

False Data Injection Attack Detection in State Estimation using Deep Learning

Rakkesh Kumar J

Department of Electrical and Electronics Engineering
College of Engineering Guindy, Anna University, Chennai, Tamil Nadu, India
rakkeshkumar18@gmail.com

Dr.V.Gomathi

Department of Electrical and Electronics Engineering
College of Engineering Guindy, Anna University, Chennai, Tamil Nadu, India
gomesceg@gmail.com

Arun Jeas

Department of Electrical and Electronics Engineering
College of Engineering Guindy, Anna University, Chennai, Tamil Nadu, India
arunjeas@gmail.com

Article Info

Volume 81

Page Number: 1677– 1684

Publication Issue:

November-December 2019

Abstract

Power grids are a complex system with a number of substations and transmission lines. Existing grids are converted into the smart grid due to the advent of digitization. For efficient control of the system, accurate state estimation is mandatory for which accurate readings of Phasor Measurement Unit becomes the base. State estimators are prone to False Data Injection(FDI) attack as it can pass through bad data detection mechanism. Covert cyber assaults framed by hackers, who have a deep understanding of power system, are dangerous as it cannot be detected by state estimators which results in the catastrophic effect. Thus an unsupervised attack detection algorithm is developed using autoencoder, which identifies the attack by examining the latest historical data and detect the state vectors which are not identical to the normal state vectors. This method is tested on IEEE 3 bus system with a wide range of attack plot.

Article History

Article Received: 5 March 2019

Revised: 18 May 2019

Accepted: 24 September 2019

Publication: 09 December 2019

I. INTRODUCTION

State estimation is the key function for obtaining such a real-time network model. In order to have successful power system operation its real-time model is required. As power systems are becoming a cyber-physical system, the need to look at it from cyber point of view and to secure the system from cyber-attack is a must. There are many types of attack that can happen in a power system such as denial of service attack (data availability violation), False Data Injection (FDI) attacks, and disclosure attack. FDI attacks

estimated states of the system are manipulated, having high destructive effect on the power system stability and control. In this paper first we look FDI attack model, then followed by the algorithm to detect FDI attack using autoencoder and then we discuss the result obtained from using the algorithm in R-software followed by conclusion.

II. FDI ATTACK MODEL

Here we discuss how FDI attack is modelled using DC state estimation, as in state estimation we estimate the state variables let x

denote vector of state variable $x = [\theta_2, \dots, \theta_n]$, n is number of state variables, bus 1 is considered as reference bus $\theta_1 = 0$. z denote measurement vector $z = [z_1, \dots, z_m]$ m is number of measurements ($m > n$). DC state estimation are formalised as

$$z = Hx + e \quad [1]$$

H is $(m \times n)$ jacobian matrix obtained from structure of power system, e is $(m \times 1)$ vector of measurements errors following Gaussian distribution with zero mean. As per weighted least square approach

$$\hat{x} = (H^T R^{-1} H)^{-1} H^T R^{-1} z \quad [2]$$

R is $m \times m$ matrix, which is error covariance matrix. Bad data is detected using euclidean norm of residual by comparing it with threshold τ .

$$r = z - H\hat{x} \quad [3]$$

$$||r|| < \tau \quad [4]$$

$$z_{bad} = z + a$$

where,

$a = Hc$, c is arbitrary vector, $c = [c_1, \dots, c_n]^T$ which is injected maliciously. Using z_{bad} will result in false estimates $\hat{x}_{bad} = \hat{x}c$, where $\hat{x}$ is the estimate of x when using original measurement z .

$$r_{bad} = z_{bad} - H\hat{x}_{bad} = z + a - H(\hat{x} + c) \quad [6]$$

$$z - H\hat{x} + (a - Hc) = z - H\hat{x} = r \quad [7]$$

this results in no variation in residual r making it similar to normal one, so detection mechanism cannot detect FDI attack.

