

Computational Intelligence Methods Applied to the Fraud Detection of Electric Energy Consumers

B. de Araujo, H. de Almeida, and F. de Mello

Abstract— Brazilian Energy Distributors need to pay special attention to electric energy loss due to non-technical irregularities. Fraud and robbery represents a loss of revenue of billions of dollars every year. The energy utilities first approach to this problem was the creation of campaigns for prevention, raising awareness of the population and encouraging the denunciation of fraudsters. Unfortunately, this was not enough, and proactive actions were needed. The current effective strategy is conducting technical inspections at consumers to whom there is some kind of evidence, based on customer data, consumption patterns, reader's notes, etc. This paper the application of neural network models, *nnet* and *avNNet*, on selecting a set of consuming units for which there is a high probability of frauds and irregularities. It also includes data enrichment to enhance the results of such neural networks. First, it presents a brief introduction about the problem, scope of this article and related work. The data, to which the companies usually have access, is analyzed and the classification model construction is described. Then, the resulting models are compared considering the precision, accuracy and duration of training. Finally, there is a reasoning concerning such results and remarks about how energy distributors may use this information.

Index Terms—Non-technical energy loss, Neural networks, Classification, Machine learning.

I. INTRODUÇÃO

O tema deste trabalho é a aplicação de algoritmos de aprendizado de máquina para a detecção de fraudes e irregularidades no uso de energia elétrica, por parte dos consumidores. Deste modo, o problema a ser resolvido é realizar a classificação de um conjunto de consumidores de energia elétrica segundo o critério de suspeição de fraude para melhorar o resultado estatístico das inspeções efetuadas por concessionárias de energia.

A quantidade de energia elétrica perdida por motivos de irregularidades, furtos e fraudes no Brasil vem crescendo nos últimos anos.

Este trabalho foi apoiado pelo Fundo Poli 19.257 da Fundação Coppeltec da Universidade Federal do Rio de Janeiro.

B. S. de Araujo estava na Universidade Federal do Rio de Janeiro e agora está na TechnipFMC, Rio de Janeiro, Brasil (e-mail: breno-7-tj@poli.ufrj.br).

H. L. S. de Almeida é membro do Laboratório de Inteligência de Máquina e Modelos de Computação da Universidade Federal do Rio de Janeiro, Rio de Janeiro, Brasil (e-mail: heraldo@ufrj.br).

F. L. de Mello é chefe do Laboratório de Inteligência de Máquina e Modelos de Computação da Universidade Federal do Rio de Janeiro, Rio de Janeiro, Brasil (e-mail: fmello@poli.ufrj.br).

Segundo a Associação Brasileira de Distribuidores de Energia Elétrica [1], o percentual de perdas comerciais (furtos, fraudes, problemas de medição e de faturamento) em relação à energia injetada no sistema global passou de 3,99% em 2000 para 5,91% em 2012.

Isso se reflete não só em uma perda de receita para as empresas distribuidoras, como também em um aumento no faturamento dos consumidores não-fraudadores. O combate a esse tipo de perdas tem se tornado um foco estratégico para a garantia de receita dessas empresas e para a criação de normas e regulamentos por parte de órgãos reguladores. Contudo, para identificar se um consumidor está de fato furtando ou fraudando energia ainda é preciso que um especialista seja enviado a campo, no local de consumo, para que seja feita uma inspeção técnica. Com o intuito de melhorar o resultado estatístico destas inspeções, as distribuidoras têm voltado sua atenção para o comportamento das unidades de consumo, utilizando técnicas de mineração de dados e análise de dados na tentativa de identificar previamente possíveis fraudadores, e assim melhor direcionar e aperfeiçoar inspeções em campo.

Desta forma, o presente trabalho apresenta um comparativo entre os algoritmos *nnet* e *avNNet* de inteligência computacional [15] aplicados à detecção de consumidores suspeitos de furto/fraude de energia. Além disso, há também uma tentativa de melhorar os resultados individuais com ajuda de enriquecimento dos dados a partir de notas de leituristas. É considerado para este estudo um conjunto de consumidores que já recebeu, em algum momento, uma visita de um inspetor para verificação de fraude. Esta visita e, principalmente, o resultado da inspeção (fraude ou não-fraude) são essenciais para o treinamento e validação dos modelos. Com base no histórico de consumo e outros dados característicos, os algoritmos de aprendizado de máquina buscam classificar o consumidor como suspeito ou não. Neste sentido, a principal contribuição é determinar uma composição eficiente de algoritmo de predição com vetores de características dominantes capaz de prover os melhores resultados na identificação de consumidores de energia com indícios de irregularidade/fraude. As redes neurais são empregadas neste trabalho pois estas permitem modelar problemas não-lineares, e porque conseguem chegar a convergência do modelo preditivo utilizando um volume menor de medições nas Unidades de Consumo. Assim, entende-se que é oferecida às concessionárias de energia uma abordagem eficiente para determinação das unidades de consumo potencialmente interessantes de serem averiguadas.

II. TRABALHOS RELACIONADOS

Penin [2] aponta que perdas técnicas podem ter várias causas e correspondem à parcela de energia perdida devido a fenômenos físicos dos materiais utilizados nos sistemas elétricos durante o processo de transmissão de energia. Já as perdas não técnicas ocorrem quando parte da energia distribuída não é faturada pela concessionária, e ainda ressalta que não é possível calcular diretamente essas perdas através de métodos matemáticos e científicos, porém uma estimativa geral pode ser feita, subtraindo as perdas técnicas das perdas globais de energia.

