

Risk-based Monitoring Using AI/ML Detection of Systematic Bias in Clinical Data

Kostiantyn Drach, University of Barcelona/Intego Group, Barcelona, Spain
Sergey Glushakov, Intego Group, Maitland, USA

ABSTRACT

To effectively plan and execute clinical trials, it is essential to oversee data integrity as part of risk-based monitoring (RBM). Inaccuracy and incompleteness in clinical data can be explained by systematic bias in multiple metrics considered for RBM execution. Systematic bias is hard to detect using standard statistical methods. Thus, we propose an approach to reveal this type of inconsistency in data using AI/ML methods together with graph-based machine learning algorithms. To perform the experiment, we use clinical datasets with synthetically generated bias as proof of concept. Our approach relies on pure geometric properties of data combined with AI/ML algorithms, such as neural networks, community detection, etc. The experiment unveils the systematic bias hidden in data with no prior knowledge of its source. This approach improves RBM in terms of error detection and may optimize some or all on-site monitoring visits, mitigate risks related to data inaccuracy, and improve data quality.

INTRODUCTION

Recently, the growth and complexity of clinical trials have presented challenges in oversight due to increased variability in investigator experience, site infrastructure, treatment choices, and healthcare standards. Nevertheless, advancements in computer systems, electronic records, statistical assessments as well as AI/ML usage offer opportunities for alternative monitoring approaches, such as centralized monitoring, to enhance the quality and efficiency of sponsor oversight. According to [1], RBM is the process of ensuring the quality of clinical trials by identifying, assessing, monitoring, and mitigating the risks that could affect the quality or safety of a study. The Food and Drug Administration (FDA) encourages sponsors to develop monitoring plans that address challenges emphasizing a risk-based approach focused on preventing or mitigating significant risks to data quality and processes critical to human subject protection and trial integrity.

The key components of RBM, as outlined in [2], are:

- key risk indicators (KRIs);
- centralized monitoring;
- of-site/remote-site monitoring;
- reduced *source data verification* (SDV);
- reduced *source document review* (SDR).

One of the key components of RBM is its use of centralized monitoring techniques. As opposed to on-site monitoring based on 100% SDV, centralized monitoring offers a number of advantages:

- **fewer errors** due to less manual work;
- **lower cost due to** reducing the frequency and extent of on-site monitoring and auditing only the sites where problems are most likely to occur;
- **better analysis and cross-site comparison**, since centralized monitoring also allows us to compare data between sites and to identify potentially fraudulent, inaccurate, or biased data.

The FDA has provided some detailed guidance on how to prepare a monitoring plan [1], although execution of the plan for specific clinical trials is still challenging and requires the search for new approaches. We propose some solutions in RBM based on *topological data analysis* (TDA) approach in combination with other AI/ML techniques (see Section 1 for solutions ideas and Sections 2 and 3 for methods description). These solutions can help to supplement the statistical by-site data processing, identify some KRIs, as well as can be used to reveal the problematic sites and provide alerts to perform targeted on-site investigation. These solutions are flexible enough to account for any equity and diversity within the study.

To demonstrate the working prototypes of solutions, we use a publicly available NIDA dataset [3]. The original NIDA experiment investigated if the buprenorphine/naloxone combination tablet can be effectively used to treat patients with opiate dependence. A total of 582 persons on 38 sites were recruited. Study participation lasted 9 to 12 weeks for patients who successfully achieved detoxification and up to 12 months for patients requiring longer buprenorphine treatment. On baseline and during the trial, various data concerning general health, vital signs, drug addiction, family history, psychological health, employment status, and other information (total over 280 features) were gathered, as well as treatment results were indicated. The variety of data allowed us to carry out two experiments, which are described in Sections 4 and 5.

1. RISK BASED MONITORING: CHALLENGES AND SOLUTIONS

RBM is a multi-faced procedure designed to effectively identify and address risks within a clinical trial at its early stages before they become serious issues. Compared to other forms of monitoring focused on past events (such as SDV, SDR, etc.), RBM mainly focuses on the present and the future (including predictive modeling, etc.).

Within RBM, we have singled out two major challenges to address:

- Clinical trial site monitoring.** Controlling activities on the sites is an important task to ensure correct protocol implementation and safety of patients and the resulting drugs. We propose an approach based on TDA that enables the detection of systematic bias on sites, which is challenging to identify using classical approaches. The approach is explained in Figure 1.
- Ongoing risk assessment related to inclusion of new patients.** Many trial designs imply including one or several new cohorts of patients into the study, for example, when the study has several phases, or when the required primary study objectives are not met due to incorrect/inaccurate sample size determination. We propose an approach that predicts completion of the protocol by newly added patients based on their baseline data and the data available for the patients who has already finished the study (either completed the protocol or dropped out due to various factors). Our approach is based on a combination of TDA and *Graph Neural Networks* (GNNs) and is illustrated in Figure 2.

