

Machine Learning and Fraud Detection

February 2020

Tamsin Crossland – Senior Architect

@CrosslandTamsin



World Class **Payment** and **Enterprise Solutions**
for the global financial sector

Two main types of article on AI



architectures essentially facilitate diffusion of information across the graph. The more GCN layers an architecture contains, the more the information is smoothed. It is therefore important to recognize that the performance of a traditional GCN relies heavily on a high degree of intraclass clustering [7]. This is a result of the convolutional kernel being determined by the adjacency matrix A , which is inherently fixed. Also note that all vertices are either binary encoded or represented by a single weight. This is a limiting aspect when it comes to graphs containing more detailed information between vertices.

Frameworks

Standard GCN (as displayed in Equation 1) is a special case of a message-passing framework. The difference in notation is that an edge is represented at the node level. This is more in line with the architectures proposed in this paper and will therefore be used for the remainder of the paper.

Given a binary adjacency matrix, Equation 1 can be written as:

$$h_i^{(l+1)} = \sigma \left(\sum_{j \in N_i} h_j^{(l)} \right) W^{(l)}, \quad (2)$$

where N_i is the set of nodes v_j in layer l , N_i denotes the set of neighbors of node v_i , and c_i is a node-specific normalization constant, typically takes a value of $|N_i| + 1$ (including self-connections), in case of an undirected graph.

Given weights w_{ij} , Equation 2 can be expressed as:

$$h_i^{(l+1)} = \sigma \left(\sum_{j \in N_i} w_{ij} h_j^{(l)} \right) W^{(l)}, \quad (3)$$

where W is a weight matrix. To allow vertices to retain their original properties

in the context of message passing frameworks, this leads to the following propagation rule:

$$h_i^{(l+1)} = \sigma \left(h_i^{(l)} W_s^{(l)} + \sum_{r \in \mathcal{R}} \sum_{j \in N_i^r} \frac{1}{c_{i,r}} h_j^{(l)} W_r^{(l)} \right), \quad (5)$$

where $W_r \in \mathbb{R}^{|\mathcal{F}^{(l)}| \times |\mathcal{F}^{(l+1)}|}$ now are *relation-specific* weight matrices, W_s denotes the weight matrix associated with self-connections (both trainable), and $c_{i,r}$ are relation-specific normalization constants, typically of value $|N_i^r|$. In order to facilitate the directional nature of some interactions, the set of relations $\mathcal{R}$ was expanded to contain a version in both edge directions (incoming and outgoing), by means of both canonical and inverse variations of each relation.

3 LATENT GRAPHS

3.1 Pseudo-relations

We can extend the multi-relational approach of R-GCNs to the case of edges being represented by a vector w_{ij} containing multiple weights across different edge attributes. These attributes can effectively be regarded as non-binary *pseudo-relations* $\mathcal{R}$. The propagation rule now takes a form akin to Equation 3:

$$h_i^{(l+1)} = \sigma \left(\frac{1}{c_i} \sum_{r \in \mathcal{R}} \left[w_s^r h_i^{(l)} + \sum_{j \in N_i} w_{ij}^r h_j^{(l)} \right] W_r^{(l)} \right), \quad (6)$$

with c_i now taking the form of:

$$c_i = \sum_{r \in \mathcal{R}} \left(w_s^r + \sum_{j \in N_i} w_{ij}^r \right). \quad (7)$$

Note that $W_s^{(l)}$ from Equation 5 has been absorbed into $W_r^{(l)}$. The weight of self-connections can now be learned by means of a self-weight vector w_s . In essence, this method obtains an updated node representation for every pseudo-relation separately, after which a weighted summation takes place in order to arrive at the final embedding. The original publication by Schlichtkrull et al. [10] alludes to the possibility of replacing the weighted sum with a more elaborate learning mechanism, such as a fully connected layer, but this was left for future work.

Machine Learning and Fraud Detection

- Payments
- Demonstration
- Thoughts

Machine Learning and Fraud Detection

- Payments
- Demonstration
- Thoughts



The increasing scale, diversity, and complexity of fraud.



- Vulnerabilities in payments services have increased as the shift to digital and mobile customer platforms accelerates.

The increasing scale, diversity, and complexity of fraud.



- New solutions have also led to payments transactions being executed more quickly, leaving banks and processors with less time to identify, counteract, and recover the underlying funds when necessary.

The increasing scale, diversity, and complexity of fraud.



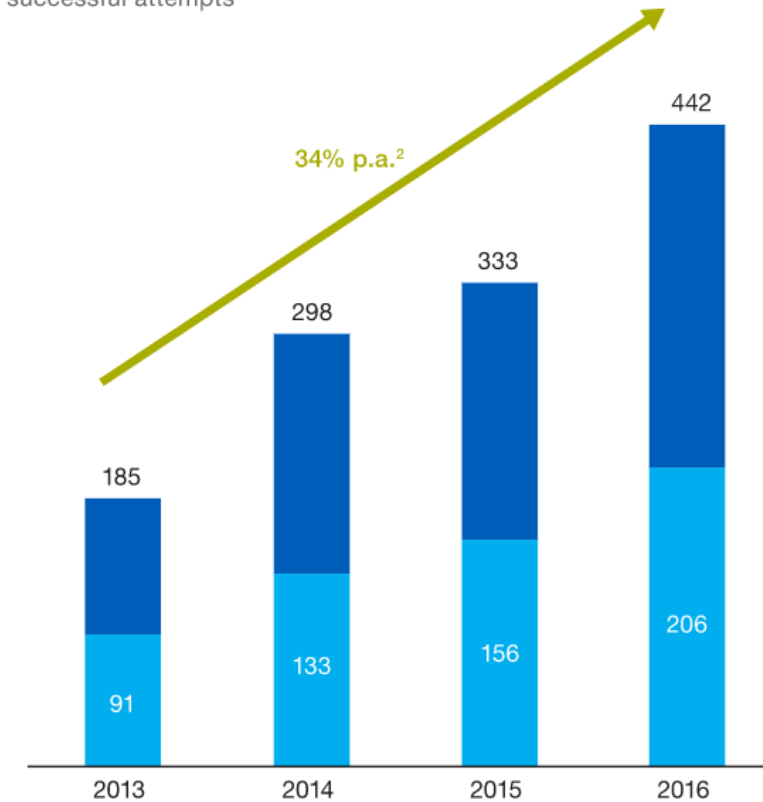
- The sophistication of fraud has increased:
 - greater collaboration among bad actors, including:
 - the exchange of stolen data,
 - new techniques, and
 - expertise on the dark web.