Um estudo [3] do *Northeast Group* aponta que o Brasil é o segundo país que mais perde energia anualmente (USD 10bi) devido a fraudes e furtos, atrás apenas da Índia (USD 16,5bi). Rong *et al.* [4] utilizam *wavelets* combinadas com múltiplo classificadores para identificar fraudes em redes de distribuição elétrica. O artigo confirma que existem padrões escondidos em séries temporais, o que sugere a possibilidade de aplicar classificadores para a identificação de fraudes. Nagi *et al.* [5,6] apresentam uma análise de perdas não técnicas empregando aprendizado de máquina através de *Support Vector Machines* (SVM). Neste caso, a busca por parâmetros ótimos a serem utilizados pelo SVM levou um tempo proibitivo para ser empregado à realidade das empresas de distribuição de energia brasileiras, o que nos influenciou a escolha por classificadores com tempo de treinamento mais rápidos, tais como os empregados no presente artigo. Nizar e Dong [7] comparam o emprego das SVMs com *Extreme Learning Machine* (ELM) no estudo de caso de perdas de energia elétrica em países desenvolvidos e em desenvolvimento. Novamente, para poderem lidar com o custo computacional de ajuste dos hiperparâmetros do SVM, os autores restringem o tamanho de seu *dataset*.

Filho *et al.* [8] criam uma abordagem para selecionar clientes sujeitos a inspeções a fim de identificar fraudes e falhas em equipamentos de medição de consumo elétrico. Eles utilizam árvores de decisão com sucesso, porém os dados utilizados em nosso trabalho parecem representar um problema não linear, o que sugere uma impossibilidade de uso de tais árvores de decisão e reforçam o emprego de redes neurais. Cabral *et al.* [9] comparam dois sistemas de detecção de fraudes para clientes de energia elétrica de alta e baixa voltagens utilizando mapas auto organizáveis, e indicam a necessidade de um ajuste mais elaborado nos casos em que fraudes nas Unidades de Consumo geram apenas pequenos furtos de energia. Spirić *et al.* [10] analisam séries temporais de consumidores de energia de baixa voltagem para identificar fraudes e consumos ilógicos por porte destes consumidores.

Tais autores apresentam uma discussão bastante madura sobre a natureza destes furtos. Jokar *et al.* [11] utilizam *smart grids* para identificar padrões de roubo de energia. Faria *et al.* [12] também utilizam *smart grid* e criam procedimentos forenses para detectar furto de energia baseado em dispositivos eletrônicos adulterados. No momento da publicação deste artigo, foi divulgado um artigo [16] que emprega aprendizado de máquina para detecção de anomalias e fraudes a partir de dados coletados de *smart meters*, com especial atenção ao classificador XGBoost. Neste caso, os *smart meters* fornecem um volume de dados significativamente maior de medições que

possibilitam traçar um perfil bastante ajustado dos clientes, um número de medições incompatíveis com a realidade brasileira.

III. ANÁLISE EXPLORATÓRIA DOS DADOS

A concessionária de energia elétrica considerada neste estudo tem acesso a uma diversidade de dados dos consumidores. Pode-se dividi-los nos seguintes subgrupos principais: dados do consumidor, dados da instalação elétrica e dados mensais de consumo de energia, incluindo notas de leitura.

Os dados do consumidor são informações que identificam, caracterizam e localizam geograficamente a Unidade Consumidora (UC). São informações que podem ser obtidas através do formulário de cadastramento, tais como: nome da pessoa física ou jurídica; localização geográfica; classe de consumo, podendo ser principalmente residencial, comercial, industrial, iluminação pública, poder público, entre outros; ramo de atividade.

Os dados da instalação elétrica são dados técnicos, importantes para o projeto e dimensionamento da rede ou para a tarifação da energia. Dentre os mais relevantes pode-se citar: fase da ligação; dados gerais sobre o medidor utilizado na UC, indicando se ele é eletrônico ou eletromecânico, se seu mostrador é de 4, 5 ou 6 dígitos; grupos tarifários [13]; subgrupos tarifários.

Os dados mensais de consumo são importantes para a tarifação feita pelas empresas e são de extrema importância na escolha do vetor de características para o modelo de aprendizado de máquina. Estes dados são formados pela data da leitura ou data de referência, pelo consumo medido e consumo faturado. Diferente do que acontece nos artigos publicados na literatura envolvendo regiões com *smart meters*, a massa de dados empregada neste trabalho contém apenas uma medição mensal da UC. Este período mais longo entre medições dificulta a identificação de alguns padrões de comportamento, algo que não acontece em trabalhos que utilizam taxas amostrais mais elevadas.

Dentre estes dados disponíveis, foi necessário selecionar aqueles com maior correlação com o resultado das inspeções. No estudo em questão foram analisadas apenas UCs do Grupo B, composto por de unidades consumidoras de baixa tensão. Normalmente as concessionárias traçam estratégias especiais para UCs do Grupo A, isto é, consumidores de alta energia, onde a receita recuperada tende a ser maior, e por isso recebem uma atenção especial da área de perdas das distribuidoras.

