

A machine learning-based framework for predicting candidate drug side effects from biological networks

Pietro Cinaglia¹, Giulia Pingitore¹, Marianna Milano², and Mario Cannataro³

¹ Department of Health Sciences, Magna Graecia University of Catanzaro, Italy

² Department of Experimental and Clinical Medicine, Magna Graecia University of Catanzaro, Italy

³ Department of Medical and Surgical Sciences, Magna Graecia University of Catanzaro, Italy.

cinaglia@unicz.it, giulia.pingitore@studenti.unicz.it, m.milano@unicz.it, cannataro@unicz.it

Abstract. Artificial Intelligence (AI) has emerged as a powerful and effective tool with several applications in health science. An inherent drawback of drug therapies is the potential for side effects, which are adverse reactions that negatively impact human health. In recent years, AI has been applied in pharmacology and pharmacovigilance, e.g., for studying and analysing drug side effects. Likewise, network science has become widely used as an effective and efficient tool for modelling interactions between biological objects. In this paper, we presented a framework for predicting candidate drug side effects by using Machine Learning (ML) techniques applied to biological multilayer networks. Experimentation supports the application of the ML-based models implemented in the proposed framework for predicting novel (candidate) drug side effects from biological multilayer networks.

Keywords: drug side effect · machine learning · artificial intelligence · bioinformatics · pharmacology.

1 Introduction

Drug side effects represent adverse reactions (i.e., unwanted undesirable effects) having a negative impact on human health, that usually emerge in clinical trials or clinical practice and need to be investigated thoroughly.

Artificial Intelligence (AI) has proven to be a valid and effective tool with multiple applications in medicine, biology, pharmacology, and more generally in health sciences [1, 5]. In recent years, the interest in this technology has grown enormously and rapidly, especially due to the ever-increasing availability of multi-modal data, e.g., from genomic, economic, demographic, clinical, and phenotypic studies [18]. Additionally, it has been furthered both by increased computing power, and the development of novel analysis methods that increasingly support parallelization and cloud computing [12].

In this context, bioinformatics played a crucial role, by supporting the application of AI for several purposes in many areas of health sciences, including pharmacology [3, 17, 19]. The vast majority of methods presented and discussed in the scientific literature, which to date also represent the state of the art in the use of AI in bioinformatics, can be associated with Machine Learning (ML) and Deep Learning (DL). The meaning of the latter is often used interchangeably by non-experts but is instead characterized by notable differences: ML is explicitly used as a means to extract knowledge from data by formulating predictions based on patterns processed in its training data; DL algorithms are ML ones which performs a deeper data processing based on Neural Networks (NN) [6], instead of more traditional linear regression or a decision tree.

Currently, AI is applied in supporting drug development, prediction, and knowledge mining, establishing itself as a parallel line of study which mainly concerns the analysis and interpretation of data from heterogenous sources that need to be collected, managed and analysed through methods and instruments of computer science [26]. It is effectively applied in pharmacovigilance for investigating adverse drug reactions [32], as well as for surveillance and signal detection, classifying individual case safety reports to adverse event profiles, extracting Drug-Drug Interactions (DDI), identifying high-risk populations for drug toxicity, predicting drug side effects, and simulating clinical trials [20, 37].

To give a purely illustrative example, the early detection and possible prevention of drug side effects can be supported by AI algorithms which allows analysing large sets of data from electronic databases of spontaneous reports, electronic medical records, databases, and digital devices aims to enhance the effectiveness and safety of medicines [24]. This led to consideration of several issues of pharmacology and pharmacovigilance as data-driven fields, which can therefore benefit from AI [13, 21]. Galeano et al. [16] presented an ML framework for predicting unknown side effects for drugs. It uses a set of side effects identified in randomized controlled clinical trials for training. Specifically, this framework learns a drug similarity matrix and a side effect similarity matrix for generating the scores for each drug-side effect pair by linearly combining these models. The frequencies of drug side effects has been computed in according to their own previous work [15] and applied to obtain side effects from clinical trials. Similarly, Liang et al. [23] proposed a method for predicting drug side effects based on the transductive matrix co-completion and leverage the low-rank structure in the side effects and drug-target data. Authors handled the impact of unobserved data by incorporating a positive-unlabelled learning into the model, while the drug chemical information are modelled by using a graph model. Beyond the complexity of the data collection operation and the determination of the hyperparameters, authors treated the prediction of side effects as a multi-label learning problem with missing features and labels, where the drug-target associations (i.e., feature matrix) and side effect labels are binary and linearly dependent. Therefore, the scope of drug targets is limited to the known ones in drug datasets, that is, if a drug causes a side effect by interacting with a protein target that is not in the known target space, this association will probably not

