



Universiteit
Leiden
The Netherlands

Applying data mining in telecommunications

Radosavljevik, D.

Citation

Radosavljevik, D. (2017, December 11). *Applying data mining in telecommunications*. Retrieved from <https://hdl.handle.net/1887/57982>

Version: Not Applicable (or Unknown)

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/57982>

Note: To cite this publication please use the final published version (if applicable).

Cover Page



Universiteit Leiden



The handle <http://hdl.handle.net/1887/57982> holds various files of this Leiden University dissertation.

Author: Radosavljevik, D.

Title: Applying data mining in telecommunications

Issue Date: 2017-12-11

Chapter 2

The Impact of Experimental Setup in Prepaid Churn Prediction for Mobile Telecommunications: What to Predict, for Whom and Does the Customer Experience Matter?

Radosavljevik, D., van der Putten, P., & Larsen, K.K.

Based on a publication in Transactions on Machine Learning and Data Mining 3 (2), pp. 80-99 (2010)

Prepaid customers in mobile telecommunications are not bound by a contract and can therefore change operators ("churn") at their convenience and without notification. This makes the task of predicting churn both challenging and financially rewarding. This chapter presents an exploratory, real world study of prepaid churn modeling by varying the experimental setup on three dimensions: data, outcome definition and population sample. First, we add Customer Experience Management (CEM) data to data traditionally used for prepaid churn prediction. Second, we vary the outcome definition of prepaid churn. Third, we alter the sample of the customers included, based on their inactivity period at the time of recording. While adding CEM parameters and varying the sample did not influence the predictability of churn, a particular change in the outcome definition had a substantial influence.

2.1 Introduction

The problem of churn, or loss of a client to a competitor, is a problem facing almost every company in any given industry. This phenomenon is a major source of financial loss, because it is generally much more expensive to attract new customers than it is to retain and sell to existing ones. Therefore, churn is important to manage, especially in industries characterized by strong competition and saturated markets, such as the mobile telecom industry. Prepaid mobile phone subscribers are not bound by a contract; therefore, they can churn at their convenience and without notification, which makes the task of predicting the likelihood and moment of churn very important.

The first step in managing churn is identifying the customers with high propensity to churn. Published papers on churn modeling on private data sets in mobile telecommunications are relatively scarce, due to the commercially sensitive nature of the problem. This is even more evident for papers based on European mobile telecom operators' data. Most of the available research relates to postpaid churn (Au, Chan and Yao, 2003; Datta, Masand, Mani and Li, 2000; Ferreira, Vellasco, Pacheco and Barbosa, 2004; Hung, Yen and Wang, 2006; Hwang, Jung and Suh, 2004; Kim and Yoon, 2004; Lemmens and Croux, 2006; Lima, Mues and Baesens, 2009; Mozer, Wolniewicz, Johnson and Kaushansky, 1999; Neslin, Gupta, Kamakura, Lu and Mason, 2006; Wei and Chiu, 2002). Even fewer studies are available for prepaid churn in this industry (Alberts, 2006; Archaux, Martin and Khenchaf, 2004; Dasgupta et al., 2008). The majority of these assumed a fixed, single experimental setup in terms of outcome definition, population and data parameters. Their focus was mostly on the data mining algorithm used or algorithmic tuning. In order to understand prepaid churn modeling better, we decided to take an end-to-end view and test different experimental setups by varying on three dimensions: data, population sample and outcome definition (Radosavljevik, van der Putten and Kyllesbech Larsen, 2010b). We used standard data mining algorithms, decision trees (in their standard form) and logistic regression (Pegasytems, 2008; Witten and Frank, 2005), because it was

not our objective to determine the impact of the algorithm; research on data mining algorithms is abundant. Furthermore, from a business perspective, the output of either of these algorithms (the model) is very easy to interpret and communicate to parties that do not have extensive experience with data mining (e.g. business managers). This quality is even more evident in the case of decision trees, which have a very intuitive graphical representation. Finally, as our experiments have shown, the choice of algorithm is a minor factor of influence compared to the other dimensions mentioned.

Making new types of data available for modeling may improve model performance. We provide an overall framework for measuring customer experience, and in the experiments we investigate the added value of customer level service quality metrics (dropped calls, call setup duration, SMS delivery failure rate etc.).

The definition of population sample and outcome is primarily driven by the business objectives, i.e. how the model will be used. In general, operators discontinue the service and consider prepaid customers churned from an administrative perspective if they have not had an activity (e.g. made a call, sent an SMS, etc.) for a period of six months. However, from a churn prediction and marketing perspective such a long period is not very useful. Usually, the customer has churned long before that time. Therefore, it is much better to use a shorter period of inactivity as an outcome definition of churn. Our variations in population sample were based on the customers' inactivity period at the moment of recording. The dilemma here is, if customers have been inactive already for one month at the time of recording, are they retainable, or have they simply thrown away the SIM card? What should be the maximum period of inactivity allowed at the time of recording?

Our choice for the research question for this chapter was driven by the multiple possibilities and choices that exist when creating an experimental setup for prepaid churn modeling:

Which one of the following variations in the experimental setup has the highest influence on the performance of prepaid churn prediction models: adding Customer Experience Management data, altering the characteristics of the sample or changing the outcome definition?

The main contributions of this chapter are as follows. First, we provide a novel theoretical business framework for measuring customer experience in mobile telecommunications. Second, to our knowledge we are the first to investigate the added value of service quality oriented customer experience data for predicting prepaid churn. Third, we explore the impact of changing the criterion for the population sample. Fourth, we propose different outcome definitions of prepaid churn and measure the impact of changing the outcome definition on model performance. Finally, this research was conducted by one of the leading mobile telecom operators in Europe, driven by high priority business needs and using large amounts of customer data, which gives it a real world dimension. Given the highly competitive, confidential and strategic nature of mobile telecommunications churn management, real world results (based on the European telecom market) are not often available in literature.

The remainder of this chapter is structured as follows: In Section 2.2 we present how (prepaid) churn modeling has been performed in the past. Section 2.3 provides an overview of Customer Experience Management. In Section 2.4, the research setup and the experimental design are discussed, followed by results in Section 2.5. We end the chapter with a discussion (Section 2.6) and a conclusion (Section 2.7).