É importante citar também que, por se tratar de uma aplicação de redes neurais supervisionadas, serão analisadas apenas UCs que tiveram alguma inspeção, com um resultado válido. Isto pode introduzir certo viés, pois uma UC que recebe uma inspeção provavelmente já sofreu tal averiguação porque apresentava algum tipo de comportamento suspeito. Infelizmente esta é uma limitação prática que não pode ser contornada, no caso de algoritmos supervisionados.

Entende-se que as variáveis fase de ligação e classe da UC, caracterizam indiretamente o consumo. UCs trifásicas comerciais tendem a ter um consumo de energia mais elevado do que UCs bifásicas ou monofásicas residenciais. Além disso, indústrias deveriam ter um consumo mais elevado do que residências comuns. Se isso não acontece, então pode ser que

haja alguma irregularidade. É claro que pode haver outras explicações para esse consumo diferenciado, como por exemplo o estabelecimento estar fechado ou desocupado por um período de tempo. Contudo, mesmo nesta situação cabe o indicativo de investigação destas circunstâncias.

Uma estratégia muito utilizada para detecção de irregularidades é a observação da curva histórica de consumo faturado das UCs e a procura por quedas de consumo. Para poder detectar tais variações incomuns de consumo, foram selecionadas as seguintes variáveis estatísticas dos valores históricos de energia faturada:

- Mediana de consumo de 6 meses antes da inspeção até o último consumo faturado anterior à inspeção;
- Amplitude interquartil de consumo de 6 meses antes da inspeção até o último consumo faturado anterior à inspeção;
- Mediana de consumo de 18 meses antes da inspeção até 6 meses antes da inspeção;
- Amplitude interquartil de consumo de 18 meses antes da inspeção até 6 meses antes da inspeção;
- Mediana de consumo de 30 meses antes da inspeção até 18 meses antes da inspeção;
- Amplitude interquartil de consumo de 30 meses antes da inspeção até 18 meses antes da inspeção.

A mediana e a amplitude interquartil foram utilizadas, em detrimento da média e desvio padrão, por representarem melhor o conjunto de dados, especialmente quando se trata de conjuntos com poucas observações, como é o caso de valores com mais de 6 meses de consumo. A partir destas estatísticas pode-se observar, por exemplo, se a UC teve uma variação muito grande na mediana de consumo, entre um período e outro, o que poderia indicar uma queda. Outros tipos de observações podem ser feitas, mas isto fica a cargo do algoritmo de aprendizado, que fará isso implicitamente durante o processo de treinamento.

Além destas variáveis, acrescenta-se também o resultado da inspeção para efetuar a supervisão do algoritmo de treinamento do modelo de classificação. Não foi realizada distinção entre casos de irregularidade e fraude, pois ambas situações demandam averiguação. Portanto, os dois possíveis resultados de uma inspeção são: regular e irregular.

O valor de *target* é o resultado da inspeção que pode ser regular ou irregular. Existem 6.425 registros, sendo 4.307 regulares e 2.118 irregulares. O *dataset* empregado possui como atributos os 23 itens a seguir: nome da pessoa física ou jurídica, localização geográfica, logradouro da UC, bairro, localidade, município, região ou base administrativa à qual a UC pertence, classe de consumo, ramo de atividade, fase da ligação, tipo de medidor, grupo tarifário, subgrupo, data da leitura, data de referência, consumo medido, consumo faturado, mediana de consumo de 6 meses antes da inspeção até o último consumo faturado anterior à inspeção, interquartil de consumo de 6 meses antes da inspeção até o último consumo faturado anterior à inspeção, mediana de 18 meses antes da inspeção até 6 meses antes da inspeção, interquartil de consumo de 18 meses antes da inspeção até 6 meses antes da inspeção, mediana de consumo de 30 meses antes da inspeção até 18 meses antes da inspeção, interquartil de consumo de 30 meses

antes da inspeção até 18 meses antes da inspeção. Ao se tentar realizar uma redução da dimensão dos dados, através de PCA (*Principal Component Analysis*) por método de covariância, obteve-se um vetor de 19 dimensões. Como a redução não foi relevante, optou-se neste projeto trabalhar com as 23 dimensões originais, em detrimento das 19 do PCA.

IV. MODELOS DE CLASSIFICAÇÃO

Na fase de treinamento, os dados dos consumidores são fornecidos ao modelo e seus parâmetros são ajustados, fazendo a comparação direta da saída do modelo (suspeito/não-suspeito) com o resultado da inspeção (regular/irregular). Foi feito o uso de um subconjunto de validação para evitar um sobreajuste (*overfitting*). Assim, foi definido um critério de parada para o aprendizado, isto é, no momento em que a acurácia do modelo sobre o conjunto de treinamento aumenta, mas a acurácia sobre o conjunto de validação começa a diminuir, o modelo é considerado como ajustado e o treinamento é encerrado. Caso este modo de treinamento não fosse adotado, esta situação resultaria em um modelo com viés direcionado ao conjunto de treinamento e que não teria um bom desempenho em conjunto de consumidores qualquer.

Finalizado o treinamento do modelo, seu desempenho é avaliado quando submetido ao conjunto de teste. Este procedimento é repetido para todos os algoritmos de classificação utilizados e, enfim, é comparada a efetividade de todos eles a partir de indicadores como taxa de acerto.