The fraud threat facing banks and payments firms has grown dramatically in recent years.

Estimates of fraud's impact on consumers and financial institutions vary significantly but losses to banks alone are conservatively estimated to exceed \$31 billion globally by 2018.

Average number of fraudulent transactions attempted per merchant per month¹

■ Average successful attempts



¹Weighted merchant responses to LexisNexis survey questions: In a typical month, approximately how many fraudulent transactions are prevented by your company/successfully completed by fraudsters? What is the average value of successful fraud transactions?

²Per annum.

McKinsey&Company | Source: 2016 LexisNexis True Cost of Fraud Study; US Department of Commerce

Instant Payments



IPF



Rule Based Systems



Example: if a credit card transaction is more than ten times larger than the average for this customer

Allow the human experts to apply their subject matter expertise.

Difficult and time-consuming to implement well.

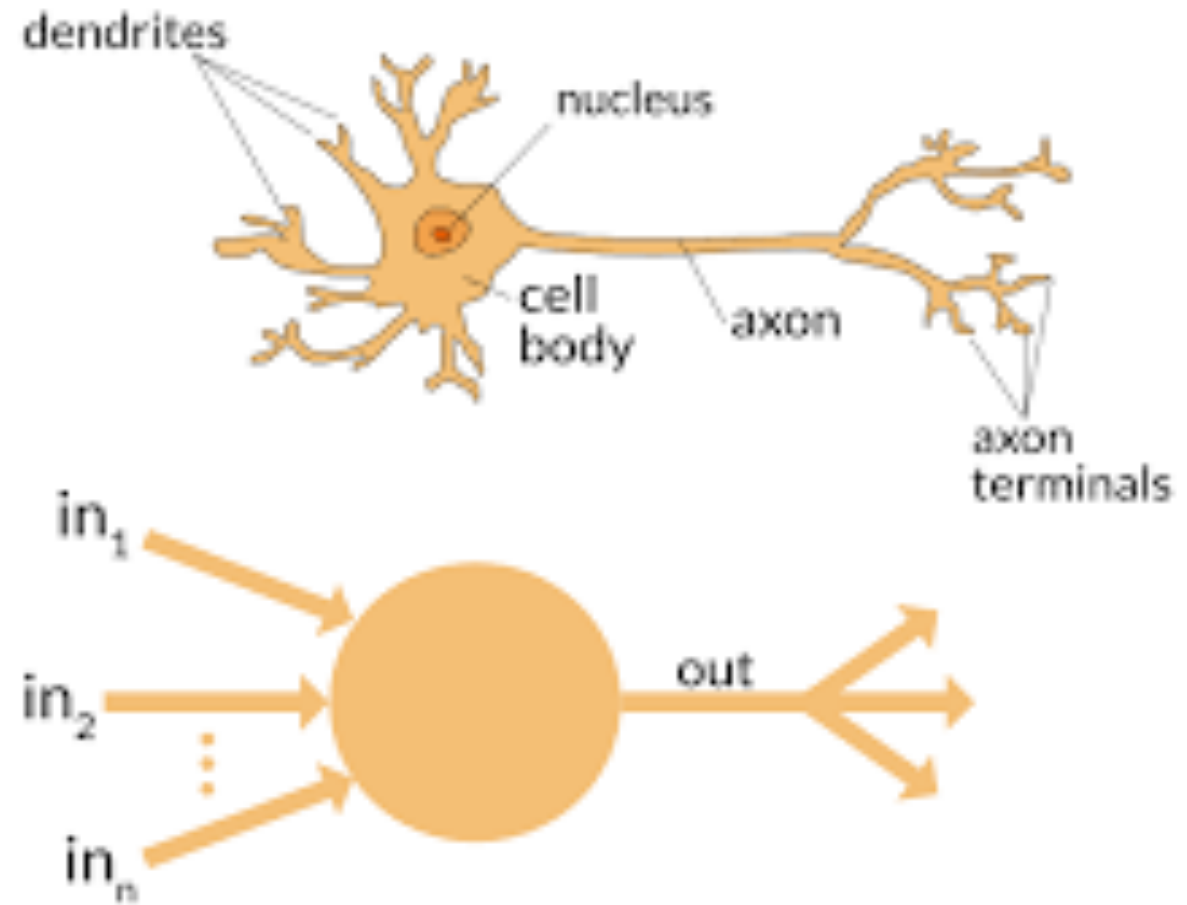
Includes the painstaking definition of every single rule for anomaly possible

If experts make an omission, undetected anomalies will happen and nobody will suspect it.