be identified; as also declared by authors. This issue had already been addressed by Campillos et al. [4], which used phenotypic side effect similarities to infer whether two drugs share a target. According to what authors stated, this approach had relevant limitation in reference set where several drugs had a low probability of sharing targets, having though other targets that could override the side effects that might be caused by the drug-target relations. Instead, Liu et al. [25] tried to overcome the issue by first computing all possible interactions between drugs and proteins, from protein structures, and subsequently predicting the side effects.

In the field of interaction modeling, graph network models are widely used in different fields, including the modelling of chemical and biological objects (i.e., biological networks). Traditionally, these are used in their simplest form for handling objects of the same type that interact on a single layer of interest, however, these are able to support heterogeneous objects on multiple layers [10, 11] (i.e., multilayer network), as well as structuring temporal evolutions [8] (i.e., temporal networks). The biological networks on which we focused are based on a multilayered model (i.e., multilayer network) which allows representing multiple biological objects in independent layer that can be also interconnected. In this context, AI is generally applied for addressing a link-prediction problem [7, 9].

In this paper, we propose a framework for predicting candidate drug side effects by using ML on biological networks.

Our contribution concerned the design of a novel framework which integrates a set of predictive methods for the inference of candidate drug side effects, by treating the issue as a link-prediction problem applied to a multilayer network which models *drug-drug*, *drug-side effect* and *chemical drug-gene* interactions.

2 Materials and Methods

2.1 Machine Learning Models in our Framework

The proposed framework integrates a set of well-known ML models, specifically:

- Random Forest Classifier (RFC)
- Support Vector Classifier (SVC)
- Decision Tree Classifier (DTC)
- k-Neighbors Classifier (kNC)
- Logistic Regression (LR)

For descriptive purpose only, we report a brief description for each one, by including their own mathematical definition of the salient elements, as follows:

- *RFC* [22, 34, 35] is an ensemble method that constructs a multitude of decision trees during training and outputs the class that is the mode of the classes (classification) or mean prediction (regression) of the individual trees. It builds multiple decision trees and merges them together to get a

more accurate and stable prediction. Formally, for a set of T decision trees $h_1, h_2, \dots, h_T$, the predicted class $\hat{y}$ for an input $\mathbf{x}$ is:

$$\hat{y} = \text{mode}\{h_t(\mathbf{x}) : t = 1, 2, \dots, T\}.$$

- *SVC*, also known as Support Vector Machine (SVM) [2, 27], is a supervised learning algorithm. It works by finding the hyperplane that best separates the classes in the feature space. SVC is effective in high-dimensional spaces and is particularly useful when the classes are not linearly separable. Given training data $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$ where $\mathbf{x}_i \in \mathbb{R}^d$ and $y_i \in \{-1, 1\}$, the optimization problem is:

$$\min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^n \xi_i,$$

subject to:

$$y_i(\mathbf{w}^\top \mathbf{x}_i + b) \geq 1 - \xi_i, \quad \xi_i \geq 0, \quad \forall i,$$

where $C > 0$ is a regularization parameter and ξ_i are slack variables.

- *DTC* [29, 39] is a tree-like model where an internal node represents a feature (or attribute), the branch represents a decision rule, and each leaf node represents the outcome. It processes data, recursively, based on features that maximize information gain. For a split at node t into child nodes t_L and t_R , the information gain $IG(t)$ is:

$$IG(t) = I(t) - \left[\frac{N_{t_L}}{N_t} I(t_L) + \frac{N_{t_R}}{N_t} I(t_R) \right],$$

where $I(t)$ is the impurity measure (e.g., Gini index or entropy), N_t is the number of samples at the node t , and N_{t_L} , N_{t_R} are the samples in the child nodes.