Para o treinamento e a aplicação do modelo foi empregado um processador Intel® Core™ i7 4500U@1.8GHz, 64 bits, 8GB RAM, sistema operacional Ubuntu 16.04 LTS, linguagem de programação R, em conjunto com o pacote *caret* (*Classification And REgression Training*)[14]. Foram utilizados os modelos de redes neurais *nnet* e *avNNet* (*Model Averaged Neural Networks*) disponíveis no pacote *caret*. O primeiro é um modelo tradicional de redes neurais *feed-forward* com duas camadas escondidas. Já o segundo modelo utiliza múltiplos modelos de redes neurais e produz um resultado global que é calculado a partir da média das saídas de cada rede neural. Em um primeiro momento, foram empregadas redes com uma única camada, pois a literatura indicava o uso de mapas auto organizáveis de uma só camada [9] e árvores de decisão [8]. Entretanto, os resultados não foram satisfatórios com redes de uma só camada, o que sugeriu uma característica não linear da massa de dados que foi utilizada. Esta característica foi comprovada quando alterado os modelos de rede para duas camadas escondidas, modelos estes que obtiveram resultados significativamente melhores.

V. DESCRIÇÃO DO EXPERIMENTO

Neste experimento foi analisada uma região da concessionária que compreende cerca de 370.387 UCs, cuja base de dados contém 21.421 inspeções, com 7.062 resultados ditos irregulares e 14.359 regulares. Esta base de dados corresponde a um banco de inspeções realizadas em consumidores do grupo B (baixa tensão). A Fig. 1 ilustra que a região analisada é composta majoritariamente por UCs monofásicas e que o número de UCs bifásicas é praticamente desprezível. Além disso, também mostra a predominância de UCs residenciais, rurais e comerciais na região.

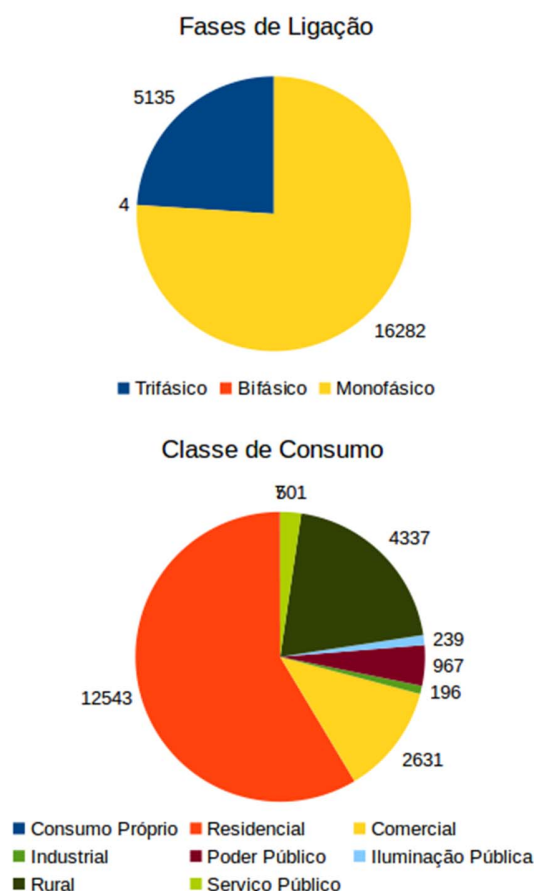


Fig. 1. Distribuição das UCs por fase de ligação e por classe de consumo.

As UCs encontram-se distribuídas em 46 municípios, identificados por seus códigos (ID_MUNICIPIO). Os dez municípios com maior número de ocorrências de inspeções na base de dados estão representados na Fig. 2, bem como os quinze municípios com maior número de irregularidades. Nota-se que todos os dez municípios com maior volume de inspeções constam da lista de quinze municípios com maior quantidade de irregularidades. Os três municípios com maior número de inspeções na base lideram a lista de municípios com maior quantidade de irregularidades.

O conjunto de dados foi dividido em um conjunto de treinamento e um conjunto de teste, contendo 70% e 30% dos dados, respectivamente. Isto corresponde a 14.996 observações no conjunto de treinamento e 6.425 observações no conjunto de teste. O treinamento dos modelos *nnet* e *avNNet* utilizou validação cruzada de 5-fold, com 10 repetições. Com esta validação 5-fold, o conjunto de treinamento foi dividido em 5 partições de aproximadamente 3.000 observações, que pode ser considerado um número satisfatório de observações para esta aplicação. O fato de serem apenas 5 partições fez com que o tempo de duração do treinamento não se estendesse na prática, visto que a escolha de um número maior de partições para o conjunto de treinamento implica em um tempo de duração mais longo para o processo de treinamento do modelo.

O tempo de duração (hh:mm:ss) para o treinamento do modelo *nnet* foi de 00:40:00,55. Já o modelo *avNNet* teve um tempo de treinamento notadamente mais alongado, de

02:34:28,47. Em seguida foram empregados os modelos de classificação então construídos para realizar a predição. A Tabela 1 apresenta a matriz de confusão resultante da aplicação dos modelos dos algoritmos de classificação. Na Tabela 1 as siglas TP, FN, FP e TN representam, respectivamente, as classificações verdadeiro-positivas, falso-negativas, falso-positivas e verdadeiro-negativas.