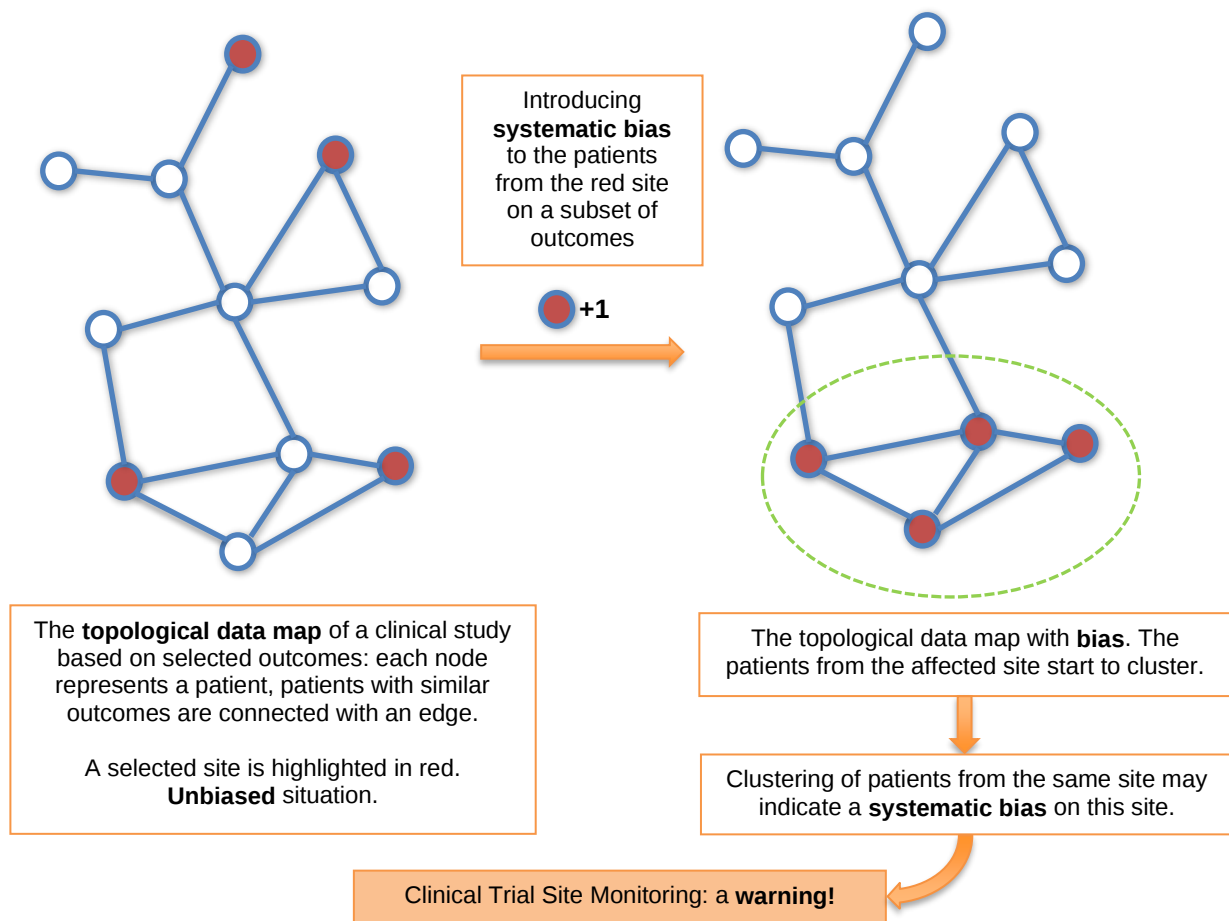


Figure 1. Bias detection in clinical trial site monitoring

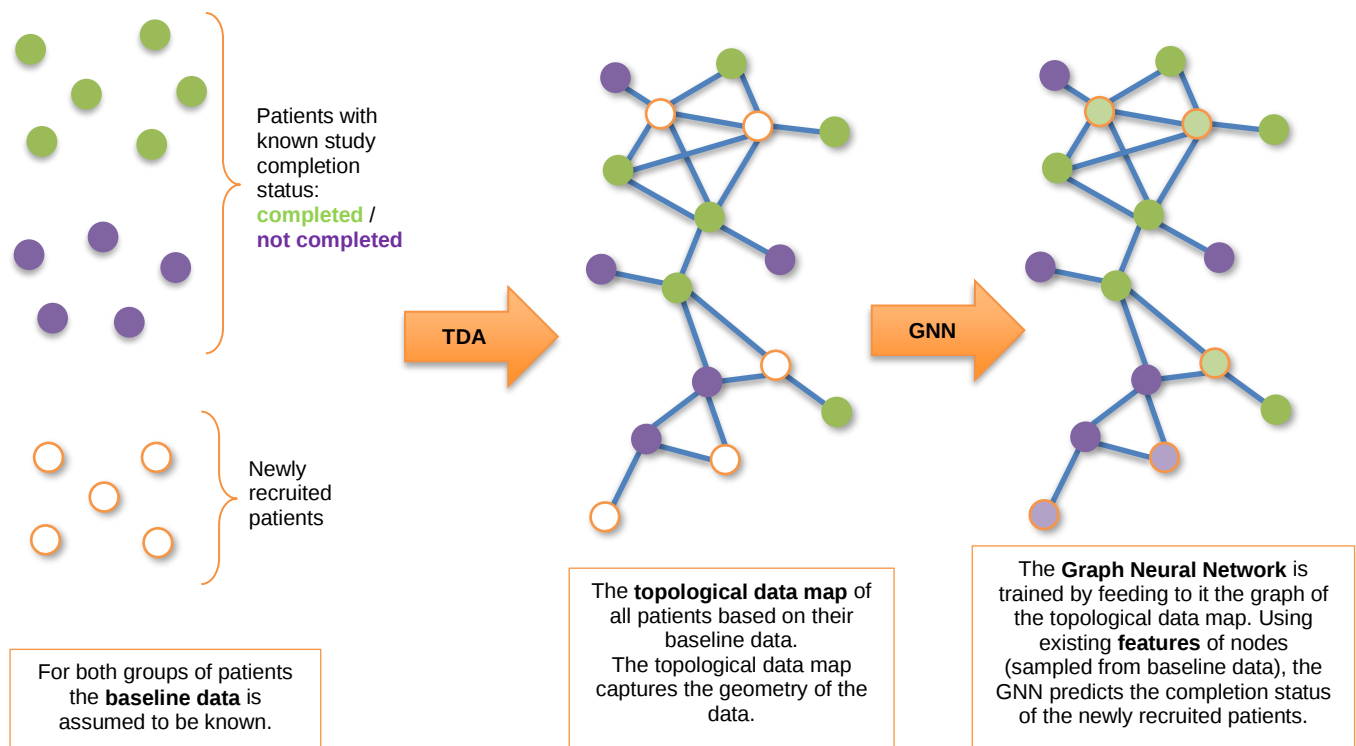


Figure 2. Ongoing risk assessment related to inclusion of new patients

2. TOPOLOGICAL DATA ANALYSIS

TDA is a novel approach of building visual representations of complex datasets. This analysis yields the extraction of comprehensive graphs from a dataset to provide a compressed graphical representation of a multidimensional set of interrelated outcomes. When applied to clinical data, this graph consists of nodes corresponding to patients participating in a study and edges connecting those patients who share similarities in terms of study outcomes or other indicators of interest.

TDA aims to extract topological data, namely, qualitative information, from finite sets of data points. TDA involves exploring datasets (viewed as finite clouds of points in a multidimensional space) at multiple scales or resolutions, from fine- to coarse-grained. Given a complex dataset, TDA can be used to extrapolate the underlying topology and build a compressed yet comprehensive topological summary of the dataset. TDA exploits various methods and algorithms stemming from computational topology and geometry, statistics, and data mining. For detailed expositions of the mathematical theories that underpin TDA and certain applications in biology, see [4, 5] and the references therein.

In the context of clinical research, the dataset under study is typically a table of outcomes in a particular clinical trial or study. The table rows correspond to individual participants in the clinical trial, and the columns contain information on specific outcome measures of interest, such as lab tests, vitals, questionnaires, etc. Given a table of clinical outcomes, two types of elements should be specified to generate a graph using TDA. The first element is a projection, which serves as a function used to stratify patients into subpopulations (bins). The second element is a distance function, which measures the proximity between patients in the multidimensional space. The distance function makes it possible to connect patients with similar outcomes within each bin.