- *kNC* [14, 33, 38] is a type of instance-based learning or lazy learning where the function is only approximated locally, and all computation is deferred until function evaluation. In k-NN classification, an object is classified by the majority score of its k nearest neighbours. Formally, it predicts the class of a sample $\mathbf{x}$ based on the majority class of its k nearest neighbours in the training data. Let $\mathcal{N}_k(\mathbf{x})$ denote the set of k nearest neighbours of $\mathbf{x}$. The predicted class $\hat{y}$ is:

$$\hat{y} = \text{mode}\{y_i : \mathbf{x}_i \in \mathcal{N}_k(\mathbf{x})\}.$$

- *LR* [30, 36] is a linear model used for binary classification problems. It uses the logistic function to transform a linear combination of input features into a normalized score (i.e., rescaling via min-max normalization), which is interpreted as the probability that a link exists. For input $\mathbf{x}$ and parameters $\mathbf{w}$, the probability of class $y = 1$ is:

$$P(y = 1 \mid \mathbf{x}) = \frac{1}{1 + \exp(-\mathbf{w}^\top \mathbf{x})}.$$

The predicted class $\hat{y}$ is determined by a specific threshold:

$$\hat{y} = \begin{cases} 1, & P(y = 1 | \mathbf{x}) \geq 0.5, \\ 0, & \text{otherwise.} \end{cases}$$

2.2 Inference of Drug Side Effects

The proposed framework allows inferring drug side effects by applying ML on a set of heterogeneous biological objects and their own interactions modelled through a multilayer network graph model (hereinafter referred to as multilayer network).

It extracts meaningful features from a multilayer network by sectioning it layer by layer, including interlayers. These include metrics like the number of connections each node has (degree), the shared neighbours between two nodes, and clustering coefficients that measure how interconnected the neighbours of a node are. It also calculates the shortest path length between the two nodes, which helps capture the global structure of the network. These features serve as the foundation for making predictions.

Specifically, for each pair of nodes (u, v) within a same layer, it extracts the following set of features: (i) number of edges adjacent to the node u , (ii) clustering coefficient for u , (iii) number of edges adjacent to the node v , (iv) clustering coefficient for v , (v) common neighbours between u and v , and (vi) the shortest path length between u and v .

Once the features are extracted, the data is split into training and testing sets to prepare for ML. This ensures the models are trained on one portion of the data and evaluated on another, helping to prevent overfitting.

The framework is versatile and supports several ML models, including RFC, SVC, DTC, kNC, and LR, as described in Section 2.1. By leveraging graph theory and ML, this framework offers a systematic and insightful approach to understanding drug interactions and their potential side effects. It is a powerful tool with applications in drug development, personalized medicine, and beyond.

2.3 Training Dataset

The reported ML models have been trained on an in-house real-world dataset constructed from the Stanford Biomedical Network Dataset Collection (BioSNAP) [40]; specifically, we integrated the data in the following BioSNAP datasets:

- Chemical-Gene Interaction network (ChG-InterDecagon): it contains information on interactions between chemical drugs and genes. Dataset statistics: 1,774 chemical drugs nodes, 7,795 gene nodes, and 131,034 edges.
- Drug Side Effect Association Network (ChSe-Decagon): it contains information on adverse drug reactions (i.e., side effects) caused by drugs that are on the U.S. market. Dataset statistics: 640 drug nodes, 10,185 Side effect nodes, and 174,978 edges.

- Drug-Drug Interaction Network (ChCh-Miner): it contains interactions between drugs, which are approved by the U.S. Food and Drug Administration. Dataset statistics: 1,514 nodes, 48,514 edges.