III. ANOMALOUS DATA IDENTIFICATION USING AUTOENCODER

Autoencoders are neural networks which can be designed for either shallow or deep learning characteristics. Autoencoders differ from other forms of neural network as they are trained to reproduce the inputs. Thus the hidden layers and neurons are not maps between an input and some other outcome, but are self (auto)-encoding. The best autoencoder is the network which learns meaningful structure in the training data or one that reduces noise, identifies outliers or anomalous data.

Step1: Training Data from PMU is applied to the auto-encoder to train the network

Step2: Auto-encoder process the data and tries reproduce the input at output by updating its weights

Step3: The test data is applied to the trained network, error between the output of autoencoder and its input is called as reconstruction error which gives the information about similarity between training and test data

Step 4: If the reconstruction error is high then threshold value (F-test) the state is considered as outlier

Step 5: Discard the outlier data

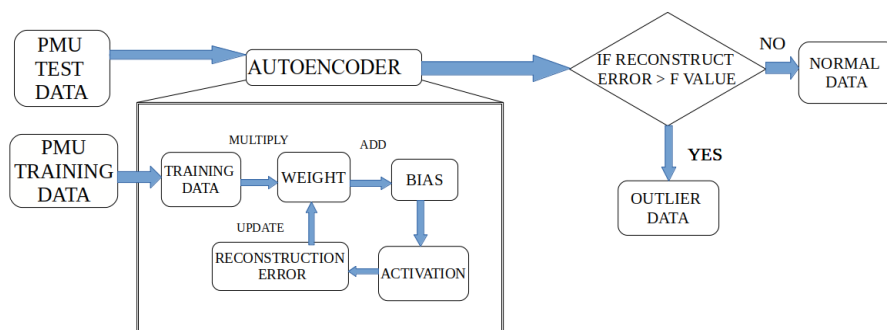


Fig.1 Anomalous Data Identification Using Auto Encoder

IV.RESULTS AND DISCUSSION

The algorithm to detect the FDI attack has been developed and it's tested against various attack scenarios by conducting FDI attack on IEEE 3bus system data from MATPOWER. Results were obtained using R-software.A

training data set with 144 data is obtained by conducting state estimation at various load condition on IEEE 3 bus system available in MATPOWER is applied to autoencoder. A sample of 24 data from the training data is taken to validate the performance of autoencoder from reconstruct error.

```
Model Details:
=====
H2OAutoEncoderModel: deeplearning
Model ID: DeepLearning_model_R_1556769610690_5
Status of Neuron Layers: auto-encoder, gaussian distribution, Quadratic loss, 12 weights/biases, 2.9 KB, 144,000
training samples, mini-batch size 1
layer units type dropout 11 12 mean_rate rate_rms momentum
1 1 2 Input 0.00 % NA NA NA NA NA
2 2 2 Tanh 0.00 % 0.000000 0.000000 0.055260 0.020029 0.000000
3 3 2 Tanh NA 0.000000 0.000000 0.068135 0.021026 0.000000
mean_weight weight_rms mean_bias bias_rms
1 NA NA NA NA
2 -0.089542 0.735054 0.568510 0.108843
3 0.137370 1.007000 -0.017521 0.711581

H2OAutoEncoderMetrics: deeplearning
** Reported on training data. **

Training Set Metrics:
=====
MSE: (Extract with `h2o.mse`) 0.0001171259
RMSE: (Extract with `h2o.rmse`) 0.01082247

H2OAutoEncoderMetrics: deeplearning
** Reported on validation data. **

Validation Set Metrics:
=====
MSE: (Extract with `h2o.mse`) 0.0001322721
RMSE: (Extract with `h2o.rmse`) 0.01150096
```

Fig 2 Training data set output of autoencoder in R

The result shows the training set metrics and validation set metrics. As the data has not been manipulated mean square error of validation

metrics are very low indicating auto-encoder reconstruct the data more accurately.