To be considered for further analysis, a graph extracted from the dataset using TDA algorithms should meet certain requirements. Namely, the graph should:

- accurately represent the original dataset;
- eliminate the features of the dataset that are not relevant to the purpose of the study;
- reduce the complexity of the features that are shown on the data map; and
- be insensitive to small noise, such as errors of measurement.

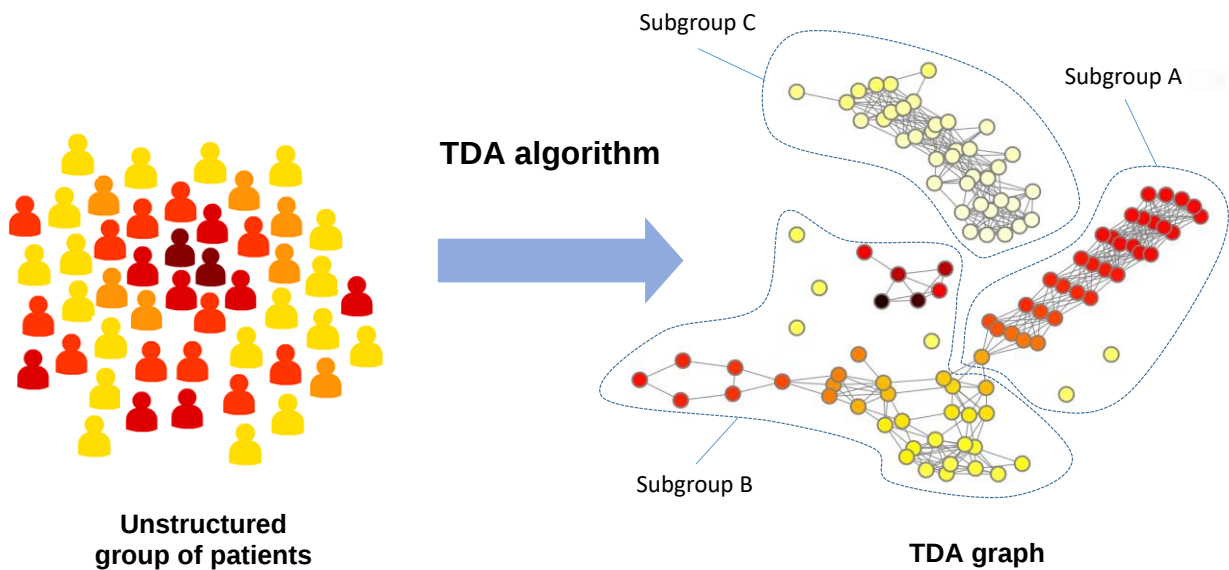


Figure 3. Discovery of multivariate patterns in clinical trial outcomes. The graph represents groups of patients structured according to the similarity of their outcomes

The core idea of clinical data mining using TDA relies on the visual/algorithmic discovery of subgroups of related patients in a graph (see Figure 3) that presents the relevant information about the dataset in a compact and efficient manner. For the clinical dataset, the following criteria have to be met to perform the analysis:

- **Each node represents a patient:** a graph extracted from a clinical dataset is actually a graphical representation of the dataset in which each node represents an individual (a trial subject).
- **Similar nodes are connected:** two nodes representing similar patients (in terms of a predefined set of clinical outcomes) are connected with an edge.
- **Coloring focused on specific outcomes/predictors:** the color of the nodes helps to highlight emerging patterns in the data and identify subgroups of patients related to the distribution of a variable of interest.
- **Visual or ML-powered discovery of subgroups:** clusters or "communities" of nodes on a graph reflect a segmentation of patients that may indicate robust patterns within the data.

TDA was successfully applied in the context of clinical studies as well as for the analysis of real-world data, e.g., in understanding the spread of COVID-19 pandemic (see [6, 7] and references therein).

3. GRAPH NEURAL NETWORKS

GNNs have attracted substantial attention in recent years owing to their ability to effectively model and analyze data with complex relational structures. Traditional machine learning algorithms often struggle to capture the intricate relationships and dependencies present in a graph-structured data, prompting the need for more advanced techniques. GNNs inherit graph topology to perform computation and feature extraction and their use has shown significant potential in addressing classification problems within graph-structured data, offering a data-driven approach to infer labels based on the underlying relational information.

GNNs are designed to operate directly on graph-structured data, enabling capturing both the local and global structural patterns within the graph. This capability makes GNNs well-suited for classification tasks, where the goal is to predict a label or class for each node in the graph based on its features and relational connectivity within a graph structure, as well as based on the existing labels. By aggregating information from neighboring nodes and propagating it through multiple layers, GNNs can effectively learn to discriminate between different classes based on the features and connectivity patterns present in the graph (see Figure 4).

There are various architectures of GNNs models. One of the most popular ones is a *Graph Convolutional Network*. The key difference between architectures of GNN models is how the GNN models aggregate information from the one-hop neighborhood (i.e., for the set of vertices connected to a given vertex with exactly one edge). For more details see [8, 9].