In order to standardize the identifiers adopted by the various datasets and therefore be able to integrate them, the identifiers in *ChCh-Miner* have been replaced; unfortunately, the porting led to discarding a portion of the nodes and associations originally reported. The final statistics of our dataset resulting from the integration of the above data are reported below:

- Layer 1 (Drugs): nodes: 1178, edges: 41,820.
- Layer 2 (Side Effects): nodes: 10,185, edges: #.
- Layer 3 (Genes): nodes: 7,796, edges: #.
- Interlayer between layers 1 and 2: edges: 174,977.
- Interlayer between layers 1 and 3: edges: 131,034.

Note that the symbol # indicates that the interactions are not defined due to lack of information. However, this information is not useful for training the models in our framework, nor have they been taken (or should be taken) into account as we foresee.

Summarizing, our dataset models drug-drug, drug-side effect and drug-gene interactions/associations. For completeness, we specify that the resulting multi-layer network in our dataset is incomplete due to the lack of side effect-side effect and gene-gene interactions/associations, that would represent the intralayers of the respective layers.

Finally, we randomly divided our dataset into train-dataset and test-dataset consisting of 80% and 20% of the drug-side effect interactions, respectively; isolated nodes were removed.

Train-dataset was used in training step, while test-dataset for experimentation (i.e., testing).

3 Results and Discussion

We performed a series of 5 tests on our test-dataset (see Section 2.3), one for each model implemented in our framework, in order to test our framework and the ML-based models envisaged by it, specifically.

After training, the models are evaluated by using a set of well-known Key Performance Indicators (KPIs): Accuracy, F1 score, Matthews Correlation Coefficient (MCC), Receiver Operating Characteristic Area Under the Curve (ROC-AUC) and Precision-Recall AUC (PR-AUC) [28,31].

These metrics provide a clear picture of how well the models perform and help identify which features contribute most to the predictions. The results offer valuable insights into the relationships within the data, shedding light on the likelihood of side effects and the underlying factors that drive these predictions.

The experimentation evaluated the discussed KPIs of interest to determine the prediction performances in the case study referred to the deduction of novel (candidate) side effects.

The results are reported in Table 1, and the related KPIs are also shown in Figure 1 as bar-plots. These reveal a set of notable observations and issues that we discuss as follow.

RFC emerges as the most effective and reliable model across nearly all evaluation indicators. It achieved the highest accuracy (0.9548), F1-score (0.9543), and MCC (0.9101), as well as the leading ROC-AUC (0.9549) and PR-AUC (0.9420). These results collectively indicate a strong balance between precision and recall, underscoring RFC's ability to generalize well and maintain consistent performance across varying classification challenges.

Closely following RFC, the DTC also demonstrated robust performance. While slightly behind RFC in overall accuracy and F1-score, DTC achieved the highest precision (0.9718), highlighting its strong confidence in positive predictions. Moreover, DTC was the most efficient in terms of runtime, completing in less than one second. This notable computational efficiency, paired with its solid predictive capability, makes it particularly attractive for deployment in real-time or low-resource environments.

The SVC, on the other hand, exhibited a markedly different performance profile. Although it achieved a reasonably high recall (0.7703), indicating its sensitivity to identifying true positives, it lagged behind in precision, accuracy, and MCC. Most critically, the model suffered from a prohibitively long runtime of over 26 minutes (26:46.293), rendering it unsuitable for time-sensitive applications unless optimized or parallelized significantly.

The kNC offered a balanced performance, especially in terms of recall (0.9352) and F1-score (0.9433), and maintained a respectable runtime. It performed slightly lower than RFC and DTC in precision and accuracy but still delivered competitive results overall. This consistency, along with its straightforward interpretability, makes kNC a viable option where moderate computational resources are available, and ease of implementation is desired.

LR, while computationally reasonable, delivered the least favourable results. Despite achieving the highest recall among all models (0.8505), it had the lowest precision (0.5706) and ROC-AUC and PR-AUC scores, which significantly diminishes its reliability. The high recall-low precision trade-off suggests that LR identified a high number of true positives but also produced a substantial number of false positives, limiting its practical application in scenarios where precision is critical.