2.2 (Prepaid) Churn Modeling

In this section, we will present how modeling churn in mobile telecommunications has been addressed in previous research.

As mentioned in the introduction, the majority of papers on wireless churn assume a single experimental setup and deal with postpaid churn (Au et al., 2003; Datta et al., 2000; Ferreira et al., 2004; Hung et al., 2006; Hwang et al., 2004; Kim and Yoon, 2004; Lemmens and Croux, 2006; Lima et al., 2009; Mozer et al., 1999; Neslin et al., 2006; Wei and Chiu, 2002). Some of these papers did not explicitly state that they are related to postpaid churn; nonetheless, this conclusion can be made from the data sets they used, which contained demographic information and contract data. These parameters are unavailable for prepaid subscribers, as there is no contract in the real sense of the word. This is the key difference between prepaid and postpaid churn prediction, even though quite a few similarities exist.

Prepaid churn modeling is typically done by performing analysis on aggregated Call Detail Records (CDRs), but additional factors, such as subscription details (e.g. duration, subscription or discontinuation of services) and handset data, are often included. Parameters used in previous papers are as follows.

Archaux et al. (2004) took into account the following data: invoicing data, such as the amounts refilled by clients or amounts withdrawn by companies for the subscribed services and options; data relating to usage, such as the total number of calls, the share of local, national or international calls, consumption peaks and averages; data relating to the subscription, such as date of beginning (age of the subscription), the current tariff plan and the number of different plans the client used; data relating to application and cancellation of services; data related to the current and the previous profitability of the clients. Alberts (2006) selected data related to usage and billing information only: average duration of a single incoming and outgoing call, ratio between outgoing and incoming call durations, sum of incoming and outgoing revenues, the current and the maximum number of successive non-usage months, cumulative number of recharges, average duration of incoming and outgoing calls over all past months, the number of months since the last voicemail call, the number of months since the last recharge, the last recharge amount, the average recharge amount of all past months, etc.

The algorithm is the focus of many of the data mining efforts related to churn. Papers related to postpaid churn have investigated the performance of standard data mining algorithms, such as decision trees (Au et al., 2003; Ferreira et al., 2004; Hung

et al., 2006; Hwang et al., 2004; Lemmens and Croux, 2006; Lima et al., 2009; Mozer et al., 1999; Neslin et al., 2006; Verbeke, Martens, Mues and Baesens, 2011; Wei and Chiu, 2002), logistic regression (Hwang et al., 2004; Lemmens and Croux, 2006; Lima et al., 2009; Mozer et al., 1999; Neslin et al., 2006; Verbeke et al., 2011), neural networks (Au et al., 2003; Datta et al., 2000; Ferreira et al., 2004; Hung et al., 2006; Hwang et al., 2004; Mozer et al., 1999; Neslin et al., 2006), Bayesian classifiers (Neslin et al., 2006) and support vector machines (Verbeke et al., 2011), and compared them to novel approaches.

In prepaid churn modeling, Archaux et al. (2004) compared the performance of models built using neural networks and support vector machines, while Alberts (2006) compared models built using survival analysis with models built using decision trees. To summarize, in research literature there is a lot of emphasis on algorithm testing and tuning, whereas in real world practice this is a smaller piece of the puzzle with relatively modest impact. Hence, we decided to focus more on experimental setup.

The influence of Social Networks on (prepaid) churn has entered the focus of interest of both business and academic communities in the period after 2008 (Dasgupta et al., 2008; Richter, Yom-Tov and Slonim, 2010; Wang, Cong, Song and Xie, 2010). These papers strive to answer the question whether the decision of a subscriber to churn is dependent on the existing members of the community with whom the subscriber has a relationship. Dasgupta et al. (2008) showed that diffusion models built on call graphs (used for identification of Social Networks) have superior performance to their baseline model. At the time of publication, this approach was considered a very progressive novelty, with relation to both algorithm and data. However, it is our opinion that these results are biased by the fact that their baseline model is constructed on CDR data alone, without including other important factors, such as the handset, prepaid account balance or inactivity period. In our opinion, if these variables would have been added in their baseline model, the improvements of their approach would not have been as high.

As we will show in our experiments, combining past behavior (extent of usage) with current factors (e.g. prepaid account balance, phone type, price plan etc.) is crucial for predicting future behavior (churn). In general, there is a gap between traditional methods that do not take the social networks into account and social networks techniques that only mine the network. Therefore, in chapter 3 we investigate the value of hybrid approaches, combining these two methods.

In addition, we consider the origin of the data used in previous research. As telecom markets differ in various parts of the world, it is expected that the churners' patterns would differ as well. The US telecom market has been addressed by several papers (Datta et al., 2000; Lemmens and Croux, 2006; Lima et al., 2009). There is also research based on Asian markets (Au et al., 2003; Hwang et al., 2004; Kim and Yoon, 2004; Wang et al., 2010). Ferreira et al. (2004) analyzed the Brazilian market. Other researchers did not reveal the location of the operator (Dasgupta et al., 2008; Richter

et al., 2010) or used publicly available data sets (Lima et al., 2009; Verbeke et al., 2011). There is a visible lack of papers focusing on the European market (except Alberts, 2006).

Last, but not least, we would like to mention the issue of defining prepaid churn, or the lack of a clear definition of a churned prepaid customer. To the best knowledge of the authors only Alberts (2006) defined in detail which prepaid customers are marked as churned. We find the outcome definition to be of high importance and we address this issue in depth in Subsections 2.4.2 and 2.4.5, where we also propose our own definition(s).

2.3 Customer Experience Management

To put our churn management activities in perspective, we see them as part of a wider ranging company effort to manage and optimize the customer experience. Products and services (e.g. telecommunication services) tend to become commodities over time; therefore, managing the customers' experience can be a source of sustainable competitive advantage (Pine and Gilmore, 1999). Pine and Gilmore (1999) stated that experiences are as distinct from services as services are from goods, thus arguing that authentic experiences are the next evolutionary step in creating customer value. Smith and Wheeler (2002) claimed that branded customer experience drives customer loyalty and profitability. Customer Experience Management (CEM) is the process of strategically managing and optimizing these experiences across all customer touch-points and channels, in interactions that are either customer initiated or company driven, direct or intermediated (Meyer and Schwager, 2007; Schmitt, 2003).

