

Evaluation of Machine Learning Models To Predict Unbound Brain: Plasma Partitioning Coefficient In CNS-Related Discovery Chemistry

Hyeonmin Cho^{1*}, Myoungwoo Kim²

¹Wissahickon High School, Ambler, PA, USA

²Blue Bell Envirobotics Lab, Ambler, Pennsylvania, USA

*Corresponding Author: blueskyhmcho@gmail.com

Advisor: Myoungwoo Kim, kimmohio@gmail.com

Received January 17, 2026; Revised June 23, 2026; Accepted July 20, 2026

Abstract

The prediction of unbound brain to plasma partitioning coefficient ($K_{p,uu}$) is highly crucial in the CNS drug discovery, as this particular parameter determines how effectively and quickly the compounds penetrate the blood-brain barrier. In this study, Random Forest (RF), Classification and Regression Tree (CART), and Extreme Gradient Boosting (XGBoost) machine learning models were evaluated to predict the $K_{p,uu}$ values using physicochemical and molecular descriptors. XGBoost has achieved the highest predictive performance with correlation coefficient R^2 of 0.9875 with 9 features selected, followed by CART with 0.9808 with 8 features, and RF with 0.8232 with 9 features. The superior handling of XGBoost suggests the robustness of handling descriptor interactions that influence brain penetration. By interpreting these correlations, compounds predicted with higher $K_{p,uu}$ are expected to penetrate the BBB more rapidly and reach faster equilibrium between the unbound plasma and brain concentrations due to three characteristics of CNS, shape-related, electronic and polarity-related, and hydrogen profile-related descriptors. This overall implies a greater exposure potential, critical for CNS-active drugs. The study also demonstrates that ensemble-based machine learning methods such as XGBoost outperform traditional tree-based approaches in modeling the unbound brain partitioning, providing a valuable computational framework for the early-stage of CNS drug design and optimization.

Keywords: Blood brain barrier, Partitioning coefficients, Machine learning

1. Introduction

In earlier studies, animal study was a prominent factor as a determinant for specifically barriers between different body parts (Di et al., 2013; Pardridge, 2005). As accurate as the data was, there were significant drawbacks, including: time and cost (Kola & Landis, 2004). With this factor in mind, the use of *In silico* modeling was popularized, aiming to reduce the continuation of animal studies and labeling it as an unethical approach, reinforced by NAMs (EPA, 2021; FDA, 2021; EFSA, 2023, ECHA, 2023). However, this emerged another research gap due to the lack of modern technology such as machine learning that can improve statistical strength. Many frequent studies using the *in silico* model have performed not an adequate amount of accuracy (Morales et al., 2017; Martins et al., 2012). This research, therefore, focuses on specifically the brain and plasma barrier for pharmacokinetic drugs to support urgent patients via discovery chemistry in both cost and time.

The brain and plasma barrier, or more exactly the permeability coefficient ($K_{p,uu}$), is a significant obstacle in the distribution of central nervous system (CNS)-related pharmacokinetic drugs, which are prescribed compounds designated for bodily function by triggering or inhibiting cellular receptors (Hammarlund-Udenaes et al., 2008; Smith et al., 2010). Several existing studies address this challenge of getting drugs to penetrate the $K_{p,uu}$ (Zhang et al., 2016; Fridén et al., 2009). Because there are limited databases on recorded $K_{p,uu}$ values for chemicals, researchers documented the use of machine learning (ML) models and artificial intelligence (AI) to predict $K_{p,uu}$ values of drug

compounds (Svetnik et al., 2003; Sheridan et al., 2016; Chen & Guestrin, 2016). This study uses these aforementioned discoveries and papers to guide the research, following closely with the experiment done by Lawrence et al. (2023).

The goal of this research is to improve the overall accuracy of *in silico* model to better predict chemical compounds that pass through the brain and plasma barrier quickly to better support discovery chemistry to aid urgent care patients who may be dealing with drug issues. By incorporating State-of-the-art technologies such as Sequential Feature Selection approach (Pearce et al., 2017; Wambaugh et al., 2015) coupled with CART, Random Forest, and XGBoost models, the research ensures a greater accuracy and dependency.