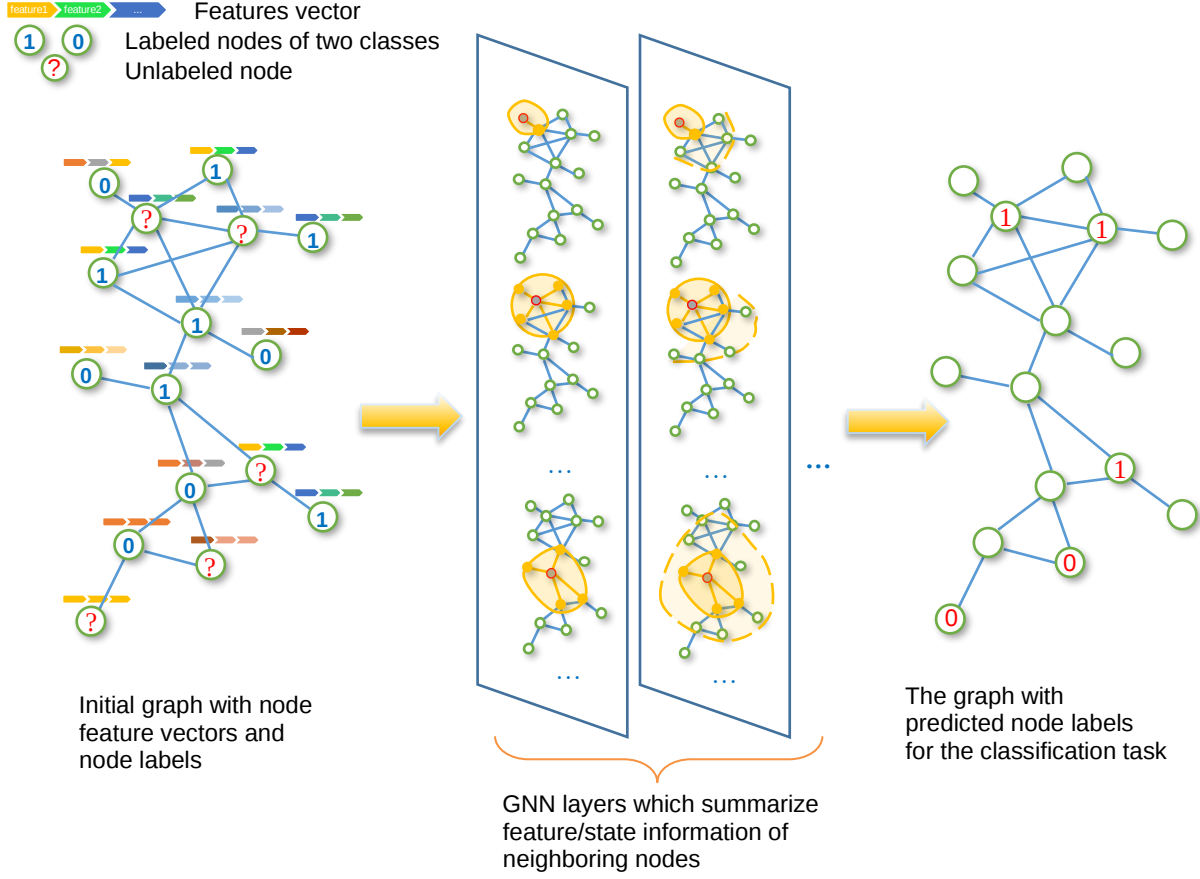


Figure 4. GNN solving a classification task

To effectively train a GNN or other AI/ML predictive model, appropriate methods are to be applied (see [10]). First, for testing purposes, some correctness metrics are to be employed (e.g., Accuracy, F1, ROC-AUC). Next, balancing between *underfitting* (inadequate fitting to training samples) and *overfitting* (capturing noise in training data) should be taken into account. Underfitting leads to poor performance on both training and test sets, often due to insufficient tuning or inadequate training. Overfitting occurs when the GNN model is overly tuned to the training set, resulting in excellent performance on training data but poor performance on the test set (see Figure 5).

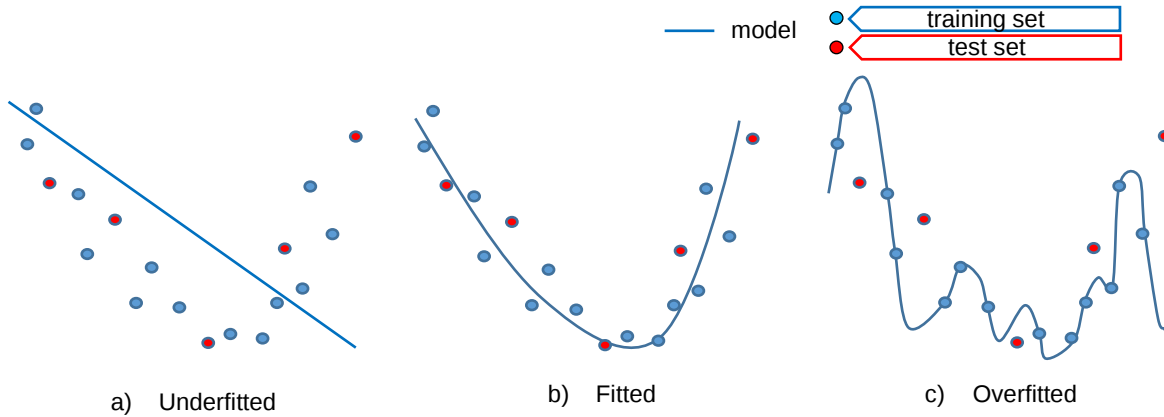


Figure 5. Illustration of a) underfitting, b) fitting, and c) overfitting GNN models trained on a training set

Selecting the GNN model with the highest accuracy on the training set does not guarantee future good performance on new data. Validation aims to estimate model's performance on future data. The drawback of a single train-test split is that the test set may not follow the same distribution of classes in general. Also, some numerical features may not

have the same distribution in the training and test set. The K -fold cross-validation addresses this issue by ensuring that every data sample appears in the test set across multiple iterations.

In our setting shown in Figure 2, the test set is fixed and explicitly given (the set of new patients). Hence, we employ k -fold cross-validation using only the training set. With k distinct training/validation splits, each of which provides scores based on a selected metric, we estimate the generalization performance of the model. Averaging these scores and analyzing deviations provides an overall evaluation. For each partition, an internal model selection procedure selects the hyper-parameters using the training data only. Thus, **test data** is never used for model selection. This process results in different ‘best’ hyperparameter configurations for each split, leading to the evaluation of a class of models. See Table 1 for the whole process.

Table 1. K -fold cross-validation

Training set	Validation set	Training set	Training set	Training set	Training set	} <i>k</i> folds
	Training set	Validation set	Training set	Training set	Training set	
	Training set	Training set	Validation set	Training set	Training set	
	Training set	Training set	Training set	Validation set	Training set	
	Training set	Training set	Training set	Training set	Validation set	
Test set	Score1	Score2	Score3	Score4	Score5	} Average

Besides, if target classes are imbalanced, the stratified k -fold modification of cross-validation should be considered. In this case, splitting the data set in k folds should provide each fold with approximately the same distribution of target classes.

4. EXPERIMENT 1: BIAS DETECTION ON SITES

Our experiment concerns the bias detection in data collected during a clinical study at a site. For this research, we select some features, which are measured for each patient during visits. For each feature, using TDA we can build a topological model (also called feature-based graph), in which nodes are patients and edges connect those patients who have similar dynamics of the feature over several visits. The grouping of patients from one site on the graph may indicate that distortions are introduced into the measurement of the feature, which may be a signal to check the data collection at this site.