I. Results of test data set 1 with FDI attack on data 1,3,9,18,24

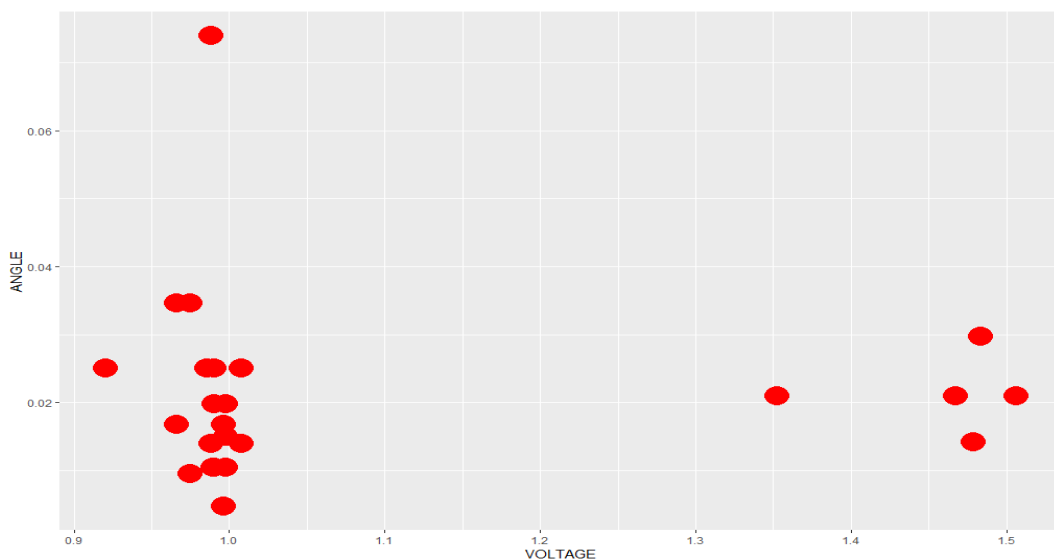


Fig 3.Test data set 1 with FDI attack on data 1,3,9,18,24

In this test case, 5 states have been modified to model FDI attack. The attacked samples are indicated by the second cluster on right side. The

attacked samples have value high then normal indicating high end attack.

Model Details:

=====

H2OAutoEncoderModel: deeplearning
Model ID: DeepLearning_model_R_1556769610690_6
Status of Neuron Layers: auto-encoder, gaussian distribution, Quadratic loss, 12 weights/biases, 2.9 KB, 144,000 training samples, mini-batch size 1

layer	units	type	dropout	l1	l2	mean_rate	rate_rms	momentum
1	1	2 Input	0.00 %	NA	NA	NA	NA	NA
2	2	2 Tanh	0.00 %	0.000000	0.000000	0.017492	0.003204	0.000000
3	3	2 Tanh	NA	0.000000	0.000000	0.052213	0.012946	0.000000

mean_weight weight_rms mean_bias bias_rms

1	NA	NA	NA	NA
2	0.234819	0.600031	0.000195	0.411829
3	1.105582	0.712223	-0.172336	0.383664

H2OAutoEncoderMetrics: deeplearning
** Reported on training data. **

Training Set Metrics:

=====

MSE: (Extract with `h2o.mse`) 0.0002679164
RMSE: (Extract with `h2o.rmse`) 0.01636815

H2OAutoEncoderMetrics: deeplearning
** Reported on validation data. **

Validation Set Metrics:

=====

MSE: (Extract with `h2o.mse`) 1.544416
RMSE: (Extract with `h2o.rmse`) 1.242745

Fig 4.Autoencoder output for test data set 1

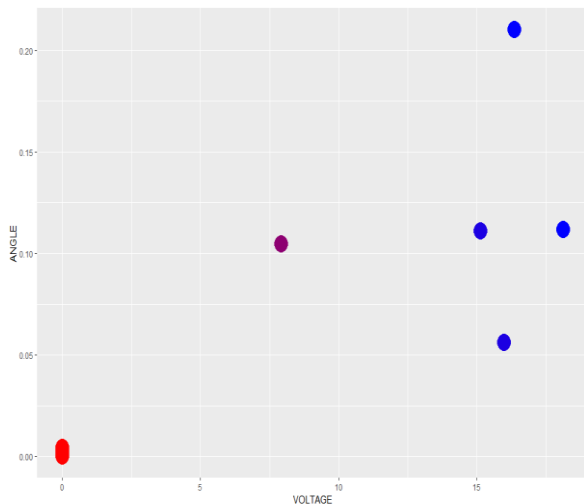


Fig 5.Reconstruct error for test data set 1

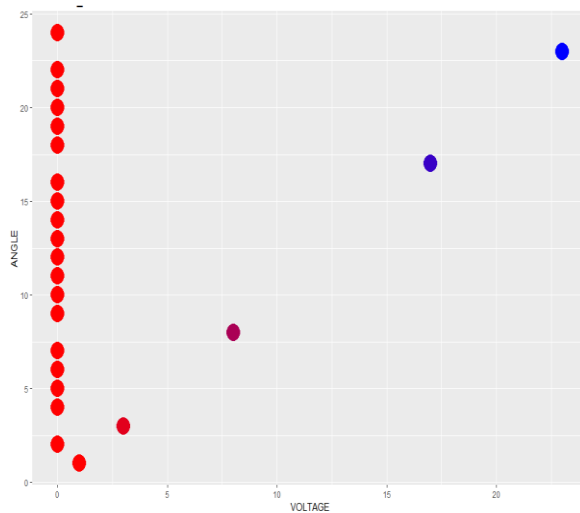


Fig 6.Outlier detection in test data set 1

The reconstruct error for data set 2 is very high in range 5 to 15 indicating the data is attacked. The outliers are the points which have reconstruction error above the threshold value. Outlier points are indicated with various colours according to their location.

II. Results of test data set 2 with FDI attack on data 4,7,16,17,25

In this test case, 5 states have been modified to model FDI attack. The attacked samples are indicated by the first cluster on left side. The attacked samples have value lower than normal indicating low end attack.