CEM is closely related to its predecessor Customer Relationship Management (CRM). The distinction is somewhat artificial, but one could argue that CRM focuses more on providing a 360 degrees view of the current relationship and managing customer processes efficiently, whereas CEM provides tools and methodologies to understand, improve and extend the relationship by optimizing the overall customer experience. For mobile telecommunication providers (and many other businesses) the key areas within an overall CEM strategy are customer acquisition, revenue stimulation for existing customers and customer retention management (churn management). Churn models, in particular, can be important components of overall strategies that make these predictions actionable to drive intelligent interactions and experiences.

2.3.1 Measuring the Customer Experience

The focus of this chapter is on building better prepaid churn models, both from a predictive power perspective, as well as from the point of view of predicting behavior that is actionable and makes business sense. For instance, predicting short term inactivity for a base that includes many already dormant accounts will result in a

powerful, yet not very useful model. Embedding the churn models in larger strategies to optimize experience in real time is more the scope of future research. From a CEM perspective, our topic of interest is measuring past customer experience as a potential input to churn models. We have created a theoretical business framework for measuring customer experience in mobile telecommunications in ideal circumstances (Figure 2.1).

To the best knowledge of the authors, this is a first attempt to create a framework of this kind and purpose. This framework provides a conceptual roadmap and direction for the coming years for the mobile telecom operator where this research was performed. In the short term, it is very challenging to measure most of these dimensions reliably. The framework contains three levels: Aspects, KPIs and Data Sources. The Aspects level represents the various aspects of customer experience in mobile telecommunications. The KPIs level defines possible measurements for the different aspects of customer experience. The Data Source level describes the potential location of various KPIs.

This framework has been designed for mobile telecommunications, but it can easily be customized for other business to consumer industries by replacing the industry specific aspects (Handset, Network Usage and Network Usability).

We are interested in the relative value of the CEM KPIs for general marketing use, as well as the business value of the software tools used to collect and analyze them. Specifically for prepaid churn prediction, the following six aspects of the CEM framework presented above can represent this added value, as they have not been used before for prepaid churn prediction: Brand, Network Usability, Customer Support, Direct Communication and Content. Network usage, Billing and Handset have been addressed by Alberts (2006), as well as Archaux et al. (2004). Peers (Social Networks) is the topic of many papers (Dasgupta et al., 2008; Richter et al., 2010; Wang et al., 2010). The Demographics aspect is not applicable to prepaid subscribers, because for most of them this data is not available.

For this research we want to particularly focus on the added value of the CEM tool data source. This tool measures service quality (dropped calls etc.) for operational purposes, but claims are made that this data is also of key importance for prepaid churn management. We want to validate whether these claims are warranted from a business case perspective, assuming that traditional usage and customer relationship data is already available.

2.4 Research Setup

In order to examine the influence of the CEM metrics, variations in population sample and outcome definitions, we constructed three separate experiments which are described in the subsections below.

A set of general rules applies to all three. The predictors were recorded at the end of March 2009 with one, three and six months of history. The random sample of

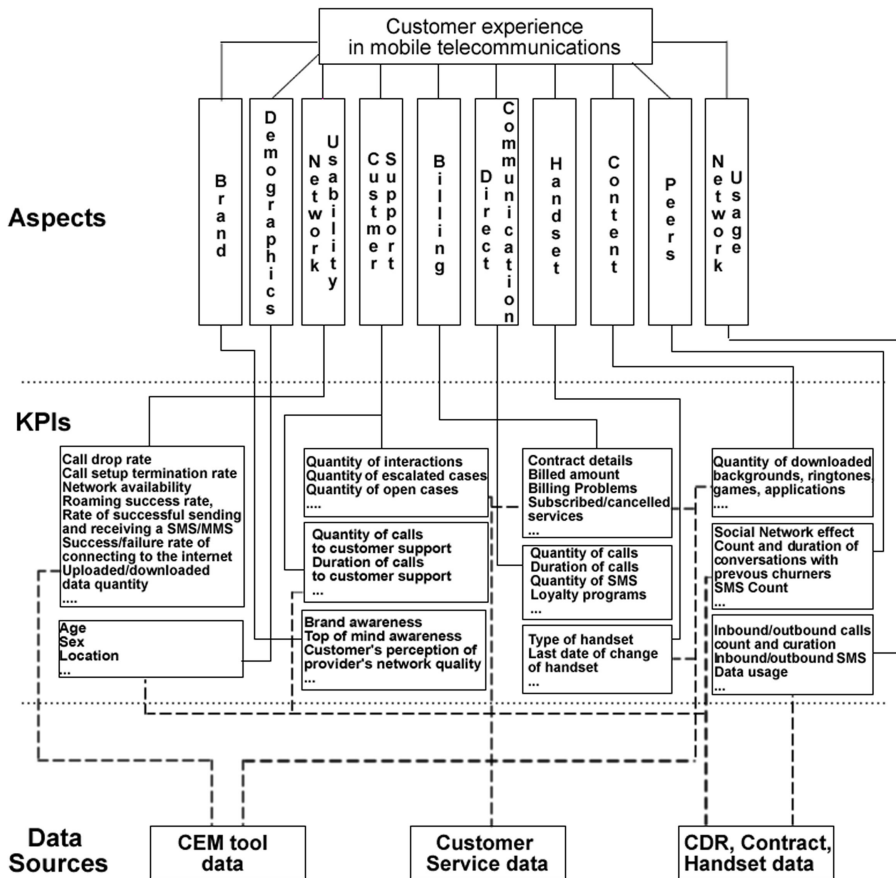


Figure 2.1: Customer Experience Framework for Mobile Telecommunications

around 10% of the prepaid customer base was divided into training, validation and test sets, using the 50:25:25 ratio. It was not necessary to oversample on churn, due to the choice of the software tool, which is able to handle unbalanced distributions of the outcome. The validation set was not used for developing the models in any automated way (e.g. model parameter setting or pruning using automated cross validation). It was used as an additional test set and for manual verification of the univariate performance of variables on an independent data set to avoid over-fitting. In other words, from an algorithm perspective it could have been called a test set, but from a methodological perspective we want to take a cautious approach and call it a validation set, because in theory the analyst himself could over-fit the data. The analyst does not have access to test set results in data preparation.

