropout: a simple way to prevent neural networks from overfitting. *J. Mach. Learn. Res.*, 15, 1929-1958. <https://doi.org/10.5555/2627435.2670313>.
- [13] Lewkowycz, A., & Gur-Ari, G. (2020). On the training dynamics of deep networks with L2 regularization. *ArXiv*, abs/2006.08643.

- [14] Paguada, S., Batina, L., Buhan, I., & Armendariz, I. (2021). Being Patient and Persistent: Optimizing An Early Stopping Strategy for Deep Learning in Profiled Attacks. *ArXiv*, abs/2111.14416.
<https://doi.org/10.1109/tc.2023.3234205>.
- [15] Bansal, M., Sharma, D., & Kathuria, D. (2022). A Systematic Review on Data Scarcity Problem in Deep Learning: Solution and Applications. *ACM Computing Surveys (CSUR)*, 54, 1 - 29. <https://doi.org/10.1145/3502287>.
- [16] Yu, Y., Si, X., Hu, C., & Zhang, J. (2019). A Review of Recurrent Neural Networks: LSTM Cells and Network Architectures. *Neural Computation*, 31, 1235-1270. https://doi.org/10.1162/neco_a_01199.
- [17] Lipton, Z. (2015). A Critical Review of Recurrent Neural Networks for Sequence Learning. *ArXiv*, abs/1506.00019.
- [18] Werbos, P. (1990). Backpropagation Through Time: What It Does and How to Do It. *Proc. IEEE*, 78, 1550-1560. <https://doi.org/10.1109/5.58337>.
- [19] Greff, K., Srivastava, R., Koutník, J., Steunebrink, B., & Schmidhuber, J. (2015). LSTM: A Search Space Odyssey. *IEEE Transactions on Neural Networks and Learning Systems*, 28, 2222-2232.
<https://doi.org/10.1109/tnnls.2016.2582924>.
- [20] Gers, F., Schmidhuber, J., & Cummins, F. (2000). Learning to Forget: Continual Prediction with LSTM. *Neural Computation*, 12, 2451-2471.
<https://doi.org/10.1162/089976600300015015>.
- [21] Prakash, N., & Devangi, P. (2021). Drug Discovery. *Encyclopedia of Molecular Pharmacology*. <https://doi.org/10.4172/jaa.1000025>.
- [22] Chopra I. (1997). Approaches to antibacterial drug discovery. *Expert opinion on investigational drugs*, 6(8), 1019–1024.
<https://doi.org/10.1517/13543784.6.8.1019>
- [23] Lavecchia, A., & Giovanni, C. (2013). Virtual screening strategies in drug discovery: a critical review.. *Current medicinal chemistry*, 20 23, 2839-60 .
<https://doi.org/10.2174/09298673113209990001>.
- [24] Batool, M., Ahmad, B., & Choi, S. (2019). A Structure-Based Drug Discovery Paradigm. *International Journal of Molecular Sciences*, 20.
<https://doi.org/10.3390/ijms20112783>.
- [25] Verma, S., & Prabhakar, Y. (2015). Target based drug design - a reality in virtual sphere.. *Current medicinal chemistry*, 22 13, 1603-30 .
<https://doi.org/10.2174/0929867322666150209151209>.
- [26] Wang, C., Xu, P., Zhang, L., Huang, J., Zhu, K., & Luo, C. (2018). Current Strategies and Applications for Precision Drug Design. *Frontiers in Pharmacology*, 9. <https://doi.org/10.3389/fphar.2018.00787>.
- [27] Gupta, A., Müller, A., Huisman, B., Fuchs, J., Schneider, P., & Schneider, G. (2017). Generative Recurrent Networks for De Novo Drug Design. *Molecular Informatics*, 37. <https://doi.org/10.1002/minf.201700111>.
- [28] Jiménez-Luna, J., Grisoni, F., & Schneider, G. (2020). Drug discovery with explainable artificial intelligence. *Nature Machine Intelligence*, 2, 573 - 584.
<https://doi.org/10.1038/s42256-020-00236-4>.

- [29] Koutroumpa, N., Papavasileiou, K., Papadiamantis, A., Melagraki, G., & Afantitis, A. (2023). A Systematic Review of Deep Learning Methodologies Used in the Drug Discovery Process with Emphasis on In Vivo Validation. *International Journal of Molecular Sciences*, 24. <https://doi.org/10.3390/ijms24076573>.
- [30] Kim, J., Park, S., Min, D., & Kim, W. (2021). Comprehensive Survey of Recent Drug Discovery Using Deep Learning. *International Journal of Molecular Sciences*, 22. <https://doi.org/10.3390/ijms22189983>.
- [31] Lipinski, C. A., Lombardo, F., Dominy, B. W., & Feeney, P. J. (2001). Experimental and computational approaches to estimate solubility and permeability in drug discovery and development settings. *Advanced drug delivery reviews*, 46(1-3), 3–26. [https://doi.org/10.1016/s0169-409x\(00\)00129-0](https://doi.org/10.1016/s0169-409x(00)00129-0)
- [32] Pollastri, M. (2010). Overview on the Rule of Five. *Current Protocols in Pharmacology*, 49. <https://doi.org/10.1002/0471141755.ph0912s49>.
- [33] Bickerton, G., Paolini, G., Besnard, J., Muresan, S., & Hopkins, A. (2012). Quantifying the chemical beauty of drugs.. *Nature chemistry*, 4 2, 90-8. <https://doi.org/10.1038/nchem.1243>.
- [34] Derringer, G., & Suich, R. (1980). Simultaneous Optimization of Several Response Variables. *Journal of Quality Technology*, 12(4), 214–219. <https://doi.org/10.1080/00224065.1980.11980968>
- [35] Abadi, M., Barham, P., Chen, J., Chen, Z., Davis, A., Dean, J., Devin, M., Ghemawat, S., Irving, G., Isard, M., Kudlur, M., Levenberg, J., Monga, R., Moore, S., Murray, D., Steiner, B., Tucker, P., Vasudevan, V., Warden, P., Wicke, M., Yu, Y., & Zhang, X. (2016). TensorFlow: A system for large-scale machine learning. *ArXiv*, abs/1605.08695. <https://doi.org/10.48550/arXiv.1605.08695>
- [36] Weininger, D. (1988). SMILES, a chemical language and information system. 1. Introduction to methodology and encoding rules. *Journal of Chemical Information and Modeling*, 28(1), 31–36. <https://doi.org/10.1021/ci00057a005>
- [37] He, Y., Zhang, X., Ao, W., & Huang, J. (2018). Determining the optimal temperature parameter for Softmax function in reinforcement learning. *Appl. Soft Comput.*, 70, 80-85. <https://doi.org/10.1016/j.asoc.2018.05.012>.