

<sup>1</sup>Stefana JANICIJEVIC, <sup>2</sup>Aleksandra MIHAJLOVIC, <sup>3</sup>A. KOJANIC

## NEW APPROACH FOR GRAPH BASED DATA MINING MODEL

<sup>1,2</sup>Information Technology School Belgrade, Savski nasip 7, Belgrade, SERBIA

<sup>3</sup>Faculty of Organisational Science Belgrade, Jove Ilica 7, Belgrade, SERBIA

**Abstract:** This paper presents new approach of Communication Network Analysis (CNA) that is interdisciplinary subfield of advanced concept of important Social Network Analysis (SNA). Objects in CNA are members of network discovered as vertices that are linked by edges. Identification of relevant vertices within connected components in telecommunication network graphs, such as influencers are proposed. Beside this result, the algorithm describes behaviour between component members, research interactions between components and telecom services usage. Algorithm is based on a combination of two important machine learning techniques - Classification technique Extreme gradient boosting (XGB) and Graph algorithm that consists from Pruning, K-Neighbourhood, Isolated islands and Centrality measure calculation. This data mining model is used in telecommunication companies as part of marketing strategies and campaign management processes since influencers are awarded for contribution in network services spreading and adopting between members.

**Keywords:** CNA, SNA, Graph algorithm, Machine learning, XGB, Pruning, K-Neighbourhood, Isolated islands, Centrality measures etc.

### 1. INTRODUCTION

#### — Graph theory

Graph theory studies mathematical structures called graphs. The graph is represented by ordered pair  $G = (V; E)$ , where  $V$  is a finite, non-empty set of vertices (tops, vertices), and  $E$  is a set of two-element subsets of set  $V$ , i.e. a set of edges (arcs, branches). Graph is a mathematical structure that consists of vertices (vertices or points) which are connected by edges (links or lines). Graph find its application in many areas, from fundamental mathematics, combinatorics, over data science and machine learning.

Graphs could be undirected and directed, depending on whether the edge connecting the vertices  $u$  and  $v$  is the same as the edge that connects the vertices  $v$  and  $u$ .

We consider neighbours of every vertex, that is, a set of vertices between which there is an edge. Two vertices are adjacent if there is an edge which connects them, that is,  $e = \{u, v\} \in E$ . For graph  $G$  it is denoted set of neighbours that is  $N(v) = \{u \in V \mid (u, v) \in E\}$ . Degree of vertex  $v$  is  $d(v)$  and it is the number of neighbours for vertex  $v$ . The loop in the graph is the edge which connects the vertex with itself. A graph that has no loop nor parallel edge is called prime. Examples of graphs that are investigated in the literature are zero graph, trivial graph, simple graph, undirected graph, directed graph, complete graph, connected graph, K-partite graph, disconnected graph, weighted graph, regular graph, cyclic graph, acyclic graph, star graph, multigraph, planar graph, etc.[1]

There are different variants of representing graphs on a computer, and each of them depends on the nature of the problem being solved and the computer resources at its disposal. Under the term computer resources, it mainly refers to the available memory space. Usually, the vertices of a graph are enumerated with  $0, 1, 2, \dots, n-1$  or  $1, 2, \dots, n$ , where  $n$  is a number of vertices in the graph. The set of edges is represented by one in two ways:

- The adjacency matrix represents the elements that provide information about whether there are edges in between vertices corresponding to the indices of these elements. If the graph is not weighted then the elements of the matrix neighbourhoods only contain information about the existence of an edge between the corresponding vertices. If it is a weighted graph, then the matrix element contains information about the weight of the edges. Formally, the neighbourhood matrix of dimension  $n \times n$  of  $G$  is written as [2]

$$A_G = (a_{ij})_{(i,j) \in V \times V}$$

Example of the adjacency matrix is:

$$\begin{pmatrix} 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 1 \\ 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 \end{pmatrix}$$

- A graph can be represented by a list of neighbours such that each of the vertices of the graph is an element to which a list is formed in which the neighbours of that vertex are placed in the graph. When working

with the sparse graph (a graph in which the number of edges is proportional to the number of vertices), neighbour lists are used because all edges of the graph are described with relatively few data.