In our situation hold out validation rather than n-fold cross validation was both sufficient and more practical. If test data sets available are large enough, a simple holdout provides a true indication of future performance (see also Witten and Frank, 2005). Please note that we had ten times more instances available in the source data set. Hold out validation is also more practical given that it results in a model along with validation and test performance estimates. Cross validation at best results in a choice of an optimal modeling approach rather than a model. Hence, one would still need to build a single model, and likely require hold out validation to estimate the test set performance for the resulting model, because it might differ from the cross validation performance.

The next subsection will provide more technical details on the end to end data mining process followed.

2.4.1 End to End Data Mining Process

The commercial data mining tool Predictive Analytics Director (Pegasystems, 2008) was used for automated data preparation (attribute discretization, grouping and selection), modeling (logistic regression, decision trees) and model evaluation. All the steps in the data mining process, including the data preparation step, are actually decoupled from each other, with the goal of making the process more manageable and providing a factory approach to model building. On average, we do not expect this to have a negative impact on model performance or robustness. For instance, supervised discretization of continuous variables prior to modeling can sometimes improve the accuracy of the model (Dougherty, Kohavi and Sahami, 1995). As another example, wrapper based approaches for predictor selection (integrated predictor selection and modeling) are not inherently better than filter based approaches (predictor selection separated from modeling) (Tsamardinos and Aliferis, 2003).

The objective of data analysis is to transform the various attributes (a.k.a. variables, columns in the flat table), and identify the most informative ones among them at this stage, from a univariate perspective. A variable is considered informative if it has a certain level of influence over the outcome variable (which in this case is churn). Both statistical and graphical tools are used to establish this degree of influence. In

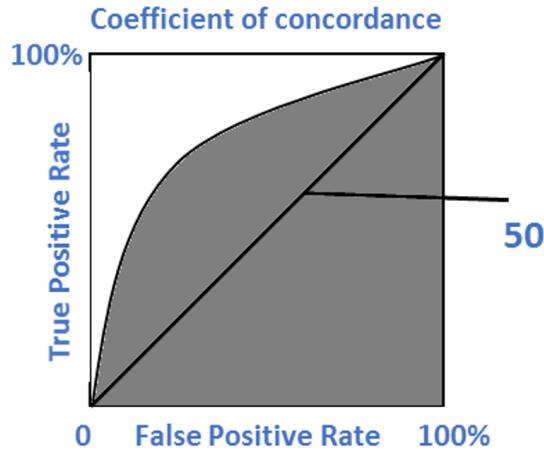


Figure 2.2: Coefficient of Concordance

this research, the chosen statistical criterion for evaluating the predictive performance of both variables and models with relation to churn is the Coefficient of Concordance (henceforth CoC), which is a rank order correlation measure, equivalent to Area under the ROC (AUC) and related to Kendall's tau (Kendall, 1938; Pegasystems, 2008). The major benefit of rank order correlation measures compared to basic measures such as accuracy is that these measures are not sensitive to skewed distributions of the outcome. For instance, assume a binary target for which one of the outcomes is very rare, e.g. churn. A majority vote model, which in practice is useless for selecting prospective churners, would have accuracy equal to the percentage of the majority class, i.e. $1 - \text{churn_rate}$ (e.g. if $\text{churn_rate}=1\%$, $\text{accuracy}=99\%$). However, in such a case, the values of CoC or AUC would only be 50 or 50%, respectively (equal to the values of random choice), which is the lowest value possible for real world models, as the prediction does not provide useful information to rank instances. One interpretation of the CoC measure is that in a scoring model it gives the probability that a randomly chosen positive case will get a higher score than a randomly chosen negative case. The CoC measures the gray area in the graph depicted on Figure 2.2 and can thus be translated to the Gini coefficient.

The data analysis process begins by discretizing continuous variables into a large number of bins. Numeric bins and symbolic values without statistically significant difference in churn rate are then grouped together. This balances the training set performance (many groups with varying churn levels) with out-of-sample (validation and test set) and out-of-time robustness (as many instances per group to provide robust estimates of churn for a group). Basically, this is a supervised, bottom-up approach to discretization of continuous variables and merging of numeric bins and

symbolic values into groups that display a significant difference in the outcome. The analyst can inspect the resulting histograms, and optionally change the parameters of the process, for instance the significance thresholds for merging bins and symbols.

After the initial data analysis, for the purpose of predictor selection, an automated procedure is used to group the variables that are correlated, independent of the target. A given predictor may have a high univariate performance, but also be correlated with other candidate predictors that are even stronger, hence not adding value to a model (subset predictor selection rather than univariate predictor selection). The user has three options when selecting predictors to be used for modeling: all predictors, only the best of each group, or manual selection. In our case we used the best predictors of each group as a starting point, and then experimented with minimal further manual selection, for instance by removing predictors from the bottom scoring groups altogether, or changing the parameters of the automated grouping process to force more or less groups. The main rationale for this was to force the use of CEM variables in order to enable a comparison.

The full data set consisted of more than 700 attributes, and the optimal number of predictors for final models typically ranges between 5 and 10 predictors. These volumes of attributes in the raw data versus the final model are quite typical for real world data mining, at least for marketing purposes; the full modeling table is reused as a basis for multiple modeling purposes, as it aims to capture all aspects relevant to customer experience (see also Figure 2.1). See the results section (Section 2.5) for more information of the types of predictors selected.

For logistic regression models, the raw predictor values are replaced with a normalized rank score concordant with the churn rate for the group (e.g. the predictor "age" is divided in various age ranges; if the group of 22-26 years shows higher churn rate on the training set, a higher score will be assigned to it). Logistic regression models are then built on the recoded data, to allow capturing of complex non-linear behavior, whilst using a stable low-variance modeler.

In our implementation, decision trees were forced to split between groups only, i.e. the grouped bins of discretized variables, thus not on the raw data (supervised discretization prior to induction (Dougherty et al., 1995)). Our motivation for this was to provide a level playing field for both algorithms (logistic regression and decision trees) and to ensure that only the "analyst approved" discretization was used. We used the CHAID splitting criterion, which selects the split points based on statistical significance as measured by the Chi Squared statistic (Kass, 1980). This criterion allows merging of both adjacent and non-adjacent bins of a discretized variable when making the splits.