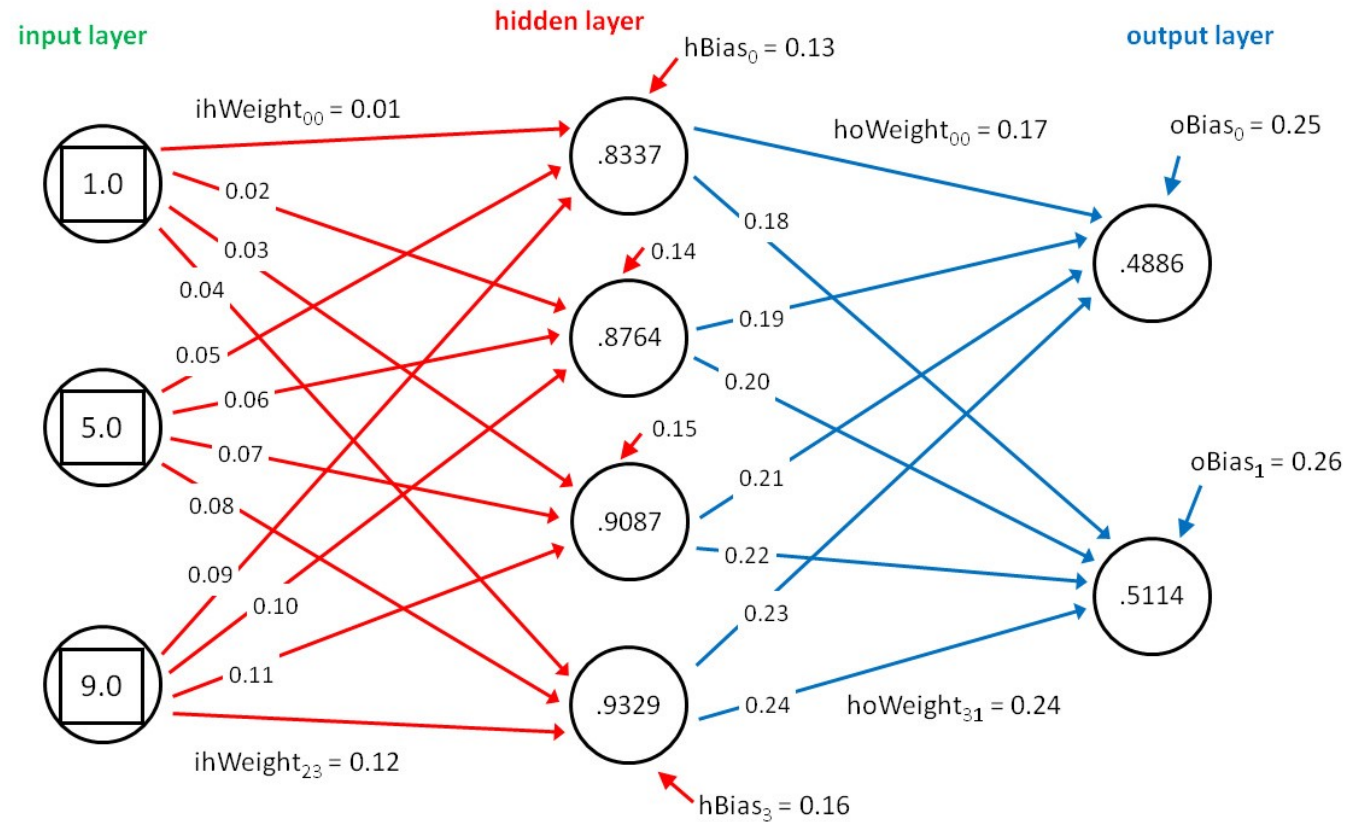
Today, legacy systems apply about 300 different rules on average to approve a transaction