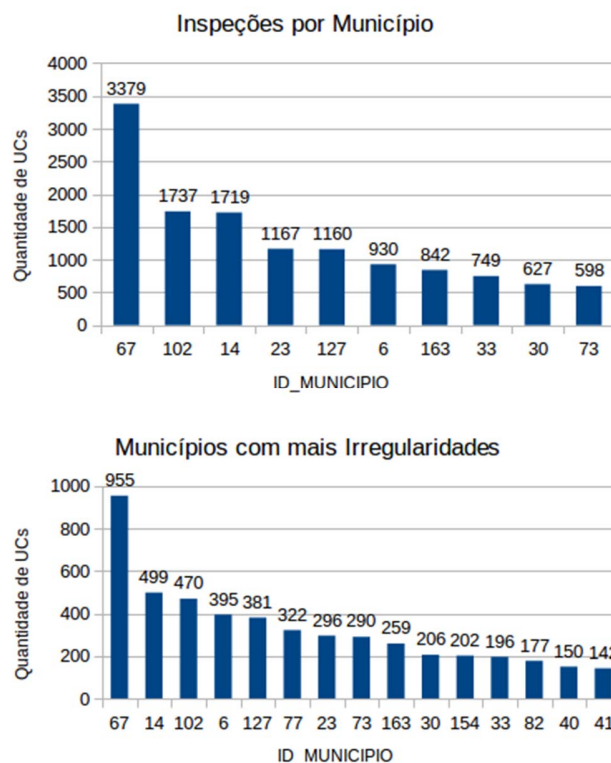


Fig. 2. Os dez municípios com maior número de inspeções com resultado e os quinze municípios com maior número de inspeções com resultado irregular, ambos identificados pelo código do município.

TABELA 1
MATRIZ DE CONFUSÃO RESULTANTE DA APLICAÇÃO DOS
MODELOS DE REDES NEURAIS (NNET E AVNNET)

Modelo <i>nnet</i>			
		Referência	
		Regular	Irregular
Predição	Regular	3624 (TP)	1202 (FN)
	Irregular	683 (FP)	916 (TN)
Modelo <i>avNNet</i>			
		Referência	
		Regular	Irregular
Predição	Regular	4096 (TP)	1752 (FN)
	Irregular	211 (FP)	366 (TN)

Ao analisar estes resultados, é importante entender que o custo de envio de um técnico para avaliação de uma UC é menor que o valor obtido pela recuperação de receita em função da energia recuperada. Assim, o valor relativo de indicadores de classificações irregulares preditas corretamente em relação ao total de predições irregulares $TN/(TN+FP)$, deve ser maximizado. Neste trabalho, tal relação será chamada de Relação Econômica, uma métrica de avaliação que não é

abordada em nenhum dos artigos publicados até o momento [5,6,7,8,9,10,11,12,16]. No entendimento dos autores desta pesquisa, isto acontece porque a literatura científica da área tem se concentrado em métricas de desempenho dos algoritmos classificadores, porém se esquece que a atividade fim, em última instância, é o aumento de lucro através de uma relação favorável entre incremento de receita e custo das inspeções. Sob esta ótica, tem-se aqui uma análise bastante alinhada com as atividades dos distribuidores de energia elétrica.

Lembre que as notas de leiturista podem ser um indicio de irregularidade/fraude, apesar de não expressarem categoricamente uma certeza de que o consumidor está irregular. O leiturista carece muitas vezes da habilidade técnica, além de não ser sua função realizar perícias na instalação. Contudo, este tipo de observação encontra-se disponível na base de dados da concessionária e, portanto, cabe um experimento de uma classificação com este tipo de parâmetro. Neste sentido, foram adicionadas novas características ao conjunto de variáveis descrito anteriormente:

- Quantidade de apontamentos de suspeita de fraude nos últimos 6 meses, anteriores à inspeção;
- Quantidade de apontamentos de medidor com defeito nos últimos 6 meses, anteriores à inspeção;
- Quantidade de apontamentos de medidor parado, com o local de consumo habitado, nos últimos 3 meses, anteriores à inspeção;
- Quantidade de apontamentos de medidor danificado, nos últimos 12 meses, anteriores à inspeção;
- Quantidade de apontamentos de medidor com display apagado, nos últimos 3 meses, anteriores à inspeção;
- Quantidade de apontamentos de medidor queimado, nos últimos 3 meses, anteriores à inspeção.

Deste modo, foi realizado um novo experimento com o modelo *avNNet*, desta vez com as informações sobre nota de leiturista (NL) adicionadas ao vetor de características. A mesma massa de dados foi empregada na avaliação, e as condições de treinamento foram mantidas. A matriz de confusão deste novo experimento é apresentada na Tabela 2.

TABELA II
MATRIZ DE CONFUSÃO RESULTANTE DA APLICAÇÃO DO MODELO AVNNet, COM A UTILIZAÇÃO DE NOTAS DE LEITURISTA NO VETOR DE CARACTERÍSTICAS.