4.1. DATA PREPROCESSING AND GRAPH CONSTRUCTION

To demonstrate how data bias affects the feature-based graph, we selected 339 patients out of 582 who completed at least 10 visits (including the baseline visit) during the first 9 weeks of the study. Among the selected 339 patients, there were those who had one missed visit during these weeks. Data on this missed visit was completed using the Last Observation Carried Forward (LOCF) method. The features measured during these 9 visits included both binary features (the presence or the absence of drugs in the urine, the presence or the absence of certain symptoms of opiate addiction) and numerical features (temperature, pulse, blood pressure, respiratory rate, etc.). Data on the presence of drugs, as well as data on the presence of symptoms of opiate dependence were aggregated in such a way as to obtain new indicators, namely the number of different drugs in the urine and the number of different symptoms of opiate dependence by visit. Numerical features were left unchanged. Using TDA, topological models were built for each of the features (viewed in dynamics over visits) and were used to track where the patients belonging to the same site were located (see Figure 6).

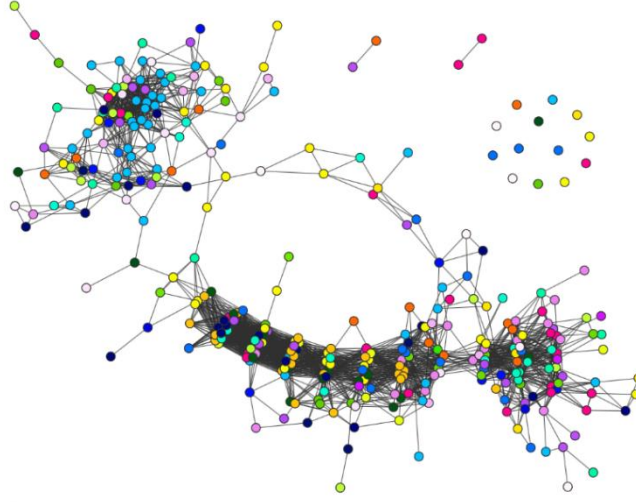


Figure 6. Topological model based on the number of opiate symptoms.
The nodes corresponding to patients are colored by sites

4.2. RESULTS AND DISCUSSION

In the experiment, we introduced the synthetic bias to the data collected for a specific feature on a specific site. Then, we built a new feature-based graph using the biased data and tracked how the position of patients from the selected site on the new graph had changed. We biased the discrete data by shifting all data on the site by a certain value. For numerical features, we used both shifting by some fixed value and addition of some Gaussian noise with certain parameters of the mean and variance. In each case, we ensured that adding noise to the data did not push the mean of the data beyond the three-sigma interval, otherwise such distortion could be detected by measuring simple statistical quantities.

To track the “degree of clustering” of patients belonging to the same site, we used the connectivity index, which was calculated as the fraction of edges connecting patients of the same site to the maximum possible number of edges connecting these patients.

Figure 7 shows the example of bias introduced to pulse dynamics at a site (indicated by a number).

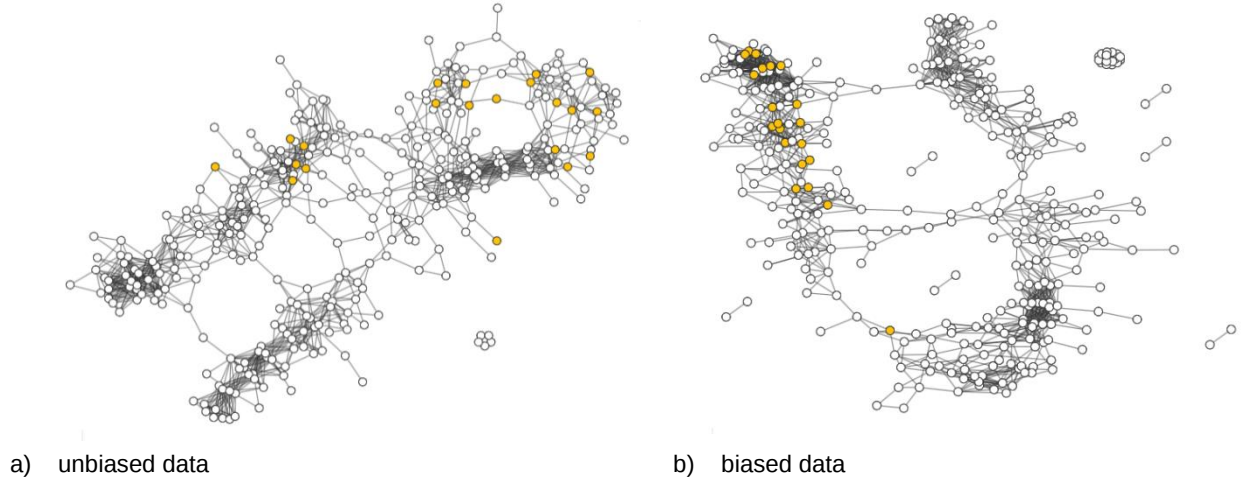


Figure 7. Feature-based graph by pulse dynamics.

The patients from site number 622 are colored. The bias introduced is a constant shift by +10, the connectivity index of unbiased site 622 data is 0.03, and the connectivity index of the biased data is 0.28. The mean over all sites is $\mu = 77.6$ and standard deviation is $\sigma = 4.3$. The mean over site 622 in the biased data is 88.9, which is less than $\mu + 3\sigma$

Thus, in the biased data, the connectivity index of patients related to the site with the bias has increased.

Note that in the clinical trial taken for the experiment there was a site (number 178) patients of which grouped on the feature-based graphs for almost all features without any synthetic bias (see Figure 8). This grouping may be caused by bias in original data, inaccurate data collection, non-random selection of patients, or other reasons which should be investigated apart.