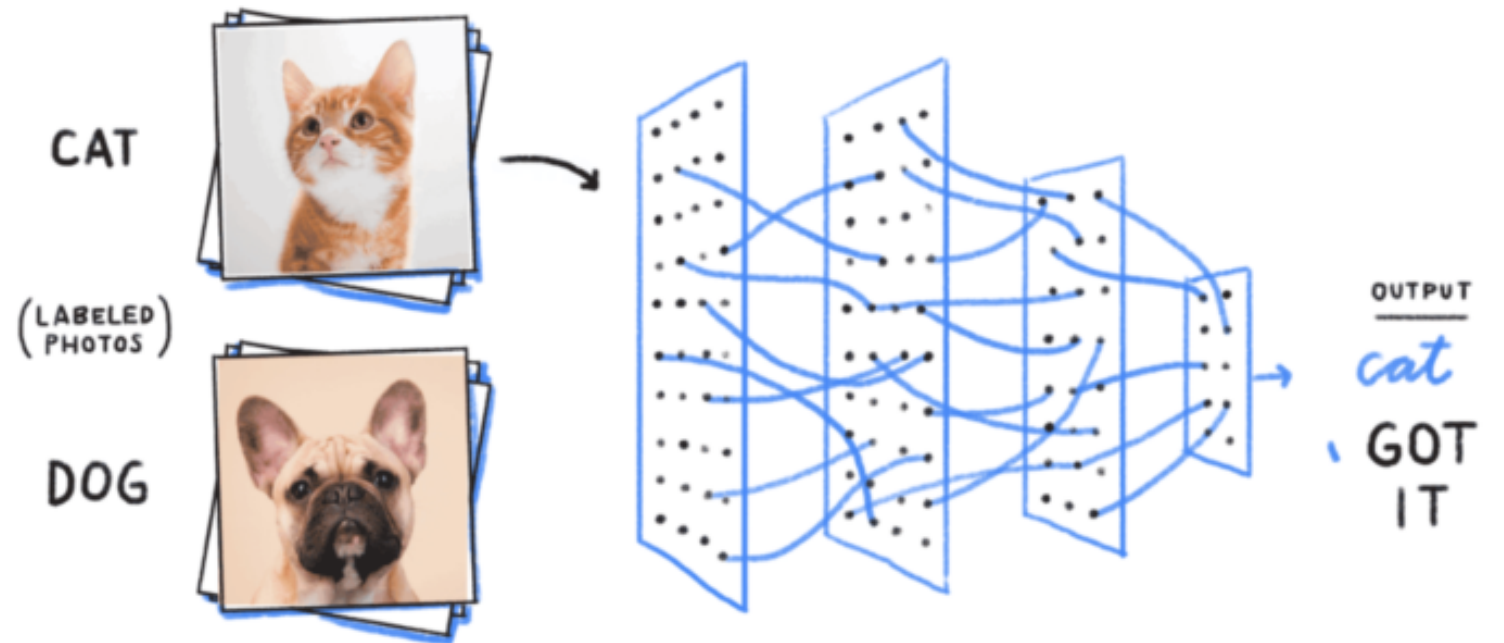
Neural Network



Weights and Biases



Training

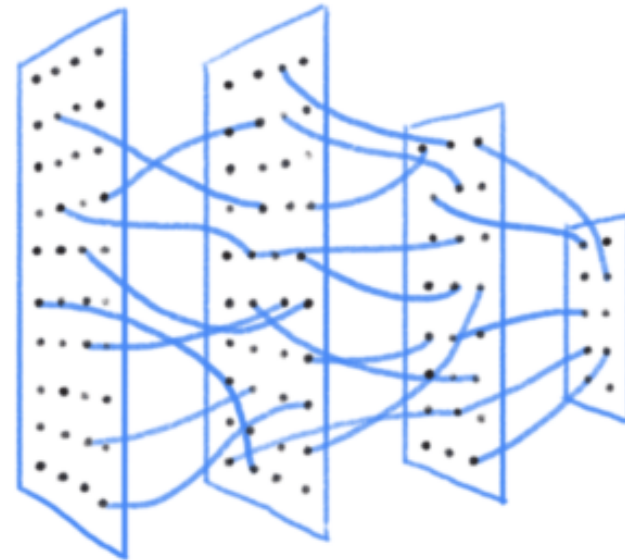


Training

Fraudulent Transaction
(Labeled PHOTOS)



Non-Fraudulent Transaction



OUTPUT

Fraud
GOT
IT

Use Case

CONTEXT

12

Data: Bank's 12-months control environment

100000000

Number of bank transactions: +10 million unique payments

32%

Previous rule-based control environment could analyze only 32% of unique payments

100%

NetGuardians' machine learning risk platform 100% - Analysis of all unique payments

RESULTS

83%

83% reduction in the number of false positives compared to the previous rule-based control environment

93%

93% reduction in fraud investigation time

118%

118% all fraud cases caught by the previous rule-based control environment also detected by NetGuardians. Plus, new fraud cases are discovered.

Rule Based versus Machine Learning

Rule Based	Machine Learning
Catches obvious fraudulent scenarios	Finds hidden correlations in data
Large amount of manual work to enumerate all possible detection rules	Automatic detection of possible fraud scenarios
Easier to explain	More difficult to explain

Machine Learning and Fraud Detection

- Payments
- **Demonstration**
- Thoughts

Install Tensorflow

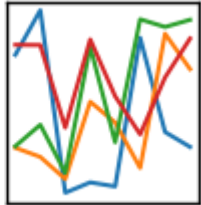
```
tamsincrossland@A318 MINGW64 ~/Documents/AI/creditcardfraud  
$ winpty docker run -it --rm -v /${PWD}:/tf/tensorflow-tutorials/wkDir -p 8888:8888 tensorflow/tensorflow:latest-py3-jupyter
```

TensorFlow

Install Libraries

pandas

$$y_{it} = \beta' x_{it} + \mu_i + \epsilon_{it}$$



Data Analysis

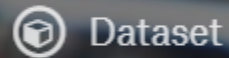


data mining and data analysis

```
winpty docker exec -i -t 07a24f61e7b6 bash
```

```
pip install pandas
```

```
pip install -U scikit-learn
```



Credit Card Fraud Detection

Anonymized credit card transactions labeled as fraudulent or genuine

Contains two days worth of credit card transactions made in September 2013 by European cardholders.

492 frauds out of 284,807 transactions (0.172%).

Contains only numerical input variables which are the result of a Principal Component Analysis transformation (a method of extracting relevant information from confusing data sets).

Due to confidentiality issues, cannot provide the original features and more background information about the data.