Several notable observations emerge from this comparison. Firstly, ensemble methods such as RFC clearly outperform single estimators, indicating the benefit of model aggregation in enhancing generalization. Secondly, a model's predictive power must be balanced against its runtime, as exemplified by SVC otherwise acceptable performance being offset by its impractical computational cost. Thirdly, precision-recall trade-offs are evident across several models, particularly in LR, highlighting the importance of selecting models not solely based on accuracy but also on the specific requirements of the application—be it minimizing false positives or maximizing detection sensitivity.

Summarizing, RFC has proven to be the most suitable model when both performance and reliability are prioritized. Nevertheless, the DTC offers an efficient and nearly as effective alternative, especially when speed is essential. Models like kNC provide a balanced middle ground, while SVC and LR may require further tuning or contextual justification for their use. These findings emphasize the need for a comprehensive evaluation framework when selecting models, one that considers not only predictive metrics but also computational demands and domain-specific constraints.

Overall, RFC stands out for performance was the best model able to predict candidate drug side effects from biological networks in the application scope and more specifically for the dataset used. One cannot help but notice the lack of efficiency in terms of running time shown by SVC; otherwise, other models perform the processing in the order of seconds.

Ultimately, we dispense with producing a deep comparison between the applied methodologies since our framework includes an assisted phase of suggestion of the best model for the case in question, consequently the interpretation of the results for the use case in which he will want to employ the solution is left to the user; a future work could certainly focus on deepening this aspect as well.

Table 1: The table reports the results related to our experimentation. Briefly, we evaluated the performances related to each method, in terms of Accuracy, F1 score, ROC-AUC, and PR-AUC. Running times (runtimes) are also reported in minutes, seconds and milliseconds (mm:ss.ms format).

Method	Runtime	Accuracy	F1 Score	MCC	AUC	ROC AUC	PR
RFC	00:14.440	0.9548	0.9543	0.9101	0.9549	0.9420	
SVC	26:46.293	0.6858	0.7116	0.3762	0.6853	0.6248	
DTC	00:00.732	0.9525	0.9519	0.9059	0.9527	0.9403	
kNC	00:02.386	0.9435	0.9433	0.8871	0.9435	0.9226	
LR	00:01.366	0.6028	0.6830	0.2337	0.6012	0.5605	

4 Conclusions

In this paper, we proposed a framework for predicting candidate drug side effects by using ML on biological networks. Our contribution concerned the design of a novel framework which integrates a set of predictive models for the inference of candidate drug side effects.