2. Methods

2.1 Data Mining and Quality Control to Form Training Data

From the full HTTK dataset, which includes pharmacokinetic parameters across multiple tissues, only brain-related data were extracted to focus on predicting unbound brain-to-plasma partitioning coefficients ($K_{p,uu}$). The raw database contained measurements for various biological organs, including liver, kidney, and adipose tissues were excluded during preprocessing. The missing or unreported values were flagged as “NS” (Not Specified) to ensure that the data-cleaning code could recognize and treat them separately from true zero or null entries, to prevent bias in model training. The final curated dataset contained approximately 93 compounds and 1,444 molecular descriptors, each representing an individual compound.

For model validation, train/test split was set to 0.8 for training and 0.2 for testing dataset with random size of 32. These testing data were used to validate the prediction model without using further observation dataset for model validation.

To find the unbound brain partitioning coefficient ($K_{p,uu}$), the fraction unbound in plasma (f_u) was multiplied by the experimental brain-to-plasma partition coefficient (K_p). In addition, $K_p = C_{\text{brain}}/C_{\text{plasma}}$, and $K_{p,uu} = K_p \times (f_{u,b}/f_{u,p})$ represents the finding of $K_{p,uu}$. This correction accounts for protein binding and enzymatic interactions in plasma that influence unbound drug availability near the blood–brain interface, providing a more relevant endpoint for machine learning model training.

2.2 Molecular Descriptor Calculation

Molecular descriptors were generated using the PaDEL-Descriptor software ver. 2 (Yap, 2010), an open-source Java-based cheminformatics tool that computes a wide range of structural and physicochemical parameters. The dataset consisted of 93 compounds, including chemicals from industrial, cosmetic, and agricultural applications, selected from the EPA HTTK database to represent diverse chemical classes to central nervous system (CNS) exposure.

For each compound, PaDEL calculated 1,444 two-dimensional (2D) molecular descriptors, encompassing topological and electronic properties. Each molecule was identified by its chemical name and SMILES (Simplified Molecular Input Line Entry System) representation, which served as the input format for descriptor computation.

The full descriptor process was highly efficient, requiring approximately five seconds to compute all descriptors for the 93 molecules. The resulting output was stored in comma-separated values (CSV) format and merged with the experimental $K_{p,uu}$ values derived from the EPA HTTK database for machine learning model development (EPA, 2023).

2.3 Sequential Feature Selection

To improve model robustness and interpretability, Sequential Feature Selection (SFS) was used as a wrapper-based feature selection technique that evaluates features in combination rather than individually. SFS works by iteratively adding the most informative descriptors to a subset based on model performance, choosing the next descriptor in context with those already selected, which better captures interactions that simple regression ranking overlooks. This greedy, sequential decision process was originally formalized to address challenges where redundant data can degrade classification performance, assessing descriptors step by step to reduce unnecessary complexity and

emphasize meaningful predictor groups that drive learning outcomes. In these $K_{p,uu}$ prediction models, SFS helped isolate compact, high-impact sets of physicochemical and molecular descriptors, thereby enhancing predictive accuracy while reducing overfitting and improving interpretability of feature contributions to blood–brain barrier penetration.

2.4 Machine Learning

Classification and Decision Tree (CART)

CART is a single decision-tree model that partitions data into branches based on feature values that best reduce the prediction error. It is easy to visualize and interpret, making it useful for identifying general trends between molecular properties and outcomes. However, because CART relies on a single decision, it tends to overfit training data and has lower predictive power on unseen data compared to ensemble models. In this study, CART served as a baseline model for comparison against more advanced ensemble techniques.

Random Forest (RF)

Random Forest is an ensemble learning method that constructs a multitude of individual decision trees and averages their results to improve predictive performance. Each tree is trained on a random subset of data and descriptors, known as bootstrap aggregation. Such randomization helps reduce overfitting and reduces variance while preserving high accuracy. RF works exceptionally well for nonlinear and high-dimensional datasets, such as molecular descriptors, where interactions between features are complex. It also provides measures of feature importance that help interpret which descriptors most influence the prediction of $K_{p,uu}$ made by the model.