Modelo <i>nnet</i> com nota de leiturista		Referência	
		Regular	Irregular
Predição	Regular	3601 (TP)	905 (FN)
	Irregular	387 (FP)	719 (TN)
Modelo <i>avNNet</i> com nota de leiturista		Referência	
		Regular	Irregular
Predição	Regular	3985 (TP)	939 (FN)
	Irregular	322 (FP)	366 (TN)

Como era de se esperar, esse experimento teve resultados melhores, devido à utilização de notas de leiturista. Estas comparações entre modelos são mais facilmente compreendidas quando se avaliam as taxas de erro total

$(FN+FP)/(TP+FN+FP+TN)$, erro na predição regular $FN/(TP+FN)$, erro na predição irregular $FP/(FP+TN)$, acurácia $(TP+TN)/(TP+FN+FP+TN)$, precisão $TP/(TP+FP)$, recall $TP/(TP+FN)$ e f1-score $2*precisão*recall/(precisão+recall)$. Estes indicadores encontram-se consolidados na Tabela 3. O que se observa é que o modelo *avNNet* combinado com as notas de leiturista apresenta uma taxa de erro total menor que os demais experimentos, assim como uma menor taxa de erro nas predições ditas regulares. Contudo, o grande diferencial está nas predições ditas irregulares, cuja taxa de erro reduziu mais sensivelmente e que é um melhor indicador para o contexto de averiguação de fraudes. O tempo gasto no treinamento dos classificadores não parece ser determinante na escolha do melhor modelo. Outro ponto importante é que a taxa de erro obtida no presente trabalho é melhor do que aquelas obtidas por trabalhos anteriormente descritos na literatura [4,8,9], ressalvada a questão da dificuldade de se conduzir uma comparação direta entre os métodos, haja vista que cada um deles emprega *datasets* distintos. O único modelo de realizar uma comparação direta com outras abordagens descritas na literatura seria utilizar um *dataset* padronizado para todos os experimentos, tornando invariante a contribuição dos dados na eficácia dos modelos. Entretanto, isto é algo que ainda não existe no momento.

TABELA III
INDICADORES DE AVALIAÇÃO DOS MODELOS DE CLASSIFICAÇÃO

Experimento	<i>nnet</i>	<i>avNNet</i>	<i>nnet</i> +NL	<i>avNNet</i> +NL
Taxa de Erro Total (%)	29,34	30,55	23,02	23,36
Taxa de Erro – Regular (%)	24,91	29,96	20,08	22,83
Taxa de Erro – Irregular (%)	42,71	36,57	34,99	25,53
Acurácia (%)	70,66	69,48	76,97	76,64
Precisão	0,8414	0,9510	0,9030	0,9252
Recall	0,7509	0,7004	0,7992	0,7717
F1-score	0,7936	0,8067	0,8479	0,8377
Relação Econômica (%)	57,29	63,43	65,01	74,46
Duração do Treinamento (hh:mm:ss)	00:40:00,55	02:34:28,47	00:49:42,51	03:05:57,00

Ainda analisando a Tabela III, observa-se que as acurácias dos dois algoritmos são muito parecidas, sendo estas incrementadas no momento em que há enriquecimento dos dados a partir da inserção de notas de leituristas. Contudo, estas mesmas acurácias apresentam valores nominais ainda baixos, algo que pode estar sendo influenciado pela assimetria do *dataset*. Os valores de precisão foram altos, sugerindo uma baixa taxa de falsos positivos, uma característica melhor do que aquelas relatadas em [5,6,7]. Com relação à métrica de recall, os valores estão acima de 0,5 o que indica bons resultados, porém não há uma clara diferenciação entre os modelos. Como o *dataset* empregado é assimétrico, a métrica de f1-score passa a ser mais confiável que a acurácia, e por esta razão pode-se dizer os treinamentos com notas de leituristas fornecem

modelos melhores que aqueles sem tais notas, o que corrobora com os resultados intuitivos de Faria *et al.* [12]. Além disso, segundo esta métrica, o algoritmo *avNNet* apresenta um melhor desempenho que o *nnet*. Por fim, cabe lembrar a necessidade de maximizar o Fator Econômico, algo que a literatura não tem observado em suas análises. Neste caso, o modelo *avNNet* também apresentou um resultado um pouco melhor que o *nnet*. Isto é um ponto importante porque não só o *f1-score* aponta o *avNNet* como melhor opção, mas também o Fator Econômico.

Na prática devem-se combinar essas estratégias de classificação de consumidores com uma avaliação da estimativa de energia a recuperar. A concessionária visa não só uma alta taxa de acerto nas inspeções, como também um alto retorno em termos de energia recuperada. Uma estratégia utilizada na prática, por exemplo, é focar as inspeções em clientes comerciais ou industriais, com alto padrão de consumo, mesmo que não apresentem muitos indícios de irregularidade. Pode ser que essa estratégia tenha uma taxa de acerto mais baixa, porém os poucos consumidores irregulares deflagrados serão obrigados a retornar uma quantia relativamente mais alta de energia recuperada à companhia. No entanto, não se deve negligenciar o retorno de receita oriundo de pequenas UCs quando analisadas em conjunto.

Para se ter comparativo, destaca-se que na região de análise, no período de janeiro de 2015 a dezembro de 2016, foram realizadas 8.367 inspeções direcionadas, com uma taxa de acerto acumulada de 30,91% (69,09% de erro). Inspeções direcionadas são inspeções a UCs selecionadas, que apresentem algum indício de irregularidade ou fraude, baseado somente na observação dos dados da UC e em notas de leiturista, assim como é feito no terceiro experimento deste projeto.

No entanto, é importante salientar que uma comparação dos resultados obtidos nos experimentos com essa taxa histórica não é totalmente justa, devido a alguns fatores. O desenvolvimento deste projeto foi baseado em modelos supervisionados de Aprendizado de Máquina, o que pressupõe necessariamente a utilização de dados para os quais se tem o resultado da classificação (valores de referência). Neste caso, isso significa dizer que somente foram utilizados dados de UCs que tiveram, em algum momento dentro do registro histórico da concessionária, uma ou mais inspeções técnicas para verificação de fraudes e irregularidades. Isto pode introduzir um viés no conjunto de dados, visto que essas UCs podem ter recebido inspeções devido a algum critério, que pode ter filtrado o conjunto de dados de forma não aleatória ou não representativa.