For example, if a split on variable "age" is performed, and the bins 22-26 and 30-34 display similar level of significance, only one split for "age" will be created, containing both these intervals (i.e. age in 22-26, 30-34 and age not in 22-26, 30-34). The end result of our procedure is a binary tree. We are aware that CHAID and other decision tree induction methods are capable of producing n-ary trees (Perner, 2002),

which could have a somewhat better performance. However, it was not our goal to test algorithm performance, thus we used a standard binary decision tree.

The modeling process results in scoring models: a rank score concordant with the probability of being a churner is allocated to each of the instances. The CoC (AUC) measure is used to measure model quality (Kendall, 1938). In addition, we use gain charts as visual representation of model performance. On the y-axis, these charts show the captured proportion of the desired class (i.e. churners in selection divided by total number of churners) with increasing selection sizes (x-axis, from highest scoring to lowest scoring) (see Figure 2.4).

2.4.2 Definition of Prepaid Churn and Initial Sample

Defining prepaid churn is more complex than defining postpaid churn, as there is no concept of contract termination. Prepaid customers are deemed as churned if they are no longer "active" for a certain period of time.

The working definition of "activity" within the company where this research was conducted was used as operational definition. This definition is based on various aspects of usage of telecom services. On one hand, it is very detailed and captures the various aspects of activity within mobile telecom industry. On the other hand, it simplified the data acquisition.

Definition 1. Operational definition of activity: Outgoing (initiated) calls, top-up (recharge) of the account, sent SMS/MMS messages and received (incoming) and answered calls are considered as an activity. Received calls without picking up, received SMS/MMS messages and bonus credit top ups awarded by the company are NOT considered as an activity.

Involuntary prepaid churn in this company occurs after six consecutive months of no activity. After this, the customer can no longer use the services of the company via that particular SIM card, and the phone number may be reissued to another client after a certain period. However, six months of no activity is too long of a period to investigate voluntary prepaid churn. Most of the internal projects investigating prepaid churn within this company consider customers to have churned if they had not been active between two and three consecutive months. Therefore, we propose the following

Definition 2. Operational Prepaid Churn Definition: Customers are marked to have churned if they are registered to have two consecutive months of no activity, or more.

It was also necessary to set certain boundaries for the population taken into account for the modeling process. The population boundaries are set in Definition 3.

Definition 3. Population definition: The population consists of prepaid subscribers that meet the following two conditions:

1. They have at least one registered activity in the last 15 days.
2. Their first activity date is at least four months before.

The first condition enables avoiding subscribers who can already be classified as churners using the definition above. Arguably, to avoid this group, it would be sufficient to set the limit in condition 1 to 59 days, but this would make the prediction task trivial. The boundary was set to 15 days as a balance between excessive reduction of the sample and limitation of the information spillover (the subscribers inactive for 30 days have already begun to display churn behavior). The second condition is useful for avoiding frequent churners (it implies loyalty) or tourists, and for avoiding bias when measuring communication with previous churners.

In summary, for the particular data set constructed, all predictors were measured in the first trimester of 2009; churn was measured two months later; inactivity at recording was limited to maximum 15 days and first activity date had to be minimum four months prior to recording. This served as baseline data set.

2.4.3 Experiment A: Addition of CEM Parameters

Capturing the service quality oriented CEM metrics was performed by using a CEM tool deployed in the company. This tool suffered from two limitations. First, it contained only 40 days of user history. In order to give the CEM parameters a fair competing chance, we only used traditional parameters with one month history. Second, this tool did not cover the entire customer experience as depicted on the CEM Framework on Figure 2.1. It is more of a network experience measurement tool due to the fact that it is focused on the more hygienic aspects of customer experience (aspects noticed by customers only if they are absent or have low quality). This database contains only the Network Usability aspect (e.g. Call Success Rate, SMS Success Rate, Call Setup Duration, etc.) and the Content aspect (e.g. count of accessing company's website for downloading ring-tones, backgrounds, etc.) of our CEM framework.

However, this tool did not capture any data related to Customer Support; therefore, Customer Support data was added from the CDRs. Other aspects of the CEM framework, such as Network Usage, Handset and Peers, which were recorded from the company's CDRs and customer subscription data, were used for benchmarking purposes (to test the added value of parameters that are viewed exclusively as CEM KPIs). Therefore, they cannot be treated as potential contributions of CEM to the model's performance, but rather as traditional parameters as described in Section 2.2. Demographic information was not available for prepaid subscribers.

The two remaining aspects, Brand and Direct Communication, could not have been analyzed because the company did not have data appropriate for data mining (Brand), or did not communicate proactively to its prepaid customers. Therefore, only three aspects were left for analysis: Network Usability, Content and Customer

Support. Hence, the only added value stemming from using CEM KPIs on prepaid churn modeling in this research can exclusively originate from one of these three aspects of the CEM framework. In order to determine the added value of CEM, we compare models consisting only of "traditional" parameters, to models consisting of both "traditional" and CEM parameters.

2.4.4 Experiment B: Variations in Population Sample

In order to test impact of the variations in the population sample we changed the activity restriction into maximum 30 and 0 days of inactivity; therefore, we change condition 1 from definition 3 into condition 1a and condition 1b, respectively.

Condition 1a: Subscribers must have at least one activity in the last 30 days.

Condition 1b: Subscribers must have at least one activity in the last day before recording.

The reasons for varying the sample on maximum period of allowed inactivity are threefold. First, our intention was to inspect the impact of these changes on model performance. Second, we wanted to test the contribution of the CEM parameters under different circumstances. Third, the typical period of user inactivity, before they become unavailable for contact and retention, is disputable. Additional motivation for experiment B can also be found in the results of experiment A.

Furthermore, using zero days of inactivity can be seen as a way to avoid information spillover. We defined churn as uninterrupted inactivity for two months. It is clear that the best predictor of this is if a user has already been inactive for a certain period. In other words, users have already started to display churn behavior. This spillover cannot happen when the inactivity period at recording is limited to zero. Additionally, these subscribers are likely to still be available for contacting. We can expect that churn is harder to predict for this subgroup of subscribers.