XGBoost

XGBoost is a powerful gradient-boosted tree algorithm that builds decision trees sequentially, with each new tree correcting the errors of previous ones. It optimizes model performance through gradient descent and regularization, making it both fast and resistant to overfitting. XGBoost efficiently captures subtle patterns and nonlinear relationships among molecular descriptors, often outperforming traditional methods. Its combination of speed, accuracy, and scalability makes it a leading tool in pharmacokinetic modeling.

2.5 Model Training, Evaluation, and Validation

No k-fold cross-validation was employed; instead, a single hold-out test set (20% of data) was used for independent model evaluation. This approach was chosen to maximize training data availability given the small dataset size (93 compounds) and to ensure reproducibility with a fixed random seed (32). Hyperparameters for all three models (CART, RF, XGBoost) were set to scikit-learn default values to maintain consistency and reduce computational complexity. Model performance was evaluated using two complementary metrics: R^2 (coefficient of determination), which measures the proportion of variance explained, and AAFE (Average Absolute Fold Error), defined as Equation (1):

$$AAFE = 10^{(\text{mean}(\text{abs}(\log_{10}(\text{predicted}/\text{observed}))))} \quad (1)$$

AAFE quantifies average multiplicative error on a fold-ratio scale. Both metrics were calculated on the held-out test set only.

2.6 Graphical analysis

Important Feature Analysis

The Feature importance bar graph illustrates which molecular descriptors most strongly influence the model prediction of $K_{p,uu}$. The x axis value shows the level of importance and contribution of each descriptor to the model's predictive performance. The y axis ranks the names of the top molecular descriptors by their importance score.

Correlation heatmap

The correlation heatmap shows a visual summary of how strongly molecular descriptors are related to each other and to the target variable, $K_{p,uu}$. It helps identify patterns of a linear relationship that influence blood-brain barrier penetration. The x and y axes show the correlation coefficient relationship, from -1 to +1, of molecular descriptors to one another as well as to the target $K_{p,uu}$. A +1 indicates a strong positive relationship while a -1 indicates an inverse relationship to one another. The 0 shows no correlation. For instance, $K_{p,uu}$ and Lf showed a correlation of -0.05, showing a very slight inverse relationship but having almost no correlation at all.

3. Results

3.1 Sequential Feature Selection

The performance of the evaluated machine learning models for predicting the unbound brain:plasma partitioning coefficient ($K_{p,uu}$) is summarized in Table 1. Each algorithm reached its optimal predictive configuration using 8–9 physicochemical and molecular descriptors, selected based on their ability to maximize the correlation between predicted and experimental $K_{p,uu}$ values. XGBoost achieved the highest correlation coefficient with $R^2 = 0.9875$ with nine descriptors, CART followed with $R^2 = 0.9808$ with eight descriptors, and RF with $R^2 = 0.8232$ with nine descriptors reflecting consistent performance in partitioning feature space despite its simpler structure. Although RF achieved an R^2 of 0.8232, its AAFE value of 548.90 is anomalously high compared to CART (1.83) and XGBoost (2.16), indicating that the model produced extreme mispredictions on outlier compounds. See Discussion for explanation.

Table 1. Best feature model and its performance from each ML method

ML Models	No. of Features Selected	Selected Descriptors and regression equation	R^2	AAFE
CART	8	Selected Features: nAcid, nArAt, nArBd, CSP1, SMDT, nHBA, nALAC, LF Regression Equation: $(0.0041 * \text{nAcid}) + (0.0002 * \text{nArAt}) + (0.2587 * \text{nArBd}) + (0.0000 * \text{CSP1}) + (0.4598 * \text{SMDT}) + (0.1446 * \text{nHBA}) + (0.1269 * \text{nALAC}) + (0.0057 * \text{LF}) + \text{Bias}$	0.9808	1.83
RF	9	Selected Features: nAcid, naAromAtom, nAromBond, ATS0m, BCUTw-11, C1SP1, IC0, nAtomLAC, 'LipinskiFailures' Regression Equation: NA*	0.8232	548.90
XGBOOST	9	Selected Features: nAcid, naAromAtom, nAromBond, C1SP1, SCH-3, FMF, IC0, nAtomLAC, MDEC-1, LipinskiFailures Regression Equation: $(0.0596 * \text{nAcid}) + (0.2114 * \text{nArAt}) + (0.0000 * \text{nArBd}) + (0.0002 * \text{CSP1}) + (0.1548 * \text{SCH}) + (0.1706 * \text{FMF}) + (0.1737 * \text{IC0}) + (0.1909 * \text{nALAC}) + (0.0387 * \text{LF}) + \text{Bias}$	0.9875	2.16

*NA: Not Applicable. No equation was found due to the bagging process that missed global minimum of bias during the bootstrapping.