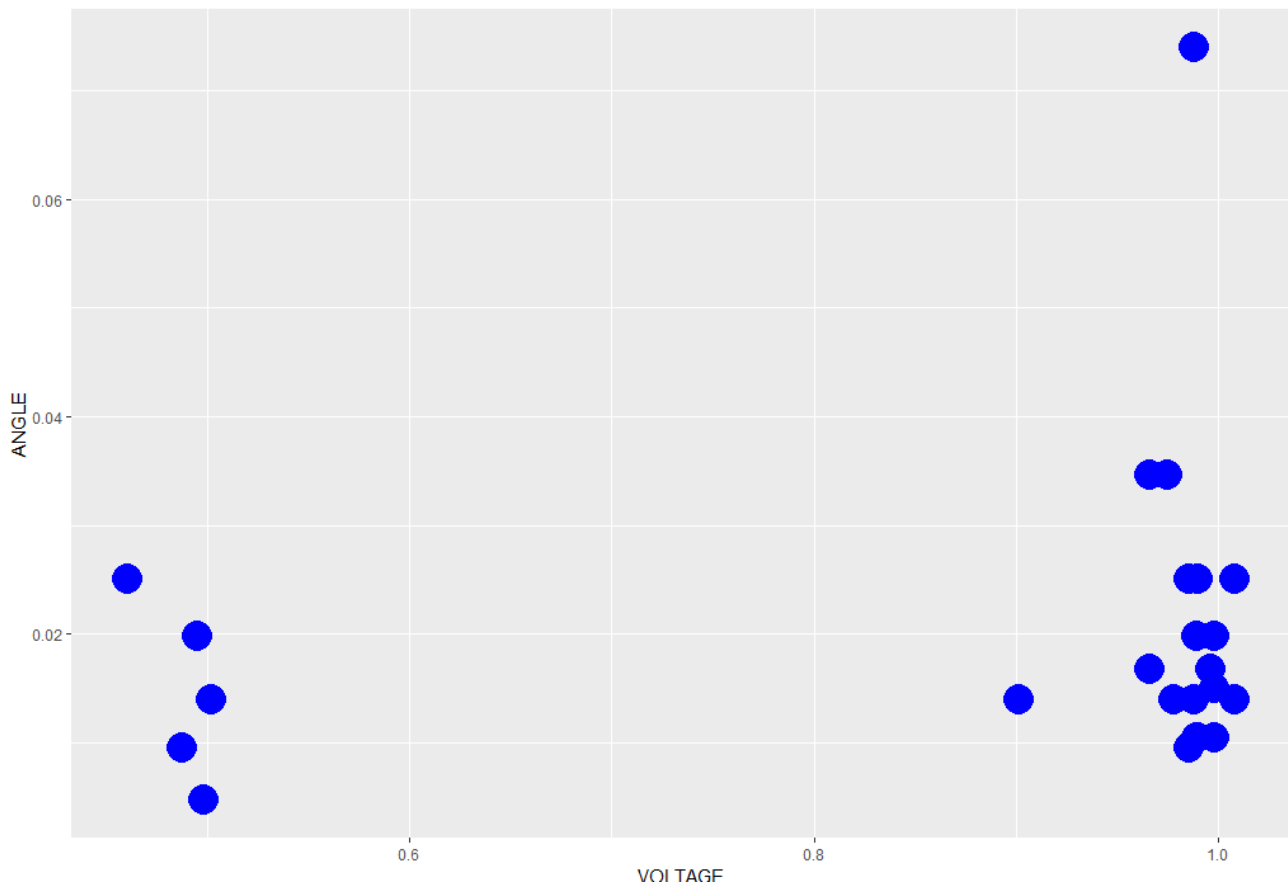


Fig 7. Test data set 2 with FDI attack on data 4,7,16,17,25

```

Model Details:
=====
H2OAutoEncoderModel: deeplearning
Model ID: DeepLearning_model_R_1556769610690_7
Status of Neuron Layers: auto-encoder, gaussian distribution, Quadratic loss, 12 weights/biases, 2.9 KB, 144,000
training samples, mini-batch size 1
layer units type dropout l1 l2 mean_rate rate_rms momentum
1 1 2 Input 0.00 % NA NA NA NA NA
2 2 2 Tanh 0.00 % 0.000000 0.000000 0.066823 0.015801 0.000000
3 3 2 Tanh NA 0.000000 0.000000 0.053424 0.010237 0.000000
mean_weight weight_rms mean_bias bias_rms
1 NA NA NA NA
2 -0.061885 0.607294 0.553300 0.150438
3 -0.021679 1.135541 0.210951 0.768226

H2OAutoEncoderMetrics: deeplearning
** Reported on training data. **

Training Set Metrics:
=====
MSE: (Extract with `h2o.mse`) 9.633452e-05
RMSE: (Extract with `h2o.rmse`) 0.009815015

H2OAutoEncoderMetrics: deeplearning
** Reported on validation data. **

Validation Set Metrics:
=====
MSE: (Extract with `h2o.mse`) 1.372536
RMSE: (Extract with `h2o.rmse`) 1.171553
    
```

Fig 8. Autoencoder output for test data set 2

As the data has been manipulated mean square error of validation metrics are high indicating FDI attacks.