Time	Amount	V1	V2	V3	V4	V5	V6	V7	V8	V9	V10	V11	V12	V13	V14	V15	V16	V17	V18	V19	V20	V21	V22	V23	V24	V25	V26	V27	V28	Class	
0	149.62	-1.35981	-0.07278	2.53635	1.37816	-0.33832	0.46239	0.2396	0.0987	0.36379	0.09079	-0.5516	-0.6178	-0.99139	-0.31117	1.46818	-0.4704	0.20797	0.02579	0.40399	0.25141	-0.01831	0.27784	-0.11047	0.06693	0.12854	-0.18911	0.13356	-0.02105	0	
0	2.69	1																												0	
1	378.66	-1																												0	
1	123.5	-0																												0	
2	69.99	-1.15823	0.87774	1.54872	0.40303	-0.40719	0.09592	0.59294	-0.27053	0.81774	0.75307	-0.82284	0.5382	1.34585	-1.11967	0.17512	-0.45145	-0.23703	-0.03819	0.80349	0.40854	-0.00943	0.79828	-0.13746	0.14127	-0.20601	0.50229	0.21942	0.21515	0	
2	3.67	-0.42597	0.96052	1.14111	-0.16825	0.42099	-0.02973	0.4762	0.26031	-0.56867	-0.37141	1.34126	0.35989	-0.35809	-0.13713	0.51762	0.40173	-0.05813	0.06865	-0.03319	0.08497	-0.20825	-0.55982	-0.0264	-0.37143	-0.23279	0.10591	0.25384	0.08108	0	
4	4.99																		282	-0.61199	-0.04558	-0.21963	-0.16772	-0.27071	-0.1541	-0.78006	0.75014	-0.25724	0.03451	0.00517	0
7	40.8																		213	-0.35822	0.3245	-0.15674	1.94347	-1.01545	0.0575	-0.64971	-0.41527	-0.05163	-1.20692	-1.08534	0
7	93.2																		977	0.11876	0.57033	0.05274	-0.07343	-0.26809	-0.20423	1.01159	0.3732	-0.38416	0.01175	0.1424	0
9	3.68	-0.55620	1.11555	1.04457	0.22215	0.45550	-0.24070	0.05150	0.00554	-0.75075	-0.50085	1.01701	0.65055	1.00084	-0.44552	0.15022	0.75345	-0.54098	0.47668	0.45177	0.20371	-0.24691	-0.63375	-0.12079	-0.38505	-0.06973	0.0942	0.24622	0.08308	0	
10	7.9	1.44004	-1.17634	0.01305	-1.37567	-1.07130	-0.65015	-1.41324	0.04846	-1.72041	1.67666	1.10064	-0.67144	-0.51305	-0.00505	0.23003	0.03107	0.25241	0.85434	-0.22137	-0.38723	-0.00093	0.31389	0.02774	0.50051	0.25137	-0.12948	0.04285	0.01625	0	
																			99	0.70766	0.12599	0.04992	0.23842	0.00913	0.99671	-0.76731	-0.49221	0.04247	-0.05434	0	
																			79	-0.68319	-0.10276	-0.23181	-0.48329	0.08467	0.39283	0.16113	-0.35499	0.02642	0.04242	0	
																			05	-0.98292	-0.1532	-0.03688	0.07441	-0.07141	0.10474	0.54826	0.10409	0.02149	0.02129	0	
																			27	2.22187	-1.58212	1.15166	0.22218	1.02059	0.02832	-0.23275	-0.23556	-0.16478	-0.03015	0	
																			99	0.43254	0.26345	0.49962	1.35365	-0.25657	-0.06508	-0.03912	-0.08709	-0.181	0.12939	0	
																			91	-0.57568	-0.11391	-0.02461	0.196	0.0138	0.10376	0.3643	-0.38226	0.09281	0.03705	0	
																			47	0.02544	-0.04702	-0.1948	-0.67264	-0.15686	-0.88839	-0.34241	-0.04903	0.07969	0.13102	0	
14	46.8	-5.40126	-5.45015	1.1863	1.73624	3.04911	-1.76341	-1.55974	0.16084	1.23309	0.34517	0.91723	0.97012	-0.26657	-0.47913	-0.52661	0.472	-0.72548	0.07508	-0.40687	-2.19685	-0.5036	0.98446	2.45859	0.04212	-0.48163	-0.62127	0.39205	0.94959	0	
																														0	
																														0	
																														0	
18	22.75	0.24749	0.27767	1.18547	-0.0926	-1.31439	-0.15012	-0.94636	-1.61794	1.54407	-0.82988	-0.5832	0.52493	-0.45338	0.08139	1.5552	-1.39689	0.78313	0.43662	2.17781	-0.23098	1.65018	0.20045	-0.18535	0.42307	0.82059	-0.22763	0.33663	0.25048	0	
22	0.89	-1.94653	-0.0449	-0.40557	-1.01306	2.94197	2.95505	-0.06306	0.85555	0.04997	0.57374	-0.08126	-0.21575	0.04416	0.0339	1.19072	0.57884	-0.97567	0.04406	0.4886	-0.21672	-0.57953	-0.79923	0.8703	0.98342	0.3212	0.14965	0.70752	0.0146	0	
22	26.43	-2.07429	-0.12148	1.32202	0.41001	0.2952	-0.95954	0.54399	-0.10463	0.47566	0.14945	-0.85657	-0.18052	-0.65523	-0.2798	-0.21167	-0.33332	0.01075	-0.48847	0.50575	-0.38669	-0.40364	-0.2274	0.74243	0.39853	0.24921	0.2744	0.35997	0.24323	0	
23	41.88	1.17328	0.3535	0.28391	1.13356	-0.17258	-0.91605	0.36902	-0.32726	-0.24665	-0.04614	-0.14342	0.97935	1.49229	0.10142	0.76148	-0.01458	-0.51164	-0.32506	-0.39093	0.02788	0.067	0.22781	-0.15049	0.43505	0.72482	-0.33708	0.01637	0.03004	0	
23	16	1.32271	-0.17404	0.43456	0.57604	-0.83676	-0.83108	-0.2649	-0.22098	-1.07142	0.86856	-0.64151	-0.11132	0.36149	0.17195	0.78217	-1.35587	-0.21694	1.27177	-1.24062	-0.52295	-0.28438	-0.32336	-0.03771	0.34715	0.55964	-0.28016	0.04234	0.02882	0	
23	33	-0.41429	0.90544	1.72745	1.47347	0.00744	-0.20033	0.74023	-0.02925	-0.59339	-0.34619	-0.01214	0.7868	0.63595	-0.08632	0.0768	-1.40592	0.77559	-0.94289	0.54397	0.09731	0.07724	0.45733	-0.0385	0.64252	-0.18389	-0.27746	0.18269	0.15266	0	
23	12.99	1.05939	-0.17532	1.26613	1.18611	-0.786	0.57844	-0.76708	0.40105	0.6995	-0.06474	1.04820	1.00562	-0.542	-0.03001	-0.21868	0.00448	-0.10355	0.04230	-0.27783	-0.17807	0.01368	0.21373	0.01446	0.00905	0.20464	-0.30507	0.08146	0.02472	0	
24	17.28	1.23743	0.06104	0.38053	0.76156	-0.35977	-0.49408	0.00649	-0.13386	0.43881	-0.20736	-0.5																		0	
25	4.45	1.11401	0.08555	0.4937	1.33576	-0.30019	-0.01075	-0.11876	0.18862	0.20569	0.08226	1.1																		0	
26	6.14	-0.52991	0.87389	1.34725	0.14546	0.41421	0.10022	0.71121	0.17607	-0.28677	-0.48469	0.8																		0	