However, the list variant is expensive because, for example, a very simple operation such as testing if some two vertices are neighbours requires reviewing the list of neighbours. In practice, the number of such tests can be very large and the performance of the program will be significantly compromised. If they have enough memory, then it is best to present the graph in both ways, to be after they used information either from the neighbourhood matrix or from a list of neighbours, depending on where they were located information that is relevant. The neighbour list is usually made on the basis of distance, so that it is first in lists the vertex neighbour that is closest, and the last vertex the neighbour that is farthest. [2]

Well known problems that are solved by graph algorithms are Hamilton loop problem, Graph  $k$ -colouring problem, Minimal partition problem, Minimum dominating set problem, Independent set problem, Maximal clique problem, Longest path problem, etc.

Graph applications are wide: Internet, telecommunication networks, scheduling, road network, railway network, power transmission system, etc. Certain problems that are successfully solved are optimal transport of goods, optimal energy transfer, protection against accidents, etc.

#### — Social Network Analysis

Social network analysis (SNA) is also a real-life graph application that finds its purpose in Internet and telecom companies. One of the first difficulties is that this problem belongs to the class of NP-complete (NP-difficult problems), since finding an optimal solution for such problem is almost impossible in real life time. Solving NP difficult problem with exact methods is expensive and often impossible to do, so it is usually found some heuristic with the help of which an approximate solution is obtained. That solution may not be optimal, but if the heuristics are good enough the solution will be close to optimal. This problem belongs to the problems of combinatorial optimization. For the most part it is possible to write them in the spirit of mathematical programming.

SNA has arisen from the broader field of graph theory and network systems. Interaction between vertices or objects provide approaches for distinguished algorithms to develop insights from communities to predict behaviour of a network.

Social network analysis is analysis of connected objects in any sort of networks and graphs. The Internet is growing and according to that, social networks are growing through the world. We could say that social networks can be viewed as a set of connected entities. Most of the time we represent it by a collection of vertices and edges. A vertex is an abstraction for a user in the network whereas edges are relations between these users. [3]

SNA applications provide relations entities and people over the globe. Most social networks are based on friendship and family linkage (Facebook, Instagram, Twitter), but there are also professionally networking (LinkedIn) or advanced expert networks (DNA App, Trading View).

#### — Communication Network Analysis

Specificity of telecommunication networks is their large-scale. It is not unusual to find networks with over a million vertices and even more edges. Number of edges is much smaller than the maximal number of edges. A common characteristic of telecommunication graphs is that they are scale free, which means that the distribution of the number of relations of every user follows a power law behaviour. One of main characteristics of a network is big data approach. It is very usual to explore networks with millions and millions of users and billions of edges. Another characteristic is sparseness, since complete subgraphs or maximal cliques are very hard to be found for a significant number of vertices. [4]

The process of defining usage communications based on both the network and the graph theory is known as Communication network analysis. Communication Network Analysis (CNA) is the subfield of Social Network Analysis where vertices and edges are common graph-based approach features such as users and relations between them, but conceptually CNA is different than SNA since, edges are developed according to telecommunication traffic such as voice and SMS system.

CNA and SNA methods have become a powerful tool for studying networks in various issues of telco field.

#### — Big Data

Big data occurs in scientific, engineering, and commercial applications. These include government, military, telecommunications, medical, biotechnology, astrology, ecological, pharmaceutical systems, as well as many others. Some of the wide range of challenges that are associated with big data sets are data storage, data warehousing, data compression, visualization, information insights, clustering, pattern recognition, algorithm performing. Addressing these issues requires special interdisciplinary efforts in developing sustainable techniques. Big databases imply with them complex and content issues and thus represent a challenging field to address. In many cases, big data sets are represented as large graphs with certain attributes associated with

vertices and edges. These attributes may contain specific information that characterizes the information. Analysing the structure of such a graph is important to understand the structural characteristics of the application which is represented, such as improving information retrieval and memory organization. Current trends in analysing large graphs are developed mostly on market graphs, telecommunication graphs and Internet graphs. [5]

## 2. PROBLEM FORMULATION

CNA objective is identification of influencers in a network based on the number of connected users and based on the traffic between them, so this paper explores the identification of influencers in the telecommunications network. Regarding this, there is methodology for identification of the influencers proposed. This research is based on CNA methodology. The main pillars of CAN are voice, SMS and data traffic.