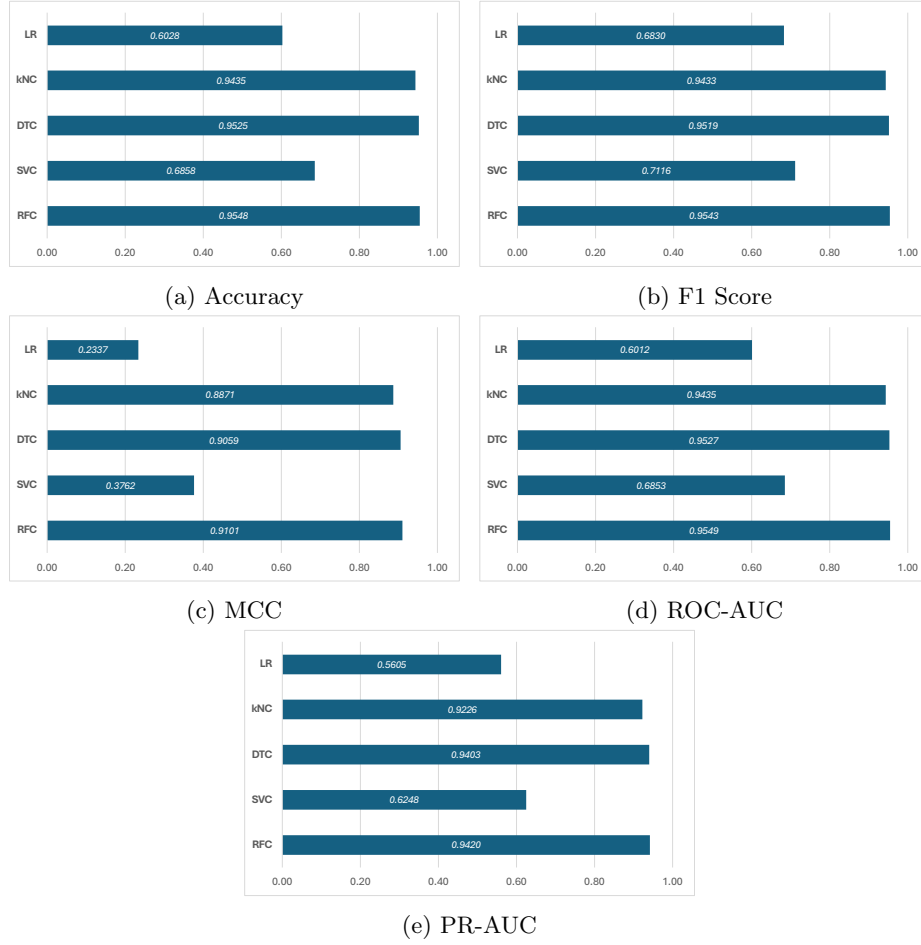


Fig. 1: The plots were built based on the results from Table 1. briefly, we evaluated the performances related to each method, in terms of Accuracy, F1 score, ROC-AUC, and PR-AUC.

By leveraging graph theory and ML, the proposed framework offers a systematic and insightful approach to understanding drug interactions and their potential side effects. It may be applied as a powerful tool with applications in drug development, personalized medicine, and beyond.

Experimentation supports the models implemented in our framework, by demonstrating their own effectiveness, in accordance with the KPIs used in the examination of the results.

The main limitation of our study is that it was limited to implement existing ML-based models, adapting them for use on multilayer networks, and evaluating the results by exploiting a set of KPIs for descriptive purposes only for each of the supported models.

Future works could be focused on improving experimentation, as well as on proving the validity of candidate side effect in real-world experimentation. Additionally, novel techniques concerning DL can also be supported.

Acknowledgements

- We acknowledge the support of the PNRR project FAIR - Future AI Research (PE00000013), Spoke 9 - Green-aware AI, under the NRRP MUR program funded by the NextGenerationEU.
- We acknowledge the support of the project “PRIN2022 - LIRA: LsrK as innovative molecular target for quorum sensing interfering agents for fighting Resistance to Antimicrobials” (Prot.: 2022BCRZK2, CUP: F53D23005060001).
- We acknowledge the support of Next Generation EU - Italian NRRP, Mission 4, Component 2, Investment 1.5, call for the creation and strengthening of 'Innovation Ecosystems', building 'Territorial R&D Leaders' (Directorial Decree n. 2021/3277) - project Tech4You - Technologies for climate change adaptation and quality of life improvement, n. ECS0000009. This work reflects only the authors' views and opinions, neither the Ministry for University and Research nor the European Commission can be considered responsible for them.
- The scientific activities of the CINI InfoLife Laboratory supported this research.

References

1. Bajwa, J., Munir, U., Nori, A., Williams, B.: Artificial intelligence in healthcare: transforming the practice of medicine. *Future Healthc. J.* **8**(2), e188–e194 (Jul 2021)
2. Bermolen, P., Rossi, D.: Support vector regression for link load prediction. *Comput. Netw.* **53**(2), 191–201 (Feb 2009)
3. Bhattamisra, S.K., Banerjee, P., Gupta, P., Mayuren, J., Patra, S., Candasamy, M.: Artificial intelligence in pharmaceutical and healthcare research. *Big Data Cogn. Comput.* **7**(1), 10 (Jan 2023)
4. Campillos, M., Kuhn, M., Gavin, A.C., Jensen, L.J., Bork, P.: Drug target identification using side-effect similarity. *Science* **321**(5886), 263–266 (2008)