2.4.5 Experiment C: Change in Outcome Definition

Since there is no general consensus on a practical definition of prepaid churn in mobile telecommunications, it is reasonable to experiment with different definitions. In practical terms, this is the first question that arises during a prepaid churn modeling project. A change in the churn definition not only affects the churn rate, but also the future retainability of prepaid consumers deemed as churners using that definition. Our motivation for changing the outcome definition was also to investigate the impact of this change on the performance of the models, while keeping the same sample and data set, and test the added value of the CEM parameters in this situation. For these reasons, we introduce a so-called "grace" period of 15 days, thus changing Definition 2 into

Table 2.1: Sample size, churn rate and CoCs in experiments A, B1a, B1b and C

Experiment	Inactivity Allowed	Sample size	Churn Rate (%)	Max CoC Train Set	Max CoC Validation Set	Max CoC Test Set
A	15 days	62565	3.88	87.9	85.6	87.5
B1a	30 days	67986	6.09	89.2	89.0	88.7
B1b	0 days	32104	0.7	85.3	77.1	79.5
C (Definition change)	15 days	62565	2.4	72.7	68.6	68.5

Definition 2b: Customers are marked to have churned if they are registered to have at least one activity in the first 15 days after recording, followed by 2 consecutive months of no activity, or more.

This outcome definition also resolves the information spillover issue discussed in Subsection 2.4.2. Namely, subscribers must have an activity in the first 15 days (“grace period”) after recording, thus breaking out of their already displayed churn behavior. Subscribers inactive in this grace period, who continue to be inactive in the future, are labeled as non-churners, because they have already churned at the time of recording, and are of no interest. In addition, subscribers recognized as future churners using this definition are more likely to be available for contacting and retention, because we are aiming at subscribers who are first active for 15 days, and then inactive for two months or more.

2.5 Results

In this section we report the results of the experiments of adding CEM data, changing the population based on inactivity duration and changing the outcome definition, or experiments A, B and C, respectively (see Table 2.1 for highlights). As explained in Section 2.4, we are using CoC as the main quantitative measure to compare models. In addition, we present gain charts to provide a visual reference for the performance of the models on the training set. These charts are used for illustrative purposes only, namely to visualize CoC performance through percentage of detected churners in a given percentage of the population scores (e.g. 78% of churners within top 20% of population scores).

It is worth mentioning that performance-wise, two models can be compared only if they have been created under the same experimental setup (e.g. Models A.Excl.CEM and A.Incl.CEM can be compared, while Models A.Excl.CEM and C.Excl.CEM cannot).

During experiment A, it became immediately clear that certain aspects of the CEM framework will not be able to add value in the case of prepaid churn prediction.

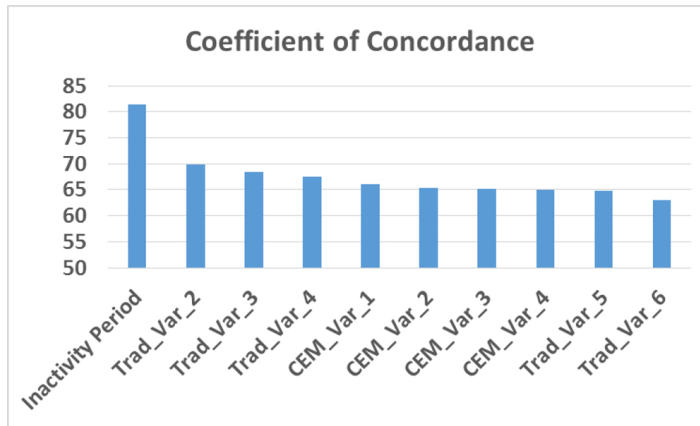


Figure 2.3: Coefficient of Concordance of predictors grouped in group 1 for experiment A

Namely, a very small percentage of the prepaid users had contact with Customer Service, or had downloaded any content from the provider's website. Therefore, none of the variables from these two aspects of the CEM Framework presented on Figure 2.1 could have been a good predictor. Furthermore, a very small percentage of these users used data services. Hence, Network Usability parameters related to mobile internet usage were also not good predictors. The only CEM predictors left to add value were the voice call and SMS related KPIs from the Network Usability of the CEM Framework. Model A.Excl.CEM contained only traditional parameters and was built on six predictors, containing only one month history. Adding parameters with longer history did not change the performance of the model substantially. The strongest predictor for churn was the inactivity period, followed by the remaining credit on the user's prepaid account. Other variables included were the handset and count and duration of calls. Model A.Excl.CEM had the following CoCs: 87.8 on the training set, 85.6 on the validation set, and 87.5 on the test set. Model A.Incl.CEM contains the same six variables, plus two CEM parameters (from the Network Usability Aspect of the CEM Framework), which were selected using a trial and error process based on their respective CoCs and the predictor group to which they belonged. Model A.Incl.CEM had the following CoCs: 87.9 on the training set, 85.6 on the validation set, and 87.5 on the test set. Thus, the results were reasonably consistent throughout the training, validation and testing sets. The only difference in the performances of these models was on the training set, which is not the best measure of model performance. Both these models were built using decision trees, but models built on the same variables using logistic regression performed just as well.

For visual reference only, we present the gain chart of the two models created for experiment A (on the training set) in Figure 2.4c. The difference in CoC of 0.1

Table 2.2: Grouping of variables of Model A.Incl.CEM

Predictors	CEM Framework Aspect	Group	Performance (CoC)
Inactivity Period	Network Usage	Group 1	81.1
Trad.var.x	Network Usage	Group 2	70.8
Trad.var.y	Network Usage	Group 2	70.6
Trad.var.z	Network Usage	Group 2	69.8
CEM.var.x	Network Usability	Group 2	68.9
Remaining credit	Billing	Group 3	67.4
CEM.var.y	Network Usability	Group 4	63.2
Handset type	Handset	Group 5	57.2

(A.Excl.CEM-87.8 vs. A.Incl.CEM-87.9) is not even visible on the gain chart. Their gain charts are identical up to 50% of the population. Both models were able to identify 78% of the churners within the top 20% of the population scores (a lift of 3.9 in the top 20% of population scores). This is representative of the test set performance as well, because the CoC on the test set of both models (87.5) is very similar.

The inability of CEM variables to substantially improve the base model is due to two reasons. First, a large number of these variables have low CoC (univariate performance). Second, even when the CoC is relatively high, in most groups there are traditional non-CEM predictors with CoC values higher than the CEM variables in the same group. Such is the case in the highest ranked group 1, depicted on Figure 2.3.