The CNA uses the msisdn of the operator which is a record of data produced by the telephone exchange. This data record documents the details of each Telecom transaction (VOICE, SMS, MMS, GPRS, internet using...) that passes through mobile devices. In addition to msisdn, CNA uses other data sources such as customer data to analyse users' relationships. Linking this information together with CNA provides better insights and values that affect the revenue and the customer satisfaction [3]. Msisdn and usage data collected for 4 months. It was implemented ETL (Extract, Transform and Load) to data.

It is used CNA to analyse relationships among interacting vertices (users, products ...etc.), therefore we discovered the structure of individuals or organizations. Relationships in a network can be directional or non-directional. When we talk about directional relationship, we could say that one person is the initiator (or (source) basis of the relationship) while the other is the one who receives (receiver) (or destination of the relationship). Weight indicating the strength of the relationship can be added.

In CNA most important is the concept that present communications over network flow, such as voice and messaging system. Networks are used to represent group of users or vertices of a network with their linkage characteristics or edges.

## 3. RELATED WORK

In the paper [6] the authors are trying to explain what exactly is SNA analysis in Telecom data. In this paper, networks and SNA concepts were applied using Telecom data such as call detail records and customers data in order to construct a weighted graph in which each relation carries a different weight, representing how close two users are to each other. SNA is used to explore the Telecom network and calculate the centrality measures. Centrality measures help to determine the vertex importance in the network.

Finding Multi – SIM users within the same operator or across different operators presents another important concern to Telecom companies because it allows to improve campaigns and churn models. The paper is based on a real dataset of 3 months msisdn and customer data provided by a local Telecom operator. Accuracy of 85% was achieved for users from different operators and 92% for users from the same operator.

In the paper [7] the authors survey recent advances in the study of influencer identifications develop from big data perspectives, and present state-of-the-art solutions of vertices whose removal would breakdown the network. They proposed survey methods to locate the essential vertices that are capable of shaping global dynamics with either continuous or discontinuous phase transitions. The solution implies recommender system in social networks.

In the paper [8] the authors argue that data science is a “concept to unify statistic, data analysis and their related methods” in order to “understand and analyse actual phenomena” with data. The principal idea in designing different marketing strategies is to identify the influencers in the network's communication. Targeting influencers usually leads to a vast spread of the centrality measures to identify and assign an influence score to each other. Higher score – better influencer.

The main goal is to find the most influential users among given pair of users. This could be scaled over the entire network. Data is used for mining model to score the test data and predict the influential user among the given pair of users

## 4. MODEL BUILDING AND VALIDATION

### — Methodology

CNA is a model that explores a user's telecommunications network through the behaviour of each user individually, based on the number of users with whom a specific user comes into contact, but also based on the intensity of communications that user has in the network. The model highlights influencer users within prominent groups of connected users. In addition, the model provides insights into the mutual relations of network members, but also insights into the relations of users towards telecommunication services.

Main methodology of influencer identification consists from combination of two approaches: Classification algorithm XGB and Graph based algorithm which considers *K*-Neighborhood, Pruning, Isolated islands and Centrality measures calculation.



Together, they are forming machine learning system for CNA. One of the goals involves modelling the mechanisms that underline human learning. It was developed learning algorithms that are generally consistent with knowledge of the human cognitive architecture and that are also designed to explain specific observed learning behaviours. This machine learning can transform training data into knowledge using algorithm. The phases of the CNA model are:

- ≡ Data preparation
- ≡ Model development
- ≡ Model evaluation

#### — Data preparation

Data preparation is a very demanding and important process. It implies scrubbing, wrangling, munging and auditing which is performed over tables for A number and for AB numbers communication. These two tables are final and they contain independent and dependent variables that enter the model, which means, selection of variables, deleting redundant variables, working with missing values, editing outliers, normalization of data, etc. Table A number consist from the predictors such as voice count, SMS count, voice duration, gprs count, long voice duration, short voice duration, etc., while the target variable is  $y$ .