Por exemplo, algumas das UCs do conjunto de dados podem ter apresentado algum indício de irregularidade, o que justificaria a inspeção. Outras podem ter sido inspecionadas por serem UCs que consomem uma grande quantidade de energia (inspeções visando a energia recuperada ou a prevenção de perdas em clientes de alto consumo). Há ainda inspeções que são realizadas por ações especiais, como varreduras a consumidores sazonais (hotéis, pousadas, restaurantes em cidades de praia, etc) ou varreduras em locais com alta repercussão midiática (como uma estratégia para alertar a população quanto à presença da concessionária no combate às perdas).

V. CONCLUSÕES

Este trabalho propôs o uso de técnicas de aprendizado de máquina para auxiliar na identificação de consumidores de energia com indícios de irregularidade/fraude. Foram utilizados dados reais de uma concessionária de energia elétrica para o treinamento e teste dos modelos. Dentre os dados dos quais a concessionária dispõe, foi selecionado um subconjunto para compor os vetores de características que foram utilizados na aplicação dos modelos. Foram realizados experimentos com e sem a utilização de notas de irregularidade de leiturista no vetor de características e foi realizado um total de três experimentos. O experimento que apresentou melhores resultados foi o modelo de média das redes neurais com estas notas de leiturista, com taxa de erro quase três vezes menor que os procedimentos adotados atualmente.

Como projeto futuro, planeja-se implementar esses algoritmos em tempo real para a otimização de inspeções técnicas de irregularidade. Além disso, propõe-se o estudo e a aplicação de outros algoritmos, como *Support Vector Machines*, na busca de classificadores com melhores desempenhos, ou ainda, a avaliação dos modelos através da análise da curva de ROC e da curva de precisão-revocação. Sugere-se também a realização de estudos exploratórios para identificar outras características que possam ter correlação com a existência de irregularidades. Além disso, propõe-se estudar o enviesamento apresentado pelo conjunto de dados, procurando formas de reduzi-lo, e a determinação de uma medida quantitativa para o mesmo. Por fim, outra consideração que vale ressaltar é que, na prática, talvez faça mais sentido que as UCs sejam entregues à concessionária na forma de uma lista ordenada por prioridade, onde as UCs com maior potencial de irregularidade estejam no topo. Ao invés de se atribuir a cada UC uma classe (regular ou irregular), utilizaria-se os valores analógicos de saída do modelo de classificação para ordenação. Como a quantidade de UCs a serem selecionadas para receber inspeções técnicas pode variar dependendo do orçamento, nível de perdas comerciais totais, entre outros fatores, uma lista dessa forma deixaria a cargo da concessionária definir o valor de corte que determina quais UCs pertencem a cada classe prevista.

AGRADECIMENTOS

Esta pesquisa foi apoiada pelo Fundo Poli 19.257 da Fundação Coppetec da Universidade Federal do Rio de Janeiro (UFRJ). O trabalho foi desenvolvido no Laboratório de Inteligência de Máquina e Modelos de Computação (IM²C) da Escola Politécnica (Poli) da UFRJ.

REFERÊNCIAS

- [1] Resende, Tatiana. Perdas: Furto e Fraude de Energia, ABRADÉE, Especial Canal Energia, 2013. <<http://www.abradee.com.br/setor-de-distribuicao/perdas/furto-e-fraude-de-energia>>, acesso em 16 de janeiro de 2017.
- [2] Penin, C. A. d. S. Combate, Prevenção e Otimização das Perdas Comerciais de Energia Elétrica, São Paulo: Escola Politécnica da Universidade de São Paulo, 2008.
- [3] Northeast Group, Emerging Markets Smart Grid: Outlook 2015. December, 2014. <<http://www.northeast-group.com/reports/Brochure-Emerging%20Markets%20Smart%20Grid-Outlook%202015-Northeast%20Group.pdf>>, acesso em 28 de novembro de 2016.
- [4] Rong Jiang, H. Tagaris, A. Lachsz and M. Jeffrey, Wavelet based feature extraction and multiple classifiers for electricity fraud detection, IEEE/PES