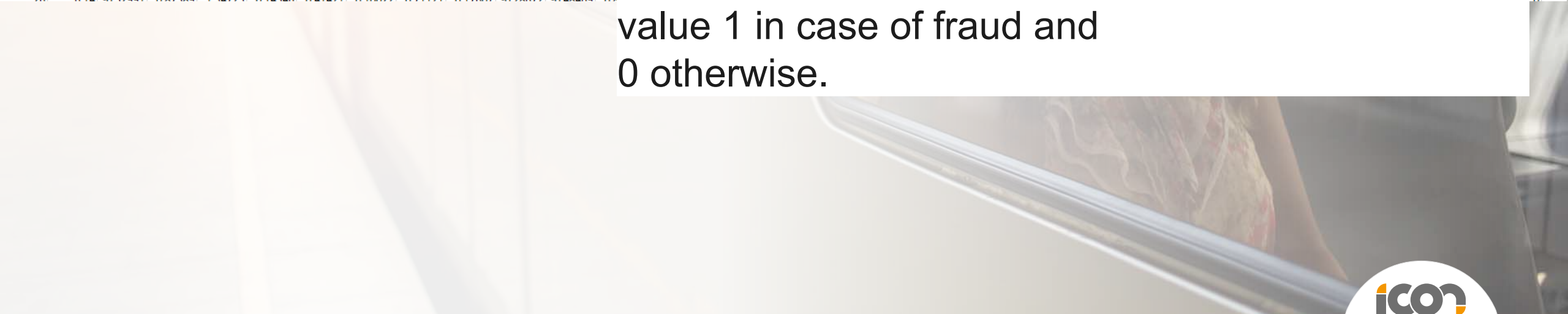
Features V1, V2, ... V28 are the principal components obtained with PCA

The feature 'Amount' is the transaction Amount

Feature 'Time' contains the seconds elapsed between each transaction and the first transaction in the dataset.

The only features which have not been transformed with PCA are 'Time' and 'Amount'

Feature 'Class' is the response variable and it takes value 1 in case of fraud and 0 otherwise.



22

Features V1, V2, ... V28 are the principal components obtained with PCA

The feature 'Amount' is the transaction Amount

Feature 'Time' contains the seconds elapsed between each transaction and the first transaction in the dataset.

The only features which have not been transformed with PCA are 'Time' and 'Amount'

Feature 'Class' is the response variable and it takes value 1 in case of fraud and 0 otherwise.

Demonstration 1

People with no idea about AI
saying it will take over the world:

My Neural Network:



Balance data

```
df = df.sample(frac=1)

# amount of fraud classes 492 rows.
fraud_df = df.loc[df['Class'] == 1]
non_fraud_df = df.loc[df['Class'] == 0][:492]

normal_distributed_df = pd.concat([fraud_df, non_fraud_df])

# Shuffle dataframe rows
new_df = normal_distributed_df.sample(frac=1, random_state=42)
```