It is defined according to formula:

$y = \{1, \text{ where minimum 6 relevant predictors are achieved for at least 50\% of value } 0, \text{ otherwise } \}$   
where:

$$x_i \in X, X \in \left\{ \begin{array}{l} \text{voi}_{\text{cnt}}, \text{voi}_{\text{dur}}, \text{sms}_{\text{cnt}}, \text{data}_{\text{cnt}}, \text{voishort}_{\text{cnt}}, \text{voilong}_{\text{cnt}}, \text{voishort}_{\text{dur}}, \\ \text{voilong}_{\text{dur}}, \text{voiin}_{\text{cnt}}, \text{smsin}_{\text{cnt}}, \text{voiout}_{\text{cnt}}, \text{smsout}_{\text{cnt}} \end{array} \right\}$$

Target variable is used for calculation potentially most valuable vertices - influencer vertices with  $y = 1$ , so we positively labelled every A number that has more communication than variable median for at least 6 variables predictors.

#### — Model development - Extreme gradient boosting (XGB)

The core of extreme gradient boosting (XGB) itself is the group algorithm based on the gradient boosting tree. Gradient boosting is an algorithm of boosting in the ensemble algorithm. XGB algorithm is an efficient implementation version of gradient boosting algorithm. Because of its excellent efficiency in application practice, it is a widely-praised technique in industry. XGB is similar to gradient boosting decision tree (GBDT) and is based on the classification and regression tree theory. It is able to build multiple weak evaluators on the data and then summarizes the modelling results of the weak evaluators. In parallel, the XGB model can effectively deal with regression and classification problems to obtain better performance than a single one. [9]

Gradient boosting involves the creation and addition of decision trees sequentially, each attempting to correct the mistakes of the learners that came before it. Most implementations of gradient boosting are configured by default with a relatively small number of trees, it is because adding more trees beyond a limit does not improve the performance of the model. The reason is in the way that the boosted tree model is constructed, sequentially where each new tree attempts to model and correct for the errors made by the sequence of previous trees. [9] The objective function is composed of two parts. The first part is used to measure the difference between the predicted value and the actual value (represents the deviation of the model), and the other part is the regularization term (the variance of the control model). The prediction accuracy of the model is determined by the deviation and variance of the model.

XGB aims to target and predict a possible influencer. It is used to reduce the user base to a size that is optimal for management, transformation, and manipulation.

1. Initialize  $f_0(x) = \arg \arg \min_{\gamma} \sum_{i=1}^N L(y_i, \gamma)$ .

2. For  $m = 1$  to  $M$ :

a. For  $i = 1, 2, \dots, N$  compute

$$r_{im} = - \left[ \frac{\delta L(y_i, f(x_i))}{\delta f(x_i)} \right]_{f=f_{m-1}}$$

b. Fit a regression tree to the targets  $r_{im}$  giving terminal regions

$$R_{jm}, j = 1, 2, \dots, J_m.$$

c. For  $j = 1, 2, \dots, J_m$  compute

$$\gamma_{jm} = \arg \arg \min_{\gamma} \sum_{x_i \in R_{jm}} L(y_i, f_{m-1}(x_i) + \gamma).$$

d. Update  $f_m(x) = f_{m-1}(x) + \sum_{j=1}^{J_m} \gamma_{jm} I(x \in R_{jm})$ .

3. Output  $\hat{f}(x) = f_M(x)$ .

The goal of XGB is to eliminate vertices with low values of total traffic, and thus to reduce the dimension of the initial graph. Based on the target variable, XGB creates a decision about which user is an influencer. The final set of users is selected in a new table which is the basis for creating a graph.

XGB is designed for speed and performance. Two reasons to use XGB are execution speed and model achievement. In addition, XGB has proven to be great combination of software and hardware optimization techniques to achieve superior results using fewer computing resources in the shortest amount of time.

XGB consists of following steps:

- ≡ Parallelism in tree construction
- ≡ Pruning using the DFS algorithm
- ≡ Computing in external memory
- ≡ Regularization due to reduced overfitting
- ≡ Efficient handling of missing data
- ≡ Built-in cross validation

Model evaluation parameters:

- ≡ Coincidence matrices
- ≡ Performance evaluation
- ≡ Evaluation metric (AUC&Gini)

#### — Model development - Graph algorithm

Graphs are a structure and a powerful tool for modelling and analysing data such as telco and social networks, websites and links, as well as vehicle locations and routes. When there is a set of objects that are interconnected, then they can be represented by graphs.