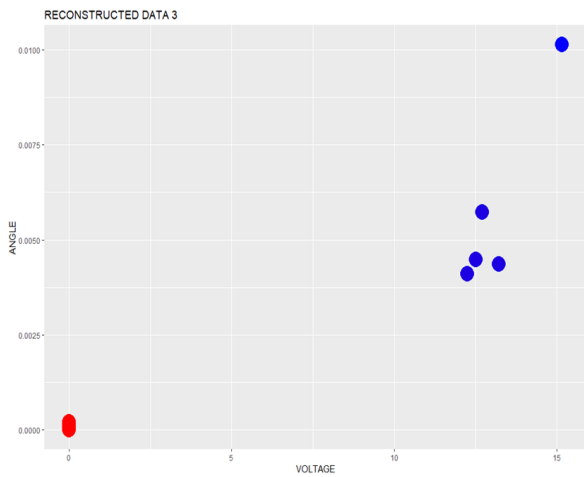


Fig 9.Reconstruct error for test data set 2

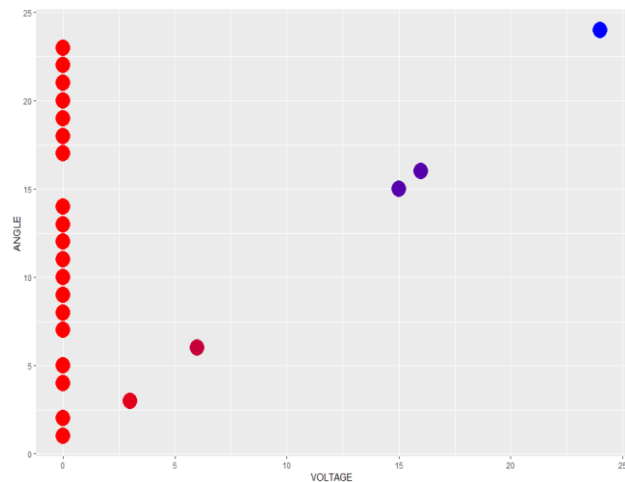


Fig 10.Outlier detection in test data set 2

The reconstruct error for data set 3 is very high in range 10 to 15 indicating the data is attacked.