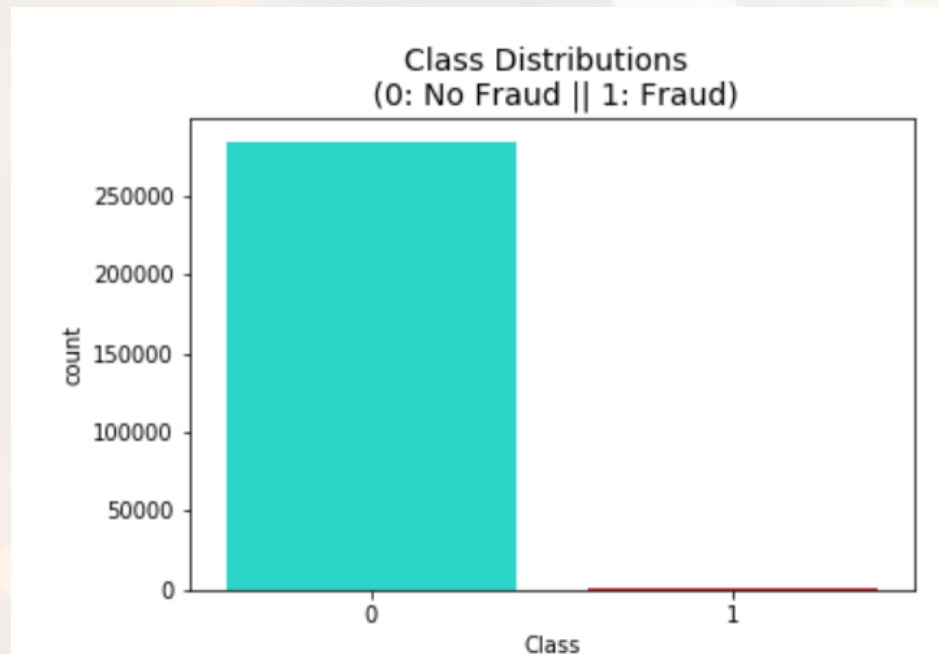
49%

```
# new X_test contains all the transaction not in X_train.
new_X_test = new_df.loc[~new_df.index.isin(new_X_train.index)]
new_train_data = new_X_train.iloc[:, 0:29]
new_train_labels = new_X_train.iloc[:, 30]
new_model = tf.keras.models.Sequential([
    tf.keras.layers.Flatten(input_shape=(29,)),
    tf.keras.layers.Dense(128, activation='relu'),
    tf.keras.layers.Dropout(0.2),
    tf.keras.layers.Dense(2, activation='softmax')
])
# %% compile model
new_model.compile(optimizer='adam',
                  loss='sparse_categorical_crossentropy',
                  metrics=['accuracy'])
new_model.fit(new_train_data, new_train_labels, epochs=5)
```

```
Epoch 1/5
787/787 [=====] - 0s 155us/sample - loss: 3847.0455 - accuracy: 0.4968
Epoch 2/5
787/787 [=====] - 0s 45us/sample - loss: 2625.2003 - accuracy: 0.4943
Epoch 3/5
787/787 [=====] - 0s 46us/sample - loss: 2379.6559 - accuracy: 0.4956
Epoch 4/5
787/787 [=====] - 0s 44us/sample - loss: 2528.0340 - accuracy: 0.4460
Epoch 5/5
787/787 [=====] - 0s 47us/sample - loss: 2153.2511 - accuracy: 0.4930
```

Data Loss

284315 -> 492



THIS IS YOUR MACHINE LEARNING SYSTEM?

YUP! YOU POUR THE DATA INTO THIS BIG PILE OF LINEAR ALGEBRA, THEN COLLECT THE ANSWERS ON THE OTHER SIDE.

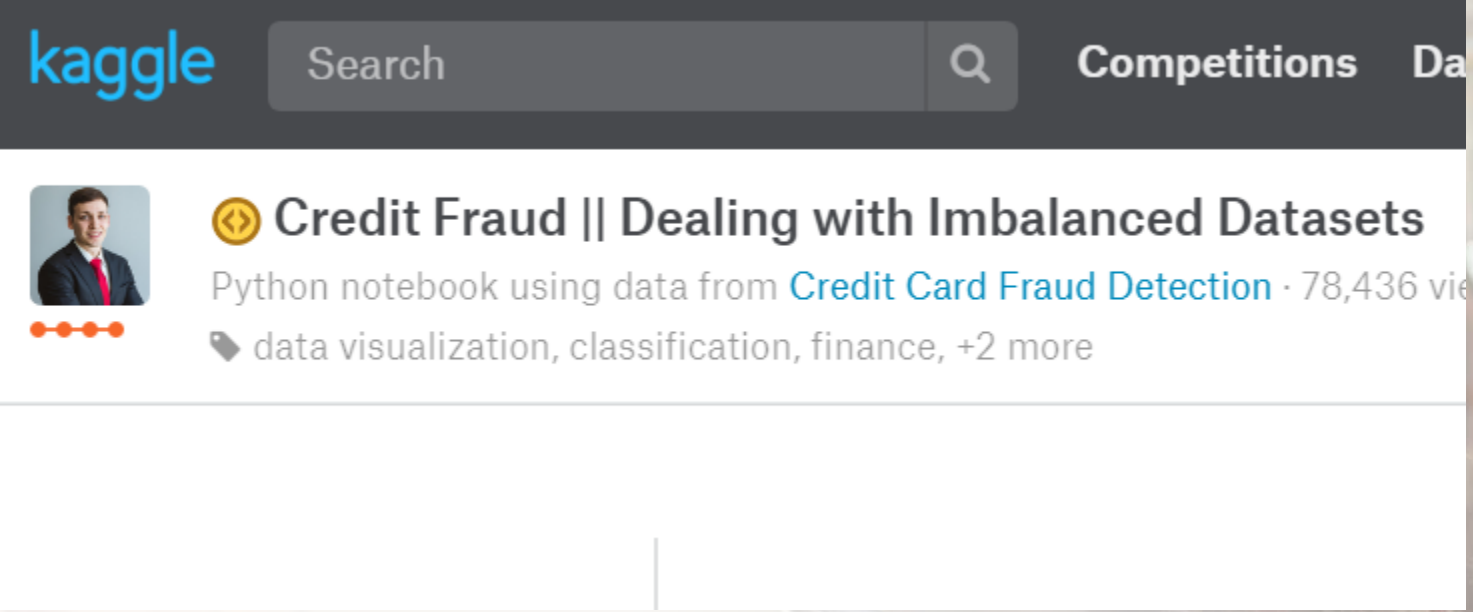
WHAT IF THE ANSWERS ARE WRONG?

JUST STIR THE PILE UNTIL THEY START LOOKING RIGHT.