The graph algorithm aims to observe the interrelationships between influencers and other users on the basis of which the final set of influencers is determined. A graph algorithm is used to highlight measures of significance of each vertex in interaction with other vertices.

The vertices with their edges that formed graph are users that the XGB model has classified as potentially influencers with their connections.

Based on the created table, the weight of the edges between A and B numbers is calculated, defining the calibration parameter for the call length, the number of calls and the number of SMS communications. The greatest weight is assigned to the length of the call, and the least to the number of SMS in accordance with the distribution of the database

$$w = (\alpha * 6/8 * \text{row}['\text{VOI\_DUR}'] + \alpha * 2/8 * \text{row}['\text{SMS\_CNT}'] + (1 - \alpha) * \text{row}['\text{VOI\_CNT}']) \quad (3)$$

where  $\alpha = 0.5$

#### ≡ K-Neighborhood

The  $K$ -neighborhood method determines the vertex degree, indegree, and outdegree of the 1st level neighbours. They are calculated for each vertex based on this initial graph.

#### ≡ Pruning

The problem of determining the proper size of an artificial neural network is recognized to be crucial, especially for its practical implementation in such important issues as learning and generalization. One popular approach for tackling this problem is commonly known as pruning and it consists of training a larger than necessary network and then removing unnecessary weights/vertices. The algorithm also provides a simple criterion for choosing the units to be removed, which has proved to work well in practice. The results obtained over various test problems demonstrate that effectiveness of the proposed approach. [10]

Data set consists from vertex-level information for all vertices that are still included in the graph after removal of weaker edges. Weak edge considers all edges smaller than median weights of all edges.

#### ≡ Isolated islands

Isolated island gives an answer to the standard question: “Counting the number of connected components in an undetected graph”. A connected component of an undirected graph is a subgraph in which every two vertices are connected to each other by a path, and which is connected to no other vertices outside the subgraph. A group of connected for example voice calls forms an island. [11]

The graph is reduced to components, eliminating weak components with the remaining weaker vertices. Using the Isolated Islands method, the graph is divided into mutually isolated components of optimal size while less isolated components are discarded. The graph is divided into connected components and further separation is performed, so that no subgraph contains less than the number of vertices defined through minimum list length. The Isolated islands method separates the components at the level of bond strength, so that the edges with the least weight are eliminated first. Within each stronger component, the central vertices (i.e. the carriers of communications in the group) are isolated. BFS algorithm reduce graph on strongly related components. In BFS, we start from a specific vertex and explore as far along each level of vertices as possible

before re-searching backwards. We also need to monitor visited vertices. When we implement BFS, we use the stack data structure to support reverse lookup.

### ≡ Centrality measures

Metrics are the final part of the model that decides who is the influencer among the vertices. Metrics are calculated for each vertex, and, we could say the metrics calculated the importance of each vertex. Typical metrics that are used the closeness and the betweenness. They are based on shortest paths and distances between vertices. [12]

$$\text{CLOSENESS} = \sum_{y \neq x} \frac{1}{d(y,x)}, \quad \text{'BETWEENNESS} = \sum_{s \neq u \neq t} \frac{\sigma_{st}(u)}{\sigma_{st}}$$

Targeting influencers usually leads to a vast spread of the centrality measures to identify and assign an influence score to each other. Higher score means better influencer.

The most important vertices are selected by sorting and comparing the highest values of these metrics, based on previously isolated components for each vertex.

The selection of influencers is made on the basis of the Table 1. These are the strongest users in the network, according to degree and who have many outgoing communications and many incoming communications. In production, the model is realized by hundreds of large groups of connected users, which are shown here.