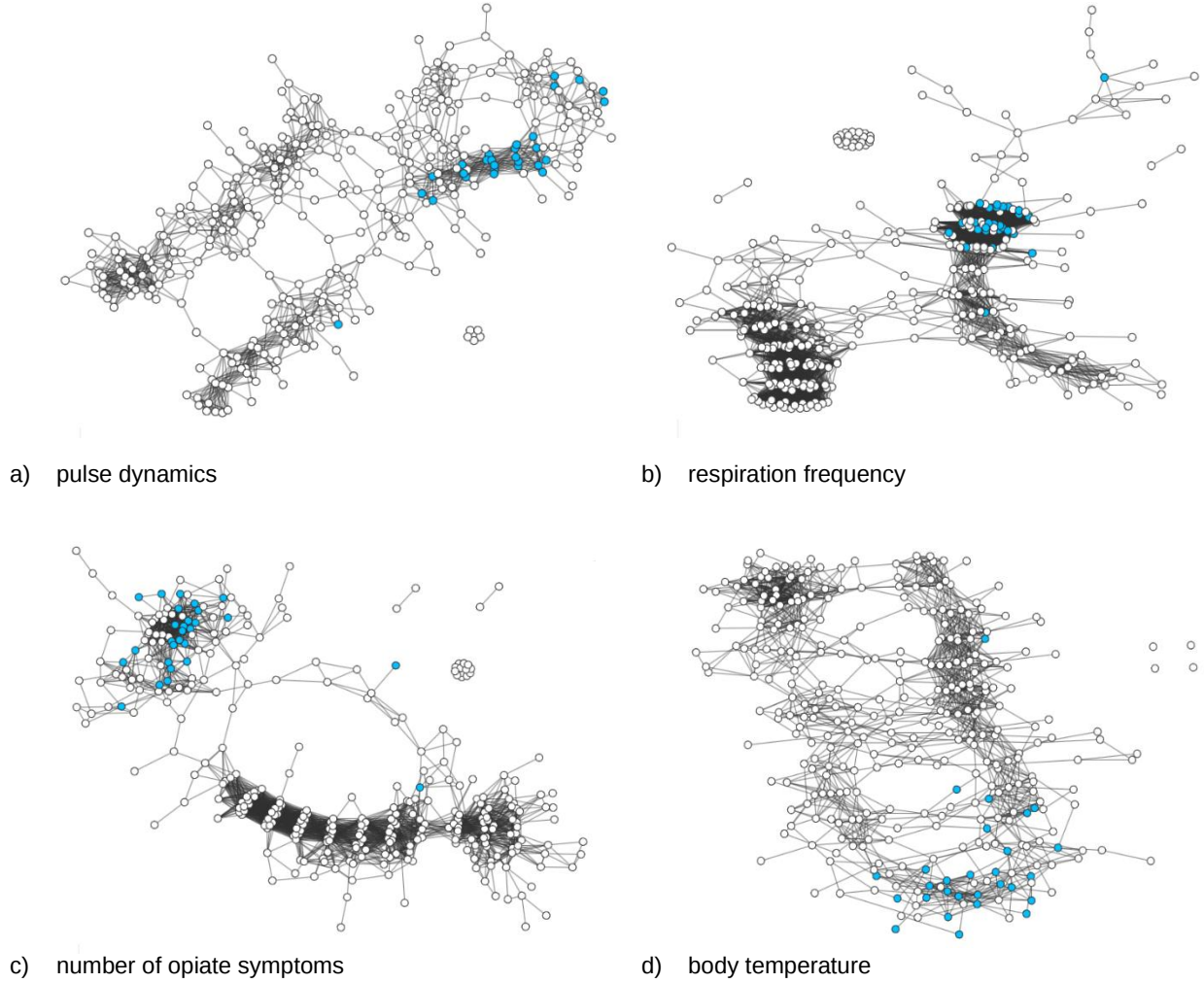


Figure 8. Feature-based graph with unbiased data for site 178.

The similarity of the patients from site 178 by a number of different features is a signal to check the data gathering and processing on this site

In all the experiments we have carried out, a simple statistical analysis does not make it possible to suspect that a bias has been introduced into the data of a certain site, since the average values of the feature after bias introducing remain within three (or even one or two) standard deviations from the mean for all sites. This highlights the advantages of the TDA approach, which provides a clear understanding of which sites and data should be checked. In the actual data, a clear warning for additional investigation should be issued with regards to site 178.

5. EXPERIMENT 2: PREDICTION OF PATIENTS' STUDY COMPLETION STATUS

The second experiment involves the task of predicting whether patients will complete the study based on data collected from them at the baseline. The data include the following information:

- demography data including gender, age, religion, and race;
- medical data on diseases;
- employment/support status (data about education, employment status, and income);
- alcohol/drugs (quantified alcohol and drug abuse indicators);
- legal status (i.e., if the patient has committed any offenses);
- family history (including alcohol and drug abuse, mental problems among family members);
- family/social relationship (data on marital status, social relationships);
- psychiatric status (including mental health data).

The objective of the experiment was a binary classification task that determines whether the patient completed the study successfully or not. Various ML algorithms, such as logistic regression, the nearest neighbor classifier, decision trees, or random forests, can address this issue. Alternatively, one can employ a neural network approach for solving

this problem. We opted for logistic regression and neural networks, subsequently comparing their scores. The overall scheme of the experiment is depicted in Figure 9.

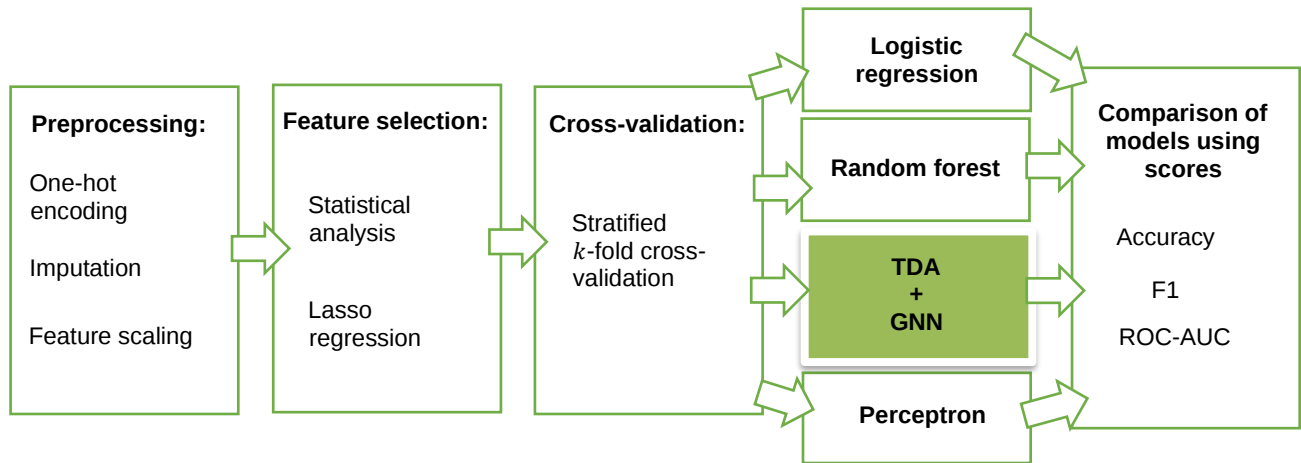


Figure 9. The general scheme of the second experiment

As the original baseline dataset comprised over 280 features (columns) and contained numerous missing values, preprocessing was necessary. For 15 patients, all essential data was absent and thus these patients had to be excluded. Additionally, certain features were categorical variables, requiring transformation before direct application of ML algorithms.