The descriptors selected for the final models are listed in Table 2 and represent the physicochemical and structural features most strongly associated with predicting $K_{p,uu}$. These include functional group counts such as nHBA and nArAt, topological indices reflecting molecular size and shape, and charge-related or electronic descriptors that influence polarity and ionization behavior. Their consistent selection across XGBoost, CART, and Random Forest indicates that these properties collectively capture the key determinants of unbound brain penetration and are essential for constructing reliable predictive models in CNS drug discovery.

Table 2. Explanation of the molecular descriptors selected in SFS.

Descriptors (acronym) in PaDEL database	Explanation	Descriptors (acronym) in PaDEL database	Explanation
nAcid (nAcid)	Number of acidic groups.	AtS0m (nBase)	Broto-Moreau autocorrelation - lag 0 / weighted by mass
naAromAtom (nArAt)	Number of aromatic atoms	BCUTw-11 (nBonds)	The lowest atom weighted BCUTS

nAromBond (nArBd)	Number of aromatic bonds	MDEC-1 (MLFER_A)	Molecular distance edge between all primary carbons
C1SP1 (CSP1)	Triply bound carbon bound to one other carbon	IC0 (IC0)	Information content index (neighborhood symmetry of 0-order)
SCH-3 (SC-3)	Simple chain, order 3	nAtomLAC (nALAC)	Number of atoms in the longest aliphatic chain
FMF (fragC)	Complexity of a molecule	LipinskiFailures (LF)	Number failures of the Lipinski's Rule Of 5
SpMax_Dt (SMDT)	Leading eigenvalue from detour matrix		

3.2 Machine Learning

Figure 1 shows the feature importance distributions for the (a) CART, (b) Random Forest, and (c) XGBoost models, highlighting clear differences in how each algorithm utilizes molecular descriptors. CART places disproportionate weight on a single descriptor, SMDT (Leading eigenvalue of detour matrix, ~ 0.47), indicating reliance on molecular size and branching. RF model (b) also emphasizes one dominant feature, ATS0m (Broto-Moreau autocorrelation, ~ 0.47), though it distributes the remaining importance more evenly across descriptors such as IC0 (~ 0.23). In contrast, XGBoost (c) demonstrates the most balanced feature usage, with meaningful contributions from nArAt (number of aromatic atoms, ~ 0.22), nALAC (longest aliphatic chain, ~ 0.19), and IC0 (information content, ~ 0.17). This broader reliance on multiple descriptors reflects XGBoost's ability to capture more complex relationships governing $K_{p,uu}$ and contributes to its superior predictive performance.

The correlation heatmaps in Figure 2 illustrate the relationships among the key descriptors used in the CART, RF, and XGBoost models, revealing distinct structural patterns within each feature set. In the CART model, a strong positive correlation between SMDT and nHBA ($r = 0.74$) indicates that molecular size and hydrogen-bonding capacity tend to co-vary in shaping its predictions. The XGBoost heatmap highlights more nuanced physicochemical relationships, including a strong inverse correlation between FMF and nALAC (Pearson correlation coefficient, $r = -0.61$), suggesting a trade-off between molecular complexity and aliphatic chain length. Additionally, the positive correlation between nArBd and FMF ($r = 0.58$) reflects the contribution of aromaticity to overall molecular complexity. Together, these patterns underscore the differing descriptor interactions each model captures when predicting $K_{p,uu}$.