Table 1. Table of final result after isolated islands and centrality measures calculations

|     | group id | msisdn    | degree | closeness   | betweenness |
|-----|----------|-----------|--------|-------------|-------------|
| 604 | 301      | 38*****63 | 1930   | 0.002142531 | 1861383.467 |
| 521 | 902      | 38*****27 | 1628   | 0.002141762 | 1324349.875 |
| 868 | 87       | 38*****30 | 1415   | 0.002141815 | 939444.9204 |
| 58  | 149      | 38*****24 | 1311   | 0.002142061 | 856551.797  |
| 871 | 406      | 38*****10 | 849    | 0.00214199  | 342381.3873 |
| 59  | 640      | 38*****00 | 767    | 0.002141245 | 293623.2    |

### ≡ Pseudocode

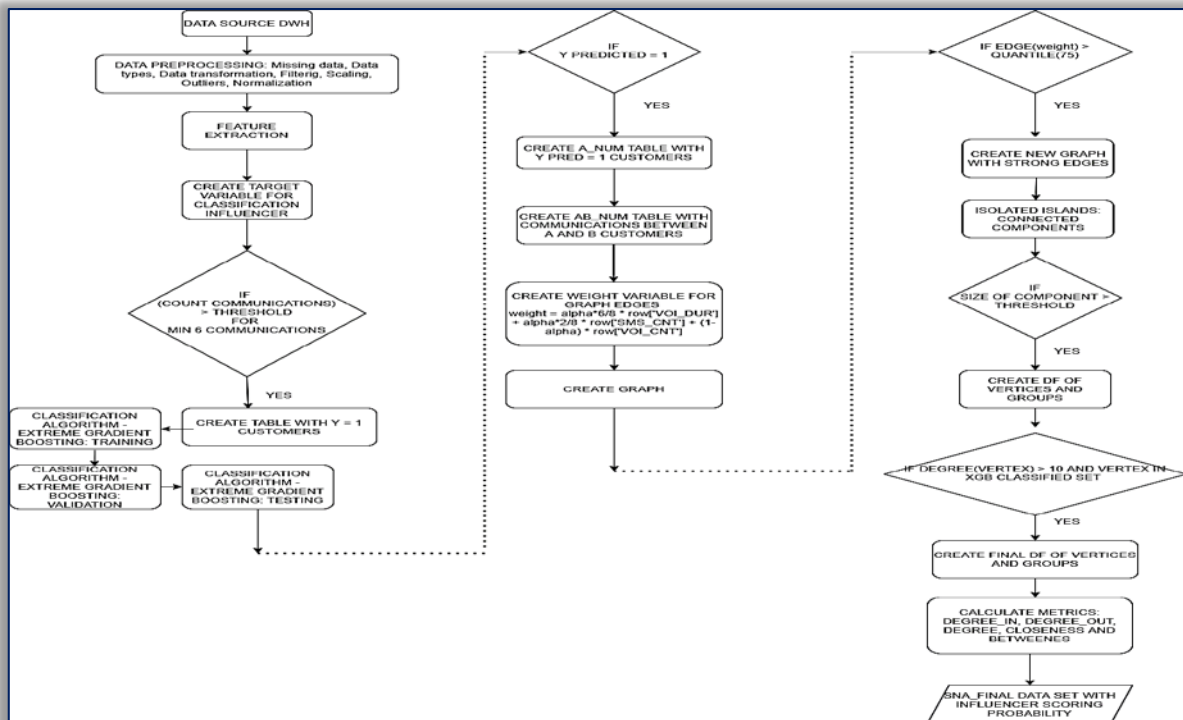


Figure 1. Algorithm CNA

1. Initialize  $f_M(x)$ , where  $f_M \in f^N$  space of classified subscribers
2. Create weighted graph and adjacency matrix  
for  $i = 1$  to  $M$   
for  $j = 1$  to  $M$   
 $x[i,j] = w(i,j)$   
iterate through all connections
3. Calculate  $K$  – Neighborhood  
for  $i = 1$  to  $M$   
for  $j = 1$  to  $M$   
if  $(G \rightarrow x[i][j] == 1)$   
degree ++
4. Pruning  
Input: A weighted graph  $G = (V, E)$



Output: Subgraph  $H \subset G$

1: Sort edges  $E$  by weights in an ascending order.