The **Preprocessing block** included the following steps:

- one-hot encoding was applied to categorical data, as a result of which instead of a column with categorical data, binary columns for each category are added to the table;
- data were imputed using the k -nearest neighbors method with $k = 5$. Missing values of each sample were imputed using the mean value from k nearest neighbors found in the source dataset. Two samples are close if the features that neither of samples is missing are close;
- since the features are scaled differently, the data were normalized using the RobustScaler method. This method uses robust estimates for the center and range of data. The RobustScaler removes the median and scales the data according to the interquartile range.

After preprocessing, quite a lot of features remained, some of which were highly correlated with each other, so the following methods were applied for **Feature selection**:

- Statistical analysis. The data were divided into two parts: those patients who successfully completed the study and those ones who did not complete the study. Next, statistically significant features were selected. For the purpose of the statistical analysis undertaken for the present experiment, continuous, mixed, binary, and categorical (non-binary) univariate predictors were differentiated according to a variable type. Continuous predictors were examined using the standard two-sample Mann–Whitney U rank test. This method verifies whether two data samples are obtained from the same distribution. The test assumes that the underlying distributions are continuous (no ties in the values of variables are allowed). To examine the statistical association between two samples within the categorical data, Fisher's exact test and the χ^2 test were used for the binary and non-binary categorical variables, respectively.
- Lasso regression is a regularized linear regression that includes a L1 penalty [11]. Regularization reduces the variance of the features estimates, reducing the absolute values of feature coefficients using penalties. Larger values specify stronger regularization. In this way, the idea of Lasso regression is to optimize the cost function reducing the absolute values of the coefficients. With this approach, features the coefficients of which are equal to zero do not affect the regression result, which means they are not significant and can be discarded.

After preprocessing and feature selection, we obtained the input dataset consisting of 567 rows and 47 columns, with no missing values. This dataset was used as an input for the ML algorithms.

In the **Cross-validation block**, while preparing the dataset for ML, a portion of the data was set aside for the test subset. This subset, following the approach outlined in Section 1, will simulate new patients entering the study that we aim to classify. The remaining data was split into training and validation sets in the 4:1 ratio, adhering to the cross-validation scheme for use with ML algorithms (see Table 1). During the data split, the class ratios were considered, approximately aligning with the 1:2 ratio, signifying that there were twice as many instances of individuals who did not complete the study compared to those who successfully completed it.

Logistic regression is a statistical technique designed for binary classification. It establishes a model for the association between a dependent binary variable and one or more independent variables by estimating the probability of the binary outcome. By employing the logistic function, logistic regression transforms linear combinations of the independent variables, making it well-suited for predicting probabilities and categorizing data into two classes based on a specified threshold.

Random forests represent an ensemble learning approach for tasks like classification, regression, and others. This method constructs numerous decision trees during training. In classification tasks, the random forest's output is the class chosen by the majority of trees. Random decision forests address the tendency of decision trees to overfit their training set.

Perceptron is a classical neural network consisting of two linear layers with the ReLU activation function, the cross-entropy was chosen as a loss function, and the Adam optimizer was used for the training. The perceptron was trained for 500 iterations (called *epochs*).

We built a **Graph Convolutional Network** (a type of GNN) with three graph convolution layers. Such architecture allowed us to take into account information from 3-hop neighbors in the graph (i.e., from vertices connected to a given one with at most 3 edges). As a loss function, we used the negative log likelihood loss. Adam optimization algorithm was used to train the neural network. The network was trained for 500 epochs. The TDA graph for the neural network was built on a prepared dataset, the similarity of patients was measured using correlation between patients' baseline data, and the projection used to construct the TDA graph was centrality. The resulting graph is shown in Figure 10.

Following the principles of cross-validation, the ML models underwent five training iterations on the training set and were subsequently estimated on the validation set. Table 2 showcases the mean values and standard deviations of the model performances post-training. The scores for the models that demonstrated the best performance are marked in bold.

Table 2. Comparison of scores of models trained on a **half of the patients**

Score	Logistic regression	Random Forest	Perceptron	Graph Neural Network
Accuracy	0.710 $\pm$ 0.072	0.686 $\pm$ 0.024	0.724 $\pm$ 0.019	0.742 $\pm$ 0.056
F1	0.526 $\pm$ 0.126	0.315 $\pm$ 0.089	0.501 $\pm$ 0.074	0.614 $\pm$ 0.059
ROC-AUC	0.733 $\pm$ 0.070	0.694 $\pm$ 0.064	0.641 $\pm$ 0.043	0.702 $\pm$ 0.050

The logistic regression model was favored based on the ROC-AUC score, yet it is important to note that Lasso regression was employed during feature selection. Lasso regression identified the most effective features for training the logistic regression model.

On the other hand, the GNN outperformed in terms of both accuracy and F1 scores. Accuracy measures the ratio of correct predictions to the total predictions, while the F1 score emphasizes a balanced approach to identifying true positive cases while mitigating false positives and false negatives. This characteristic renders the F1 score particularly appropriate as a metric in the context of imbalanced datasets. See Figure 10 for explanation.

CONCLUSION

Risk Based Monitoring is an important component of modern clinical trials. The elements of RBM are not fully standardized and it presents a significant challenge for clinical researchers and stakeholders.

In this paper, we discuss two important challenges of RBM: clinical trial site monitoring and ongoing risk assessment related to inclusion of new patients. We propose several effective AI-enhanced tools to address these challenges based on Topological Data Analysis and Graph Neural Networks.

To demonstrate the working prototypes of our solutions, we carry out two experiments. In the first experiment, we detect the systematic bias on one of the clinical trial sites by visual discovery and by numeric characteristic of site's community. In the second experiment, we predict the study completion status of a patient with the help of four models and show that Graph Neural Network, based on a topological model of the data, gives better results than usual machine learning models for two out of three scores.