4. Discussion

Across the three models, XGBoost provided the strongest and most stable performance, and the patterns observed are broadly consistent with prior QSAR literature. The finding that gradient boosting outperformed a single tree and a bagged ensemble agrees with Sheridan et al. (2016), who reported that extreme gradient boosting is competitive with or superior to Random Forest for QSAR tasks, and with the general effectiveness of tree ensembles for chemical-

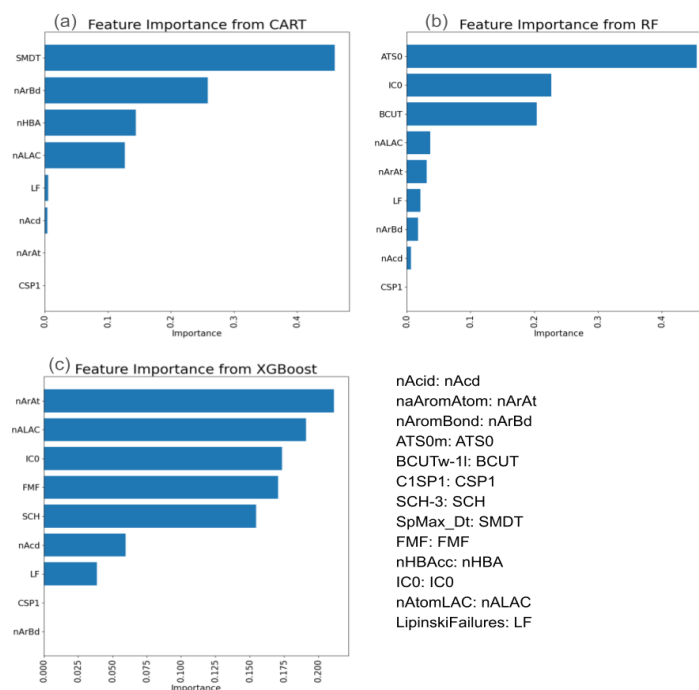


Figure 1. Feature importance graphs from (a) CART, (b) RF, and (c) XGBoost.

property precision reported by Svetnik et al. (2003). The descriptor classes identified here – aromaticity hydrogen-bonding capacity, molecular size, and drug likeness – echo the physicochemical determinants of BBB penetration emphasized in earlier Kp,uu and BBB modeling work (Friden et al., 2009; Zhang et al., 2016; Morales et al., 2017; Martins et al., 2012), supporting the view that Kp,uu is governed by an interplay of several descriptors rather than a single dominant factor.

The contrast between RF's moderate R2 (0.8232) and its very high AAFE (548.90) reflects the different properties of the two metrics. R2 measures the proportion of variance explained and is dominated by the bulk of well-predicted compounds, so a model can retain a moderate R2 even when a few compounds are poorly predicted. AAFE, by contrast, averages the absolute fold error in log space and is therefore extremely sensitive to a small number of large multiplicative mispredictions. A handful of outlier compounds predicted several orders of magnitude away from their observed values can inflate AAFE dramatically while leaving R2 relatively intact. This explains why RF, which appears acceptable by R2, is the least reliable model by AAFE.

The descriptors selected by SFS can be grouped by how they affect transport across the BBB. Structural and shape descriptors (SMDT, nALAC, SCH-3) describe molecular size, flexibility, and compactness; smaller, moderately flexible molecules tend to cross the BBB more readily. Aromaticity descriptors (nArAt, nArBd) increase lipophilicity and planarity, improving partitioning into membrane lipids. Electronic and polarity descriptors (nHBA, nAcid) govern hydrogen-bonding and ionization; molecules forming fewer hydrogen bonds, and that are less acidic generally show higher Kp,uu because they incur a smaller desolvation penalty when crossing the membrane. Drug-likeness summaries such as Lipinski Failures (LF) integrate these effects, with compounds satisfying Lipinski's Rule of Five tending to penetrate the brain more effectively. The strong XGBoost performance suggests that gradient boosting captures these interacting effects well and, beyond prediction, offers interpretable guidance for CNS drug design; coupling this approach with physiologically based pharmacokinetic (PBPK) modeling could further improve temporal prediction of brain exposure.

For medicinal chemists and drug discovery teams, these findings suggest prioritizing compounds with lower molecular complexity (lower FMF) and moderately extended aliphatic chains (higher nALAC) when designing blood-brain barrier-penetrating compounds. The XGBoost model's superior performance, combined with its balanced descriptor usage, makes it the recommended computational tool for rapid *in silico* screening of BBB permeability in early-stage CNS drug discovery pipelines. However, final validation should incorporate physiologically based pharmacokinetic (PBPK) modeling and complementary *in vitro* and *in vivo* studies to confirm predictions.