2:  $F \leftarrow E$

3:  $n \leftarrow (|E| - (|V| - 1))$

4: { Iteratively prune the weakest edge which does not cut the graph }

5:  $i \leftarrow 1, j \leftarrow 1$  {  $j$  is an index to the sorted list of edges }

6: while  $i \leq n$  do

7: if  $C(u, v; F \setminus \{e_j\})$  is not  $\infty$  then

8:  $F \leftarrow F \setminus \{e_j\}$

9:  $i \leftarrow i + 1$

10:  $j \leftarrow j + 1$

11: Return  $H = (V, F)$

##### 5. Isolated islands

for  $i = 1$  to  $M$

for  $j = 1$  to  $N$

if  $(x[i][j] \&\& !visited[i][j])$  {

// visited yet, then new island found

// Visit all cells in this island

BFS( $x, i, j, visited$ )

// and increment island count

++ count

##### 6. Calculate centrality measures

## 5. RESULTS AND INTERPRETATION

Classification method is trained on 1130925 unique MSISDN data set. Target variable  $y$  is equal 0 on 601350 MSISDN and 1 on 529575 MSISDN. After classification, it was calculated contingency matrix distributed on TP, FP, TN and FN as follows:

Table 1. Extreme Gradient Boosting  $Y$  - Train set

| XGB - Y/TRAIN | 0      | 1      |
|---------------|--------|--------|
| 0             | 537488 | 63862  |
| 1             | 58725  | 470850 |

Table 4. XGB results for different parameters combinations on train data set

| Classification                                                                 | AUC   | GINI  | Accuracy | Precision |
|--------------------------------------------------------------------------------|-------|-------|----------|-----------|
| XGB (max depth = 10, boost = 5, min child weight = 0.1, max delta step = 0.7)  | 0.698 | 0.720 | 63.3%    | 74.83%    |
| XGB (max depth = 20, boost = 10, min child weight = 1.0, max delta step = 0.5) | 0.713 | 0.731 | 65.7%    | 76.24%    |
| XGB (max depth = 25, boost = 10, min child weight = 1.0, max delta step = 0.6) | 0.762 | 0.791 | 69.4%    | 79.14%    |
| XGB (max depth = 35, boost = 5, min child weight = 1.1, max delta step = 0.4)  | 0.    | 0.837 | 74.3%    | 82.51%    |
| XGB (max depth = 40, boost = 10, min child weight = 1.1, max delta step = 0.2) | 0.904 | 0.802 | 83.16%   | 88.91%    |

Best approach is last combination of parameters for this model. Parameters are:

- ≡ Tree method: basic,
- ≡ Number of boosts: 10,
- ≡ Max depth: 40,
- ≡ Min child weight: 1.1,
- ≡ Max delta step: 0.2
- ≡ Objective function: binary logistic,
- ≡ Sub sample: 1.0,
- ≡ Eta: 0.3,
- ≡ Gamma: 0.0,
- ≡ Alpha: 0.0
- ≡ Scale pos weight: 1.0

After selection of best combination of parameters on training data, it is applied model on evaluation set. Number of unique MSISDN was 2038410.

Comparison between total population and  $y = 1$  population according to features distribution:

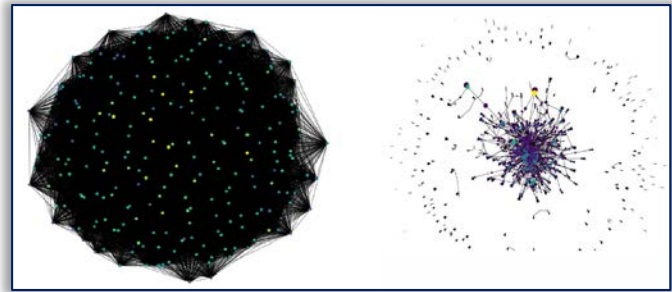


Figure 1. These are population graph and sample graph

Table 2. Extreme Gradient Boosting  $Y$  - Eval set

| XGB - $Y$ /EVAL | 0      | 1      |
|-----------------|--------|--------|
| 0               | 971640 | 132346 |
| 1               | 111081 | 823343 |

Table 3. Comparing  $XGB$  with  $y$

|         |           |        |
|---------|-----------|--------|
| Correct | 1,794,983 | 88.06% |
| Wrong   | 243,427   | 11.94% |
| Total   | 2,038,410 |        |

Table 3. Evaluation Metrics

| Model      | AUC   | Gini  |
|------------|-------|-------|
| %XGB - $y$ | 0.907 | 0.815 |

For the evaluation set it is achieved 16234 strong vertices (influencers) in 2341 isolated components.