In order to illustrate the reasons for the weak effect of the CEM variables on model performance improvement, we isolated only the eight variables from model A.Incl.CEM, and ran the automatic grouping operation, with more strict grouping parameter settings than used previously, to force further grouping. The results of this exercise are presented on Table 2.2. One of the CEM parameters, CEM.var.x, has a reasonably high performance, but is in the same group as three other traditional variables, which means it has a degree of correlation with them, and has the lowest CoC in that group. This explains why CEM.var.x does not improve the model performance substantially. The second CEM variable in this model, CEM.var.y is in a class of its own, but it does not have a very high CoC; therefore, it cannot add substantial value to model performance. Please note that Table 2.2 and Figure 2.3 do not present the same groups of predictors.

Due to the strong influence of the inactivity period in experiment A, we decided to vary the population sample by changing the inactivity limit at recording to 30 and zero days (experiments B1a and B1b, respectively). Model B1a.Excl.CEM was built on seven non-CEM variables, very similar to the ones used in model A.Excl.CEM, using decision trees. B1a.Incl.CEM was built on the same seven variables plus two

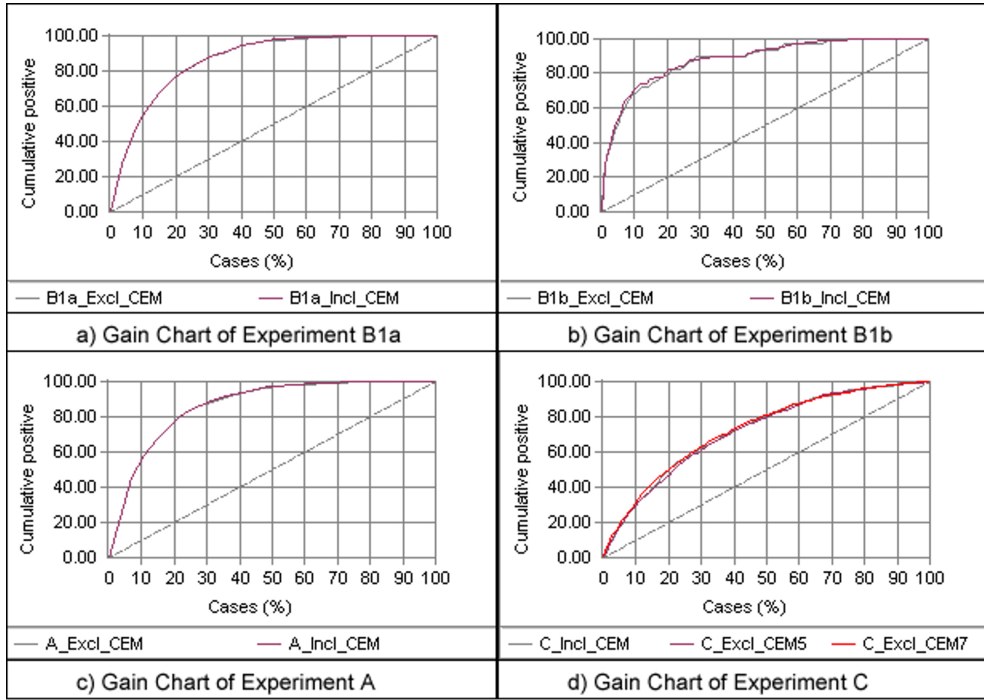


Figure 2.4: Gain chart of models for experiment A, B and C (training set)

CEM parameters (again from the Network Usability Aspect of the CEM Framework), selected based on the improvement they caused. Both models had the following CoCs of 89.2 on the training set, 89.0 on the validation set and 88.7 on the test set, which is a very stable performance. Once again, there was no visible performance difference (see Figure 2.4a). The influence of the inactivity period was even stronger, once again followed by the remaining credit. The gain chart on Figure 2.4a shows that both models identified about 78% of churners in the top 20% of the population scores.

In experiment B1b, the model B1b_Excl_CEM is built on only two non-CEM parameters (prepaid account balance and count of calls), because adding more variables caused overfitting (i.e. higher differences in CoC on the training, validation and test set). Model B1b_Incl_CEM was built on the same two variables, plus one more CEM parameter (call setup duration). This time, both models were built using logistic regression. In this situation, models built on decision trees were less stable across the three data sets, and their performance on the test set was lower, even though the performance on the training set was similar (overfitting). Nevertheless, even when using logistic regression, the models were not as stable as they had been in the previous cases.

Model B1b_Excl_CEM had the following CoCs: 84.9 on the training set, 77.1 on the validation set and 79.5 on the test set. Model B1b_Incl_CEM had the following CoCs: 85.3 on the training set, 74.0 on the validation set and 79.5 on the test set. Conclusively, adding CEM variables did not improve performance. On the contrary there is a performance drop on the validation set. Please note that the gain chart on Figure 2.4b is somewhat optimistic, because it was built on the training set. The performance on the test set is lower. Nevertheless, both models are able to identify 71% of the churners on the top 20% of the population scores in the training set, which is lower than the models in the previous experiments.

Finally, we present results for experiment C. The altered churn definition included the requirement of activity in the first two weeks of the outcome period, followed by a period of no activity of two months or more. Model C_Excl_CEM5 contained only traditional parameters and was built on five predictors. This model had the following CoCs: 72.2 on the training set, 68.5 on the validation set and 67.9 on the test set. Model C_Incl_CEM contained the same five variables, plus two CEM parameters. The CoCs of that model were 72.7 on the training set, 68.6 on the validation set and 67.9 on the test set. Both models were built using logistic regression, because these had a better performance on the test set when compared to models built on decision trees. There is an insubstantial 0.1 CoC improvement on the validation set, which is also achievable by adding two non-CEM parameters (model C_Excl_CEM7).

The gain charts of these models' performances on the training set are presented on Figure 2.4d. The maximum churners percentage achieved within the top 20% of population sample here is 50%. Note that this number would be even lower on the test set. Nevertheless, this is almost 30% identified churners less (in the top 20% of population scores) than what was achieved in experiments A and B1a.