Several limitations should be acknowledged. The small dataset (93 compounds) restricts coverage of chemical space and may limit generalizability. Data-quality issues – extreme outliers and skewed descriptor distributions dominated by zero values – contributed to error amplification and to the unusually high AAFE observed for RF. Future work should expand the dataset, incorporate three-dimensional descriptors, and apply rigorous validation.

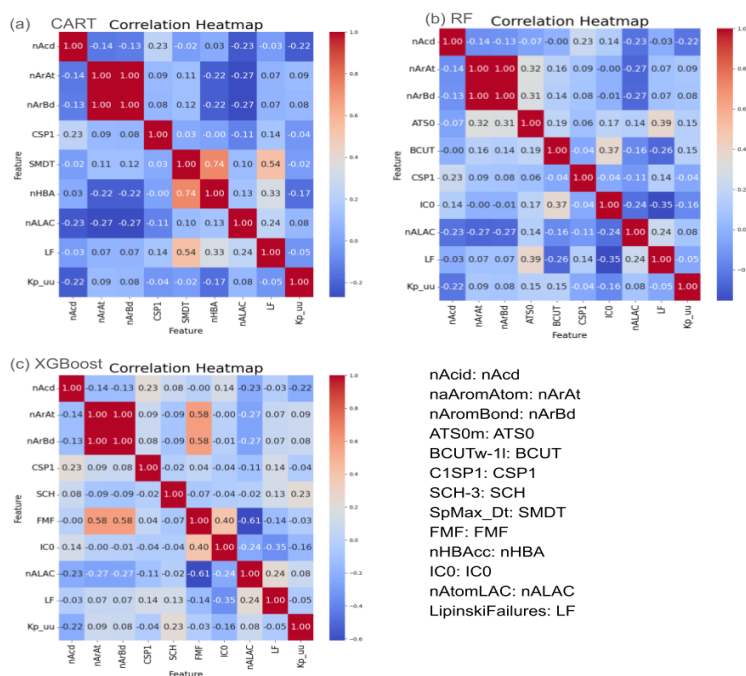


Figure 2. Correlation heatmaps from (a) CART, (b) RF, and (C) XGBoost.

5. Conclusion

Among the compared models, XGBoost yielded the highest predictive accuracy ($R^2 = 0.9875$) in modeling the unbound brain-to-plasma partitioning coefficient, $K_{p,uu}$, outperforming CART ($R^2 = 0.9808$) and RF ($R^2 = 0.8232$), therefore demonstrating the superior ability to capture nonlinear descriptor interactions governing BBB permeability. The most influential descriptors - nHBA, LF, nALAC, IC0, and FMF- represent hydrogen-bonding capacity, drug-likeness, aliphatic chain length, information content, and molecular complexity, respectively.

Correction heatmap analysis showed that FMF and nALAC were inversely related ($r = -0.61$) while nArBd and FMF were positively related ($r = 0.58$), underscoring the opposing roles of molecular size and aromaticity in modulating permeability. Despite data irregularities such as outliers and excess zero values that contributed to the elevated AAEF, XGBoost provided robust and interpretable predictions.

Overall, the study supports XGBoost as a useful framework for *in silico* $K_{p,uu}$ prediction in high-throughput virtual screening and the rational design of CNS-active compounds, provided it's applied with appropriate caution.

Acknowledgment

Authors appreciate EPA's HHTK database for the experimental partitioning coefficients between tissue and plasma that has been continuously updated since 2006.