- Transmission and Distribution Conference and Exhibition, pp. 2251-2256 v.3., 2002. doi: 10.1109/TDC.2002.1177814
- [5] Nagi, J., Yap, K. S.; Tiong, S. K.; Ahmed, S. K.; Mohamad, M. Nontechnical Loss Detection for Metered Customers in Power Utility Using Support Vector Machines, in IEEE Transactions on Power Delivery, v.25, n.2, pp.1162-1171, April 2010. doi: 10.1109/TPWRD.2009.2030890
- [6] Nagi, J., Yap, K. S.; Tiong, S. K.; Ahmed, S. K.; Nagi, F. Improving SVM-Based Nontechnical Loss Detection in Power Utility Using the Fuzzy Inference System, in IEEE Transactions on Power Delivery, v.26, n.2, pp.1284-1285, April 2011. doi: 10.1109/TPWRD.2010.2055670
- [7] Nizar, A. H.; Dong, Z. Y.; Wang, Y.. Power Utility Nontechnical Loss Analysis With Extreme Learning Machine Method, in IEEE Transactions on Power Systems, v.23, n.3, pp.946-955, Aug. 2008. doi: 10.1109/TPWRS.2008.926431
- [8] Filho, J. R.; Gontijo, E. M.; Delaiba, A. C.; Mazina, E.; Cabral, J. E.; Pinto, J. O. P. Fraud identification in electricity company customers using decision tree, 2004 IEEE International Conference on Systems, Man and Cybernetics (IEEE Cat. No.04CH37583), pp.3730-3734, v.4, 2004. doi: 10.1109/ICSMC.2004.1400924
- [9] Cabral, J. E.; Pinto, J. O. P.; Martins, E. M.; Pinto, A. M. A. C. Fraud detection in high voltage electricity consumers using data mining, 2008 IEEE/PES Transmission and Distribution Conference and Exposition, Chicago, IL, pp.1-52008. doi: 10.1109/TDC.2008.4517232
- [10] Spirić, Josif V.; Dočić, Miroslav; Stanković, Slobodan S. Fraud detection in registered electricity time series, in International Journal of Electrical Power & Energy Systems 71, October 2015. doi: 10.1016/j.ijepes.2015.02.037
- [11] Jokar, P.; Arianpoo, N.; Leung, V.C.M. A survey on security issues in smart grids, Secur. Commun. Netw., v.9, pp.262-273, 2012. doi: 10.1002/sec.559
- [12] Faria, R. A. D.; Fonseca, K. V. O.; Schneider, B.; Nguang, S. K. Collusion and Fraud Detection on Electronic Energy Meters - A Use Case of Forensics Investigation Procedures, 2014 IEEE Security and Privacy Workshops, San Jose, CA, pp.65-68, 2014. doi: 10.1109/SPW.2014.19
- [13] PROCEL. Manual de Tarificação da Energia Elétrica, Programa Nacional de Conservação de Energia Elétrica, 2011. <http://www.mme.gov.br/documents/10584/1985241/Manual%20de%20Tarif%20En%20El%20-%20Procel_EPP%20-%20Agosto-2011.pdf>, acesso em 26 de novembro de 2016.
- [14] Kuhn, Max. The caret Package, 2016. <<http://topepo.github.io/caret/index.html>>, acesso em 16 de janeiro de 2017.
- [15] Mello, Flávio Luis de; Carvalho, Roberto Lins de . Rogers' Isomorphism Theorem and Cryptology Applications. Revista IEEE América Latina, v. 13, p. 2009-2016, 2015. doi: 10.1109/TLA.2015.7164229
- [16] M. Buzau, J. Tejedor-Aguilera, P. Cruz-Romero and A. Gómez-Expósito, Detection of Non-Technical Losses Using Smart Meter Data and Supervised Learning, IEEE Transactions on Smart Grid (Early Accept), 2018. doi: 10.1109/TSG.2018.2807925.



Breno Serrano de Araujo é Engenheiro Eletrônico e de Computação (2017) pela Universidade Federal do Rio de Janeiro. Durante a graduação, realizou um ano de intercâmbio acadêmico na Universidade Técnica de Munique (Alemanha), onde teve o primeiro contato com a área de Inteligência Artificial. Atuou como estagiário em uma empresa de consultoria na área de Perdas

Não Técnicas, onde utilizava técnicas de Analytics e Business Intelligence para a otimização de inspeções em campo. Trabalha atualmente na Accenture Analytics Brasil, onde desenvolve modelos e aplicações de Machine Learning e Analytics para clientes de diversos setores de mercado.



Heraldo Luís Silveira de Almeida é Engenheiro Eletrônico (1988), Mestre em Ciências (1991) e Doutor em Ciências em Engenharia Elétrica (2003) pela Universidade Federal do Rio de Janeiro.

Foi engenheiro de desenvolvimento da IBM Brasil e tem experiência em projetos de desenvolvimento de hardware e software nas áreas de energia elétrica, telecomunicações, militar e industrial. Desde 2005 é Professor Associado no Departamento de Engenharia Eletrônica e de Computação da Escola Politécnica da UFRJ, onde tem atuado em projetos de pesquisa e desenvolvimento envolvendo sistemas de controle, instrumentação, processamento de imagens e aplicações de inteligência computacional, principalmente nas áreas de energia elétrica e petróleo. Atua desde 2016 como pesquisador do Laboratório de Inteligência de Máquina e Modelos de Computação (IM²C), desta mesma instituição, com foco de pesquisa em aplicações de Machine Learning.



Flávio Luis de Mello realizou seu DSc. em teoria da computação e processamento de imagens na Universidade Federal do Rio de Janeiro (2006), MSc. em computação gráfica pela Universidade Federal do Rio de Janeiro (2003), graduação em engenharia de sistemas e computação pelo Instituto Militar de Engenharia (1998). Desenvolveu sistemas de comando e controle, implementou sistemas de comunicação durante doze anos no Exército Brasileiro e lecionou no Instituto Militar de Engenharia (IME). Atua no desenvolvimento de aplicações baseadas em prova de teoremas, sistemas baseados em conhecimento e representação do conhecimento. É professor associado do Departamento de Eletrônica e de Computação (DEL) da Escola Politécnica (Poli) na Universidade Federal do Rio de Janeiro (UFRJ) desde 2007 e chefe do Laboratório de Inteligência de Máquina e Modelos de Computação (IM²C) na mesma instituição.