The results of experiment C were less consistent throughout the three data sets when compared to experiments A and B1a, but somewhat more consistent when compared to B1b. It is interesting that the inactivity period was not a factor in these models, even though 15 days of inactivity were allowed at recording. The most powerful predictors were the remaining credit, the handset and the call count. Similarly to experiment A, the inability of CEM variables to improve the model performance in both experiments B and C is due to either low CoC or correlation with stronger traditional predictors.

2.6 Discussion and Future Research

We conducted three experiments in order to compare the influence of new CEM parameters, as well as changes in sample population and outcome definition, on prepaid churn model's performance.

CEM is advertised in literature to have an added value in predicting churn, but this was not the case in our experiments. Models without CEM parameters performed almost the same as models which included CEM parameters in all three

experiments. However, we only tested the CEM parameters with “hygienic” nature, which are only noticed when absent or unsatisfactory. The softer aspects of CEM remain untested. However, at this point we expect that it will be hard to find new non behavioral predictors with sufficient predictive power compared to the behavioral data. Even though the CEM data we had available contained only 40 days of history, it is unlikely that longer history would change the outcome, because the non-CEM parameters used by the models in most cases also had only one month history.

The rationale behind the low added value of CEM parameters for prepaid churn modeling may be found in several factors.

The prepaid customers themselves are the first factor. Prepaid customers that were subject to this research had average call duration of around one minute. A very low percentage of prepaid users used data services or have called the Customer Service in the research period. In other words, these events are rare, which limits the potential to become an interesting predictor. Next, prepaid users are mostly interested in controlling (reducing) their mobile phone expenses; otherwise, they could switch to using post-paid services that offer less expensive calling tariffs. The interest of prepaid users to control their mobile phone expenses may have been enhanced by the Global Financial Crisis of 2007/2008. To summarize, prepaid customers are more concerned with the quantity of experiences, which is measured by traditional predictors (Section 2.2) rather than the quality of their experiences, measured by the CEM parameters. The only exception is the handset which is a parameter that deeply influences the quality of the experiences, and is also regarded as a lifestyle product.

The second factor that could explain the low added value of CEM parameters was the high network quality. The percentage of customers experiencing network problems (negative experiences) is very small. However, the network quality cannot be seriously tested on average call duration of one minute.

The third explanation for the low added value of CEM parameters is that the quality parameters are correlated (to a degree) to their quantitative counterparts (e.g. number of dropped calls is correlated to a degree with number of calls).

Changing the population sample by varying the inactivity limit at the time of recording between 15 and 30 days also did not contribute to a substantial change in model performance. Models in experiments A and B1a had CoCs of about 88 and 89, respectfully (they identified between 78% and 79% of churners in the top 20% of the population). However, the churn rate did change drastically (there are almost twice more churners in B1a), while the sample size change was less than 10%, as presented in Table 2.1. These changes are even more evident at experiment B1b, where no inactivity was allowed at recording. Here, the sample size is twice smaller when compared to experiment A, while the churn rate is five times smaller when compared to experiment A, and even 9 times smaller when compared to experiment B1a. Due to the very low churn rate and the higher complexity of the task, the performance in this experiment was lower. The maximum CoC achieved on the test set was 79.5, which is nine CoC points less when compared to the other two

experiments; this results in a lower percentage of churners in the top 20% of the population (8% less when comparing training sets, but the difference is higher on the test set). The change of allowed inactivity period to zero also influenced the model stability. Note that we used a very small number of variables (two and three) in this experiment's models, in order to avoid overfitting. Once again it is important to emphasize that customers with zero days of inactivity at the time of deployment are more likely to be available for retention than customers with 15 or 30 days of inactivity.

The most dramatic change in model performance was shown when the outcome definition was changed and a so-called grace period of 15 days was included to mirror the inactivity period allowed at the time of recording. The model performance dropped by 20 CoC points on the test set, compared to experiments A and B1a (results in an almost 30% drop in identified churners in the top 20% percent of the sample on the training set, and even more on the test set). This does not mean that models built under experimental setup C are worse than the others. It merely implies the expected performance under such conditions. This steep decline is due to the complicated churn definition we deployed (we targeted an inactivity-activity-inactivity pattern). In this case, the inactivity period, that was a dominating variable in experiments A and B1a, was not a factor at all. The benefit of using such a definition is that upon deployment, the identified churners are likely to have an activity in the next 15 days, which makes them available for retention.

The focus of the research was on the impact of the experimental setup, rather than the algorithm used to create the model. Therefore, we used standard data mining algorithms, such as decision trees and logistic regression. Having said that, we would like to emphasize that in experiments A and B1a, there was barely any difference on model performance that could be attributed to the usage of the particular algorithms. In experiments B1b and C, there was a small difference in performance of algorithms, but it was not as substantial as changing the outcome definition or changing the allowed period of inactivity at recording to zero. In these two experiments, logistic regression had a more stable performance on the training, validation and test sets, compared to decision trees.

In terms of directions for future research, we would like to investigate a richer set of customer experience data, particularly around proactive communications and brand. Additionally, it would be worthwhile to investigate the relations between duration of inactivity and availability of subscribers for retention, by inspecting their presence on the network (regardless of making calls). Last but not least, we would like to take into account the feedback from retention campaigns, in order to focus on "retainable churners." After all, the end target of churn prediction is retention.

2.7 Conclusion

In this chapter, we presented how performance of prepaid churn models changes when varying the conditions in three different dimensions: data- by adding CEM parameters; population sample- by limiting the inactivity period at the time of recording to 15, 30 and zero days, respectively; and outcome definition- by introducing a so-called grace period of 15 days after the time of recording, in which customers must make an activity in order to be classified as churners.

Looking at the research question posed in section 2.1, the answer is clearly that changing the outcome definition has a much higher influence of the performance of the prepaid churn models than adding the CEM parameters or changing the characteristics of the sample based on inactivity at the time of recording. Adding the CEM parameters into the models did not add substantial value in terms of model performance under any of these conditions. Similarly, switching the population sample on the period of inactivity at the time of recording between 15 and 30 days did not influence model performance, only the sample size and churn rate. When we changed the population sample by disallowing inactivity at the time of recording, apart from the change in sample size and churn rate there was also a drop in performance and stability of the models. However, this drop in performance was not nearly as high as the one that occurred when changing the outcome definition by setting a grace period, thus making the behavior to be predicted more complex. This change obviously influenced the churn rate as well. Nevertheless, the latter two approaches should provide more time for retention.