iii. Result of test data set 3 with FDI attack on data 2,6,15,21,23

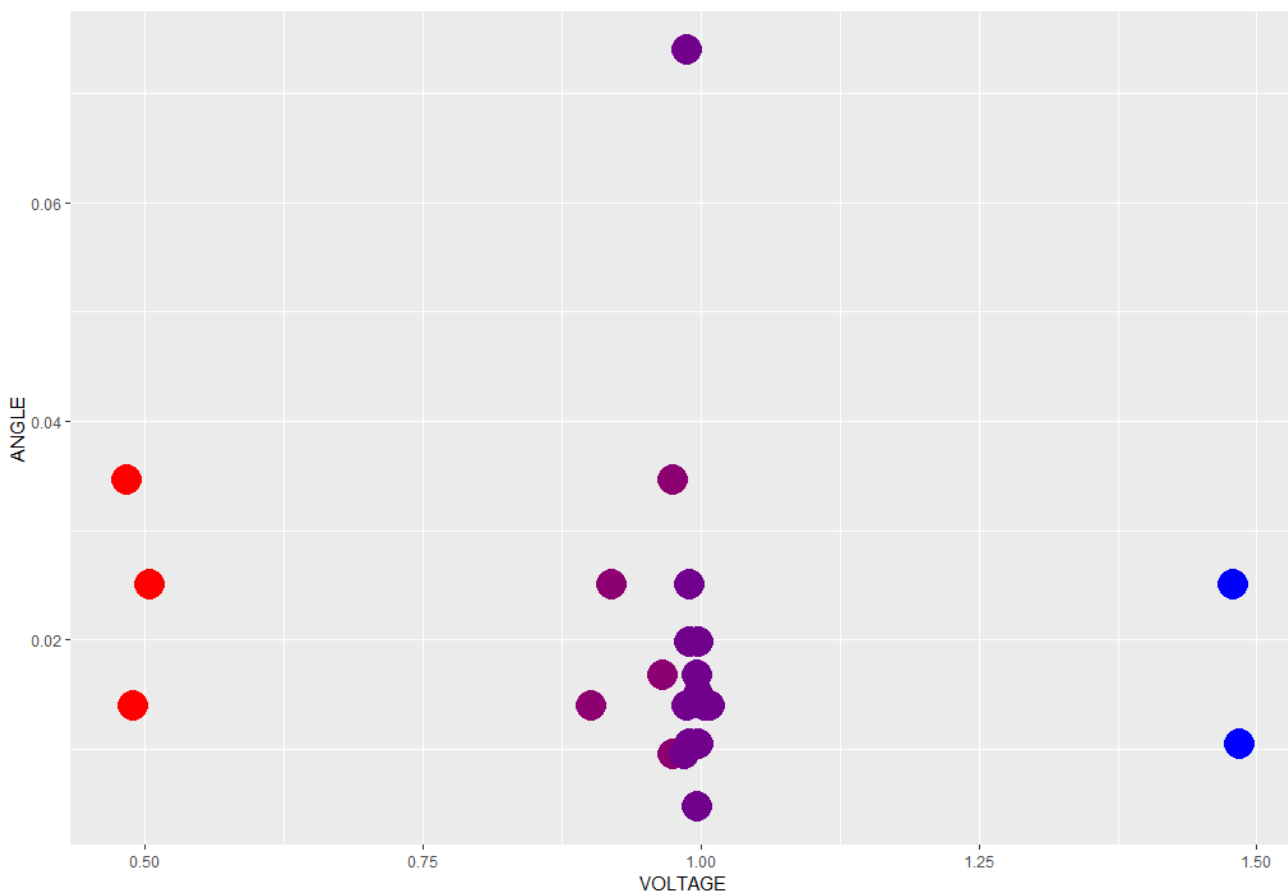


Fig 11.Test data set 3 with FDI attack on data 2,6,15,21,23

In this test case, 5 states have been modified to model FDI attack. The attacked samples are indicated by the two clusters on right and left side.

The attacked samples have value higher and lower than normal indicating both high end and low end attack.

Model Details:

=====

H2OAutoEncoderModel: deeplearning

Model ID: DeepLearning_model_R_1556769610690_9

Status of Neuron Layers: auto-encoder, gaussian distribution, Quadratic loss, 12 weights/biases, 2.9 KB, 144,000 training samples, mini-batch size 1

layer	units	type	dropout	l1	l2	mean_rate	rate_rms	momentum
1	1	2 Input	0.00 %	NA	NA	NA	NA	NA
2	2	2 Tanh	0.00 %	0.000000	0.000000	0.066173	0.026223	0.000000
3	3	2 Tanh	NA	0.000000	0.000000	0.056537	0.003861	0.000000

	mean_weight	weight_rms	mean_bias	bias_rms
1	NA	NA	NA	NA
2	0.075950	0.629848	-0.576195	0.107478
3	-0.114582	1.141761	0.015287	0.801486

H2OAutoEncoderMetrics: deeplearning

** Reported on training data. **

Training Set Metrics:

=====

MSE: (Extract with `h2o.mse`) 9.111372e-05

RMSE: (Extract with `h2o.rmse`) 0.009545351

H2OAutoEncoderMetrics: deeplearning

** Reported on validation data. **

Validation Set Metrics:

=====

MSE: (Extract with `h2o.mse`) 1.456499

RMSE: (Extract with `h2o.rmse`) 1.206855

Fig 12.Autoencoder output for test data set 3

The reconstruct error for data set 4 is very high in range 10 to15 indicating the data is attacked.