5. Caroprese, L., Cascini, P.L., Cinaglia, P., Dattola, F., Franco, P., Iaquina, P., Iusi, M., Tradigo, G., Veltri, P., Zumpano, E.: Software tools for medical imaging extended abstract. In: *New Trends in Databases and Information Systems*. pp. 297–304. Springer International Publishing, Cham (2018)
6. Choi, R.Y., Coyner, A.S., Kalpathy-Cramer, J., Chiang, M.F., Campbell, J.P.: Introduction to machine learning, neural networks, and deep learning. *Transl. Vis. Sci. Technol.* **9**(2), 14 (Feb 2020)
7. Cinaglia, P.: Pymulsim: a method for computing node similarities between multilayer networks via graph isomorphism networks. *BMC Bioinformatics* **25**(1) (Jun 2024). <https://doi.org/10.1186/s12859-024-05830-6>, <http://dx.doi.org/10.1186/s12859-024-05830-6>
8. Cinaglia, P., Cannataro, M.: Network alignment and motif discovery in dynamic networks. *Network Modeling Analysis in Health Informatics and Bioinformatics* **11**(1) (Oct 2022). <https://doi.org/10.1007/s13721-022-00383-1>, <http://dx.doi.org/10.1007/s13721-022-00383-1>
9. Cinaglia, P., Cannataro, M.: Identifying candidate gene-disease associations via graph neural networks. *Entropy (Basel)* **25**(6) (Jun 2023)
10. Cinaglia, P., Cannataro, M.: Multiglobal: Global alignment of multilayer networks. *SoftwareX* **24**, 101552 (2023). <https://doi.org/https://doi.org/10.1016/j.softx.2023.101552>
11. Cinaglia, P., Milano, M., Cannataro, M.: Multilayer network alignment based on topological assessment via embeddings. *BMC Bioinformatics* **24**(1) (Nov 2023). <https://doi.org/10.1186/s12859-023-05508-5>, <http://dx.doi.org/10.1186/s12859-023-05508-5>
12. Cinaglia, P., Vázquez-Poletti, J.L., Cannataro, M.: Massive parallel alignment of rna-seq reads in serverless computing. *Big Data and Cognitive Computing* **7**(2) (2023). <https://doi.org/10.3390/bdcc7020098>
13. Dong, Y., Yang, T., Xing, Y., Du, J., Meng, Q.: Data-driven modeling methods and techniques for pharmaceutical processes. *Processes (Basel)* **11**(7), 2096 (Jul 2023)
14. Fava, B., Marques, P.C., Lopes, H.F.: Probabilistic nearest neighbors classification. *Entropy (Basel)* **26**(1), 39 (Dec 2023)
15. Galeano, D., Li, S., Gerstein, M., Paccanaro, A.: Predicting the frequencies of drug side effects. *Nat. Commun.* **11**(1), 4575 (Sep 2020)
16. Galeano, D., Paccanaro, A.: Machine learning prediction of side effects for drugs in clinical trials. *Cell Rep. Methods* **2**(12), 100358 (Dec 2022)
17. Gao, F., Huang, K., Xing, Y.: Artificial intelligence in omics. *Genomics, Proteomics and Bioinformatics* **20**(5), 811–813 (Oct 2022). <https://doi.org/10.1016/j.gpb.2023.01.002>
18. Ginsburg, G.S., Phillips, K.A.: Precision medicine: From science to value. *Health Aff. (Millwood)* **37**(5), 694–701 (May 2018)
19. Hong, H., Slikker, W.: Integrating artificial intelligence with bioinformatics promotes public health. *Experimental Biology and Medicine* **248**(21), 1905–1907 (Nov 2023). <https://doi.org/10.1177/15353702231223575>
20. Kolluri, S., Lin, J., Liu, R., Zhang, Y., Zhang, W.: Machine learning and artificial intelligence in pharmaceutical research and development: a review. *The AAPS Journal* **24**(1) (Jan 2022). <https://doi.org/10.1208/s12248-021-00644-3>
21. Kompa, B., Hakim, J.B., Palepu, A., Kompa, K.G., Smith, M., Bain, P.A., Woloszynek, S., Painter, J.L., Bate, A., Beam, A.L.: Artificial intelligence based on machine learning in pharmacovigilance: A scoping review. *Drug Saf.* **45**(5), 477–491 (May 2022)