References

- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785–794. <https://doi.org/10.1145/2939672.2939785>
- Di, L., Rong, H., & Feng, B. (2013). Demystifying brain penetration in central nervous system drug discovery. *Journal of Medicinal Chemistry*, 56(1), 2–12. <https://doi.org/10.1021/jm301297f>
- European Chemicals Agency. (2023). *New approach methodologies in regulatory science*. ECHA. <https://echa.europa.eu/>
- European Food Safety Authority. (2023). *Guidance on new approach methodologies*. EFSA. <https://www.efsa.europa.eu/>
- Fridén, M., et al. (2009). Structure-brain exposure relationships in rat and human using a novel data set of unbound drug concentrations in brain interstitial and cerebrospinal fluids. *Journal of Medicinal Chemistry*, 52(20), 6233–6243. <https://doi.org/10.1021/jm901036q>
- Hammarlund-Udenaes, M., et al. (2008). On the rate and extent of drug delivery to the brain. *Pharmaceutical Research*, 25(8), 1737–1750. <https://doi.org/10.1007/s11095-007-9502-2>
- Kola, I., & Landis, J. (2004). Can the pharmaceutical industry reduce attrition rates? *Nature Reviews Drug Discovery*, 3(8), 711–716. <https://doi.org/10.1038/nrd1470>
- Lawrenz, M., et al. (2023). A computational physics-based approach to predict unbound brain-to-plasma partition coefficient, $K_{p,uu}$. *Journal of Chemical Information and Modeling*, 63(12), 3786–3798. <https://doi.org/10.1021/acs.jcim.3c00150>
- Martins, I. F., et al. (2012). A Bayesian approach to *in silico* blood-brain barrier penetration modeling. *Journal of Chemical Information and Modeling*, 52(6), 1686–1697. <https://doi.org/10.1021/ci300124c>
- Morales, J. F., et al (2017). Current state and future perspectives in QSAR models to predict blood-brain barrier penetration in central nervous system drug R&D. *Mini-Reviews in Medicinal Chemistry*, 17(3), 247–257. <https://doi.org/10.2174/1389557516666161013110813>

- Pardridge, W. M. (2005). The blood-brain barrier: Bottleneck in brain drug development. *NeuroRx*, 2(1), 3–14. <https://doi.org/10.1602/neurorx.2.1.3>
- Pearce, R. G., et al. (2017). htk: R package for high-throughput toxicokinetics. *Journal of Statistical Software*, 79(4), 1–26. <https://doi.org/10.18637/jss.v079.i04>
- Rückstieß, T., Osendorfer, C., & van der Smagt, P. (2011). *Sequential feature selection for classification*. In D. Wang & M. Reynolds (Eds.), *AI 2011: Advances in Artificial Intelligence* (pp. 132–141). Lecture Notes in Computer Science (Vol. 7106). Springer, Berlin Heidelberg. https://doi.org/10.1007/978-3-642-25832-9_14
- Sheridan, R. P., et al. (2016). Extreme gradient boosting as a method for quantitative structure-activity relationships. *Journal of Chemical Information and Modeling*, 56(12), 2353–2360. <https://doi.org/10.1021/acs.jcim.6b00591>
- Smith, D. A., Di, L., & Kerns, E. H. (2010). The effect of plasma protein binding on *in vivo* efficacy: Misconceptions in drug discovery. *Nature Reviews Drug Discovery*, 9(12), 929–939. <https://doi.org/10.1038/nrd3287>
- Svetnik, V., et al. (2003). Random forest: A classification and regression tool for compound classification and QSAR modeling. *Journal of Chemical Information and Computer Sciences*, 43(6), 1947–1958. <https://doi.org/10.1021/ci034160g>
- U.S. Environmental Protection Agency. (2020, June). *New approach methods work plan*. Safer Chemicals Research. <https://www.epa.gov/chemical-research/new-approach-methods-work-plan>
- U.S. Food and Drug Administration. (2025, April). *FDA announces plan to phase out animal testing requirement for monoclonal antibodies and other drugs*. Press Announcements. <https://www.fda.gov/news-events/press-announcements/>
- Wambaugh, J. F., et al. (2015). Toxicokinetic triage for environmental chemicals. *Toxicological Sciences*, 147(1), 55–67. <https://doi.org/10.1093/toxsci/kfv118>
- Yap, C. W. (2011). PaDEL-descriptor: An open source software to calculate molecular descriptors and fingerprints. *Journal of Computational Chemistry*, 32(7), 1466–1474. <https://doi.org/10.1002/jcc.21707>
- Zhang, Y. Y., et al. (2016). Integrating *in silico* and *in vitro* approaches to predict drug accessibility to the central nervous system. *Molecular Pharmaceutics*, 13(5), 1540–1550. <https://doi.org/10.1021/acs.molpharmaceut.6b00031>