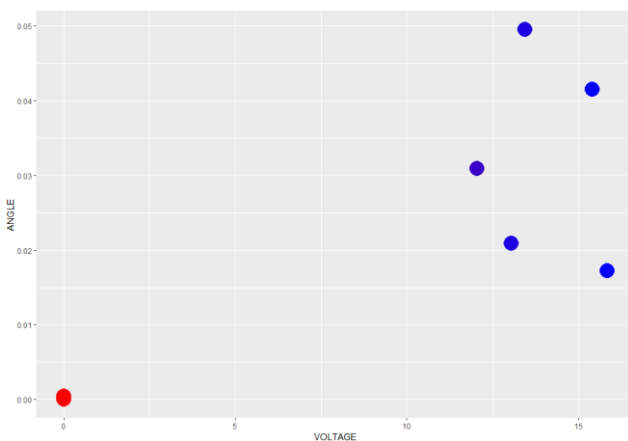


Fig 13.Reconstruct error for test data set 2

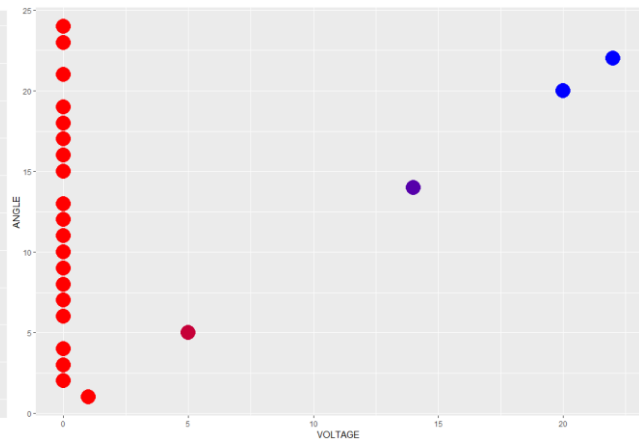


Fig 14.Outlier detection in test data set 2

Table 1. FDI DETECTION RESULT

DATA	ORIGINAL DATA*0.5 DETECTION %	ORIGINAL DATA *1.5 DETECTION %
Test set 1	99%	100%
Test set 2	99%	98%

Test set 3	100%	100%
------------	------	------

V.CONCLUSION

False data injection attack can cause catastrophic damage to the power system by creating false state estimates. The proposed method detects the FDI attack accurately, so the attacked data's can be prevented from entering the control systems. Detection accuracy can be further improved by using more than one deep learning algorithm and comparing the results.

VI.REFERENCES

1. H.Wang, J.Ruan, G.Wang and B.Zhou, "Deep learning-based interval state estimation of AC smart grids against sparse cyber attacks,"IEEE Trans.Industrial Informatics, vol.14, Nov.2018,pp,4766-4778.
2. D.Wilson,Y.Tang,J.Yan,Z.Lu, "Deep Learning-Aided Cyber-Attack Detection in Power Transmission Systems,"IEEE Int.Conf Power & Energy Society General Meeting ,978-1-5386-7703-2, 2018.
3. Mostafa.M,A.Sami,Y.Weng,"Identification of False Data Injection Attacks With Considering the Impact of Wind Generation and Topology Reconfigurations," IEEE Trans.Sustainable Energy, vol. 9, no. 3, Jul,2018.
4. A.Abdallah and X.S.Shen,"Efficient prevention technique for false data injection attack in smart grid," in Proc. IEEE Int. Conf. Commun., Kuala Lumpur, Malaysia, 2016, pp. 1–6.
5. G.Chaojun, P.Jirutitijaroen, and M.Motani, "Detecting false data injection attacks in AC state estimation," IEEE Trans.Smart Grid,vol.6,no.5, pp. 2476–2483, Sep. 2015.
6. S.Bi and Y.J.Zhang, "Graphical methods for defense against false-data injection attacks on power system state estimation," IEEE Trans. Smart Grid, vol. 5, no. 3, pp. 1216–1227, May 2014.
7. Q.Yang, J.Yang, W.Yu, D.An, N.Zhang, and W.Zhao, "On false data-injection attacks against power system state estimation: Modeling and counter measures," IEEE Trans. Parallel Distrib. Syst., vol. 25, no. 3, pp. 717–729, Mar. 2014.
8. R. Zimmerman and D. Gan, "MATPOWER: A MATLAB power system simulation package." [Online]. Available: [http://www.pserc.cornell.edu.matpower](http://www.pserc.cornell.edu/matpower), Dec. 2016