22. Li, K., Tu, L.: Link prediction for complex networks via random forest. *J. Phys. Conf. Ser.* **1302**(2), 022030 (Aug 2019)
23. Liang, X., Fu, Y., Qu, L., Zhang, P., Chen, Y.: Prediction of drug side effects with transductive matrix co-completion. *Bioinformatics* **39**(1) (Jan 2023)
24. Litvinova, O., Yeung, A.W.K., Hammerle, F.P., Mickael, M.E., Matin, M., Kletecka-Pulker, M., Atanasov, A.G., Willschke, H.: Digital technology applications in the management of adverse drug reactions: Bibliometric analysis. *Pharmaceuticals (Basel)* **17**(3) (Mar 2024)
25. Liu, T., Altman, R.B.: Relating essential proteins to drug side-effects using canonical component analysis: a structure-based approach. *Journal of chemical information and modeling* **55**(7), 1483–1494 (2015)
26. Mak, K.K., Pichika, M.R.: Artificial intelligence in drug development: present status and future prospects. *Drug Discov. Today* **24**(3), 773–780 (Mar 2019)
27. Malhotra, D., Goyal, R.: Supervised-learning link prediction in single layer and multiplex networks. *Mach. Learn. Appl.* **6**(100086), 100086 (Dec 2021)
28. Nahm, F.S.: Receiver operating characteristic curve: overview and practical use for clinicians. *Korean J Anesthesiol* **75**(1), 25–36 (Feb 2022)
29. Nair, L.S., Jayaraman, S., Krishna Nagam, S.P.: An improved link prediction approach for directed complex networks using stochastic block modeling. *Big Data Cogn. Comput.* **7**(1), 31 (Feb 2023)
30. Parsaeian, M., Mohammad, K., Mahmoudi, M., Zeraati, H.: Comparison of logistic regression and artificial neural network in low back pain prediction: second national health survey. *Iran. J. Public Health* **41**(6), 86–92 (Jun 2012)
31. Saito, T., Rehmsmeier, M.: The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets. *PLoS One* **10**(3), e0118432 (Mar 2015)
32. Salas, M., Petracek, J., Yalamanchili, P., Aimer, O., Kasthuril, D., Dhingra, S., Junaid, T., Bostic, T.: The use of artificial intelligence in pharmacovigilance: a systematic review of the literature. *Pharmaceut. Med.* **36**(5), 295–306 (Oct 2022)
33. Salvador-Meneses, J., Ruiz-Chavez, Z., Garcia-Rodriguez, J.: Compressed kNN: K-nearest neighbors with data compression. *Entropy (Basel)* **21**(3), 234 (Feb 2019)
34. Wang, T., Jiao, M., Wang, X.: Link prediction in complex networks using recursive feature elimination and stacking ensemble learning. *Entropy (Basel)* **24**(8), 1124 (Aug 2022)
35. Wu, M., Huang, Y., Zhao, L., He, Y.: Link prediction based on random forest in signed social networks. In: 2018 10th International Conference on Intelligent Human-Machine Systems and Cybernetics (IHMSC). IEEE (Aug 2018)
36. Wu, Q., Zhang, Z., Ma, T., Waltz, J., Milton, D., Chen, S.: Link predictions for incomplete network data with outcome misclassification. *Stat. Med.* **40**(6), 1519–1534 (Mar 2021)
37. Zhang, Y., Luo, M., Wu, P., Wu, S., Lee, T.Y., Bai, C.: Application of computational biology and artificial intelligence in drug design. *International Journal of Molecular Sciences* **23**(21), 13568 (Nov 2022). <https://doi.org/10.3390/ijms232113568>
38. Zhang, Z.: Introduction to machine learning: k-nearest neighbors. *Ann. Transl. Med.* **4**(11), 218 (Jun 2016)
39. Zinilli, A., Cerulli, G.: Link prediction and feature relevance in knowledge networks: A machine learning approach. *PLoS One* **18**(11), e0290018 (Nov 2023)
40. Zitnik, M., Sosić, R., Maheshwari, S., Leskovec, J.: BioSNAP Datasets: Stanford biomedical network dataset collection. <http://snap.stanford.edu/biodata> (Aug 2018)