## 6. CONCLUSIONS

Global graph search is currently the most optimal method of processing large data; all large systems such as the Internet and telecommunications use graph-based algorithms since other heuristics are almost impossible without it. The CNA model provides and regularly submits a list of users and groups sorted by the highest probability. A list of users from the most influential to the least influential according to  $K$ -Neighbourhood degree in the selected group of influencers is submitted. For related groups, information on group size and strength is provided, which is also sorted by probability from largest to smallest.

The results of the model present possibilities for telco industry and for advanced analytics.

The model approach is very well applied for the big data concept for HDFS but also for DWH concept. It is possible to research telco sparse graphs and calculate different thresholds for pruning weak edges and threshold for the number of components. Based on that table, it is possible to form reports at the aggregate level, as well as at the individual level.

The questions that the model answers are:

- Who are the most influential users - central network users?
- Which groups can be spotted in the network?
- In what way is the network divided into less loosely connected groups?
- How is the network evolving?
- Will the network be maintained or cease to exist?
- How do ideas-information spread through the network?

Note:

This paper is based on the paper presented at International Conference on Applied Sciences – ICAS 2021, organized by University Politehnica Timisoara – Faculty of Engineering Hunedoara (ROMANIA) and University of Banja Luka, Faculty of Mechanical Engineering Banja Luka (BOSNIA & HERZEGOVINA), in Hunedoara, ROMANIA, 12–14 May, 2021.

## References

- [1] Diestel R 2000 Graph Theory, Second Edition, Springer-Verlag, New York.
- [2] Janicijevic S 2016 Variable Formulation and Neighbourhood Search Methods for the Maximum Clique Problem in Graph, Ph.D. thesis, University of Novi Sad, Faculty of Technical Sciences
- [3] Amer M S 2015 Social network analysis framework in Telecom, Int J Syst Appl Eng Dev 9.1, 201-205.
- [4] Kihl M et al. 2010 Traffic analysis and characterization of Internet user behaviour, International Congress on Ultra Modern Telecommunications and Control Systems. IEEE.
- [5] Ahmed H and Ismail M A 2020 Towards a Novel Framework for Automatic Big Data Detection, in IEEE Access, vol. 8, pp. 186304-186322
- [6] Molhem Al et al. 2019 Social network analysis in Telecom data, Journal of Big Data, 6.1, 1-17.
- [7] Pham T A N et al. 2015 A general graph-based model for recommendation in event-based social networks, IEEE 31st international conference on data engineering. IEEE.
- [8] Bethu S et al. 2018 Data science: Identifying influencers in social networks, Periodicals of Engineering and Natural Sciences 6.1, 215-228.
- [9] Bhattacharya S et al. 2020 A novel PCA-firefly based XGBoost classification model for intrusion detection in networks using GPU, Electronics 9.
- [10] Zhang, Mingyang et al. 2019 Graph pruning for model compression, arXiv preprint arXiv:1911.09817.
- [11] Hanneman R A 2005 et al Introduction to social network methods, online text.
- [12] Zhang J et al. 2017 Degree centrality, betweenness centrality, and closeness centrality in social network, 2nd International Conference on Modelling, Simulation and Applied Mathematics (MSAM2017), Atlantis Press.

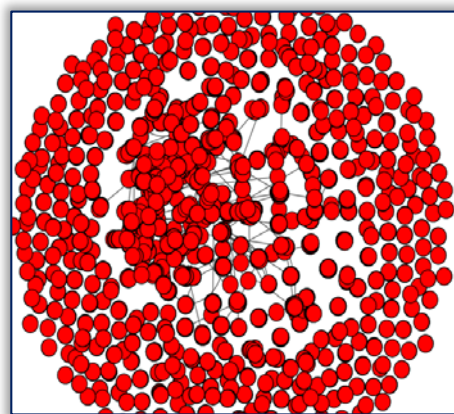


Figure 2. This is influencer graph