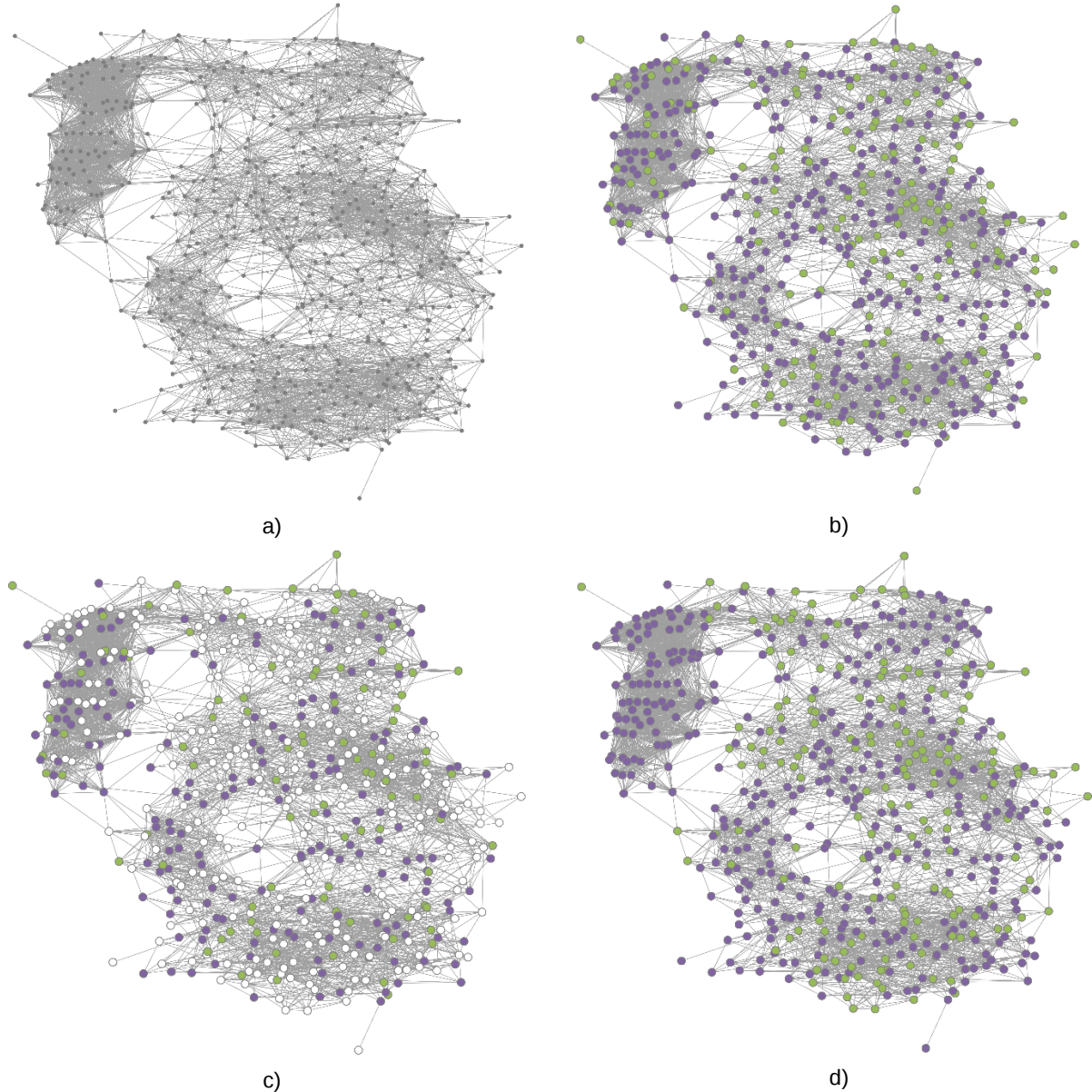


Figure 10. Training GNN on the TDA graph. a) The starting graph (topological model of data) built on the baseline data. b) The starting graph with coloring by completion: green indicates patients who completed the study, purple indicates those who did not. c) According to the proposed solution (Section 1), the white nodes model new-coming patients; 50% of initial data was randomly designated as new-coming (white) patients. d) The best performing GNN predicts labels (green or purple) of the white nodes (chosen using F1 score, see Table 2). This coloring should be compared to b).

REFERENCES

- [1] FDA, "Oversight of Clinical Investigations — A Risk-Based Approach to Monitoring. Guidance for Industry.," 2013. [Online]. Available: <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/oversight-clinical-investigations-risk-based-approach-monitoring>. [Accessed 30 November 2023].
- [2] B. Barnes, N. Stansbury, D. Brown and oth., "Risk-Based Monitoring in Clinical Trials: Past, Present, and Future," *Therapeutic Innovation & Regulatory Science*, vol. 55, p. 899–906, 2021.
- [3] "National Institute on Drug Abuse; NIDA-CSP-1018," [Online]. Available: <https://datashare.nida.nih.gov/study/nida-csp-1018>. [Accessed 30 November 2023].
- [4] G. Carlsson, "Topology and data," *Bulletin of the American Mathematical Society*, vol. 46, no. 2, pp. 255-308, 2009.

- [5] A. Zomorodian, *Topology for Computing*, Cambridge: Cambridge University Press, 2005.
- [6] K. Drach and I. Kotenko, "Using the geometric properties of datasets to analyze the accuracy of COVID-19 forecasting models," in *Pharmaceutical Users Software Exchange 2022 (Phuse US Connect, May 1-4, 2022, Atlanta, USA)*, 2022.
- [7] S. Glushakov, I. Kotenko and K. Drach, "Understanding the spread of COVID-19 in the United States using topology and machine learning for time series," in *Pharmaceutical Users Software Exchange 2020 (Phuse US Connect, March 8-11, 2020, virtual)*, 2020.
- [8] Kipf and Welling, "Semi-Supervised Classification with Graph Convolutional Networks," 2016. [Online]. Available: <https://arxiv.org/abs/1609.02907>.
- [9] W. L. Hamilton, "Graph Representation Learning," *Synthesis Lectures on Artificial Intelligence and Machine Learning*, vol. 14, no. 3, pp. 1-159, 2020.
- [10] F. Errica, M. Podda, D. Bacciu and A. Micheli, "A fair comparison of graph neural networks for graph classification," Published as a conference paper at ICLR, 2020. [Online]. Available: <http://arxiv.org/abs/1912.09893v3>.
- [11] R. Tibshirani, "Regression Shrinkage and Selection via the Lasso," *Journal of the Royal Statistical Society*, vol. 58, no. 1, pp. 267-288, 1996.

ACKNOWLEDGEMENTS

We would like to acknowledge **Oleksandr Leonov** (Kharkiv National University, Ukraine), **Lyudmyla Polyakova** (Kharkiv National University, Ukraine) and **Victoriia Shevtsova** (Intego Group, Ukraine) for being core members of the research team and handling the computational experiment. Without you, this research and the experiment would not have been possible.

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the group of authors at:

Contact: Sergey Glushakov

Company: Intego Group

Address: 2300 Maitland Center Pkwy, Suite 240, Maitland, FL 32751, USA

Work Phone: +1 407 512-1006

Email: sergey.glushakov@intego-group.com

Web: www.intego-group.com

Brand and product names are trademarks of their respective companies.