Attempt 3

Janio Martinez Bachmann



The image shows a screenshot of the Kaggle website's search results. At the top is the Kaggle logo in blue, followed by a search bar containing the word 'Search' and a magnifying glass icon. To the right of the search bar are the links 'Competitions' and 'Data'. Below the search bar, a result is displayed for a notebook titled 'Credit Fraud || Dealing with Imbalanced Datasets'. To the left of the title is a profile picture of a man in a suit and a red tie, with three orange dots below it. The title itself is preceded by a yellow icon of a double-headed arrow. Below the title, the text reads 'Python notebook using data from [Credit Card Fraud Detection](#) · 78,436 views'. At the bottom of the result, there is a tag icon followed by the text 'data visualization, classification, finance, +2 more'.

kaggle Search Competitions Data


●●●

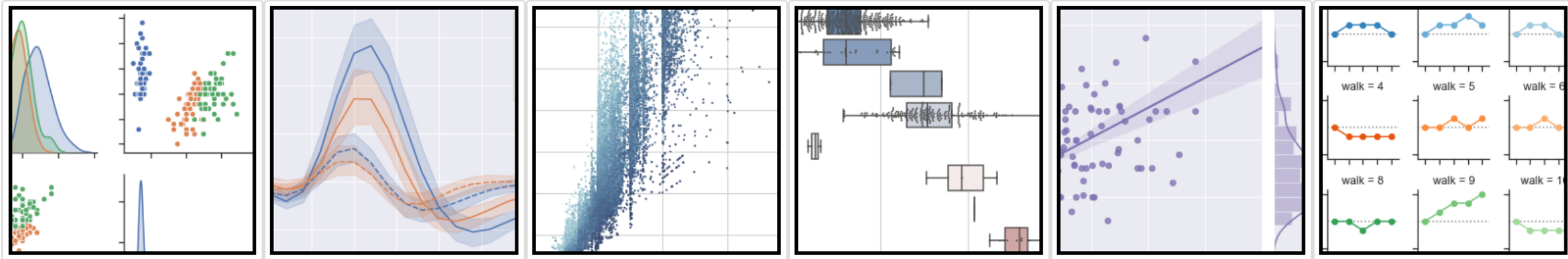
 **Credit Fraud || Dealing with Imbalanced Datasets**

Python notebook using data from [Credit Card Fraud Detection](#) · 78,436 views

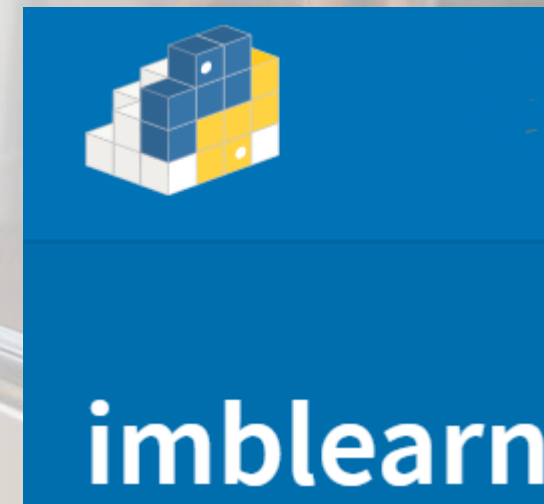
🏷 data visualization, classification, finance, +2 more

Libraries

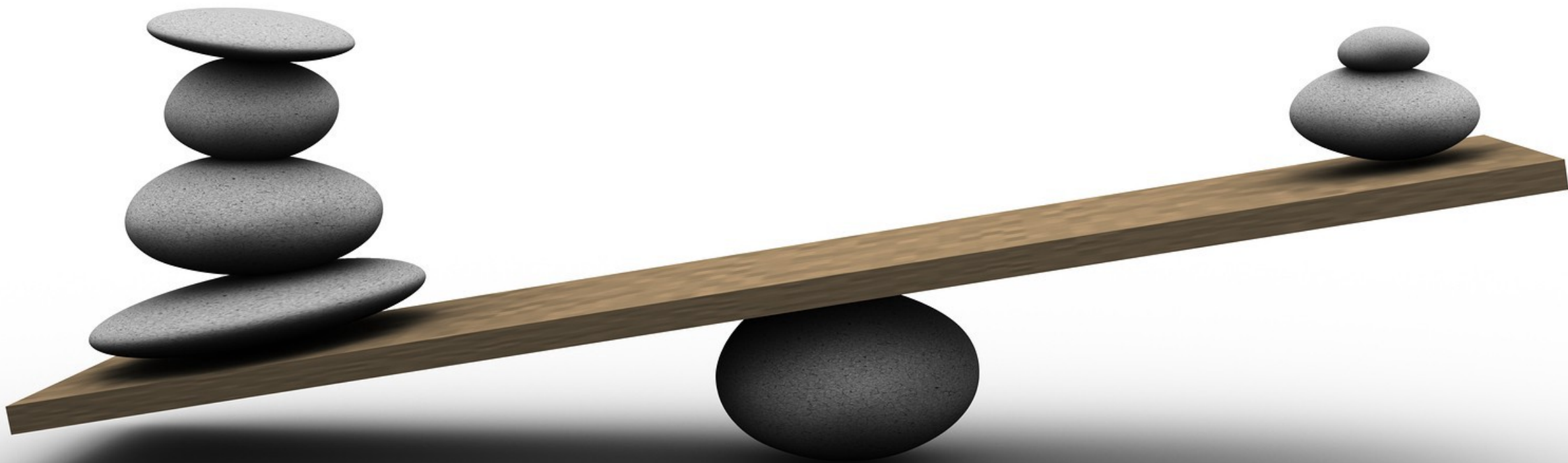
seaborn: statistical data visualization



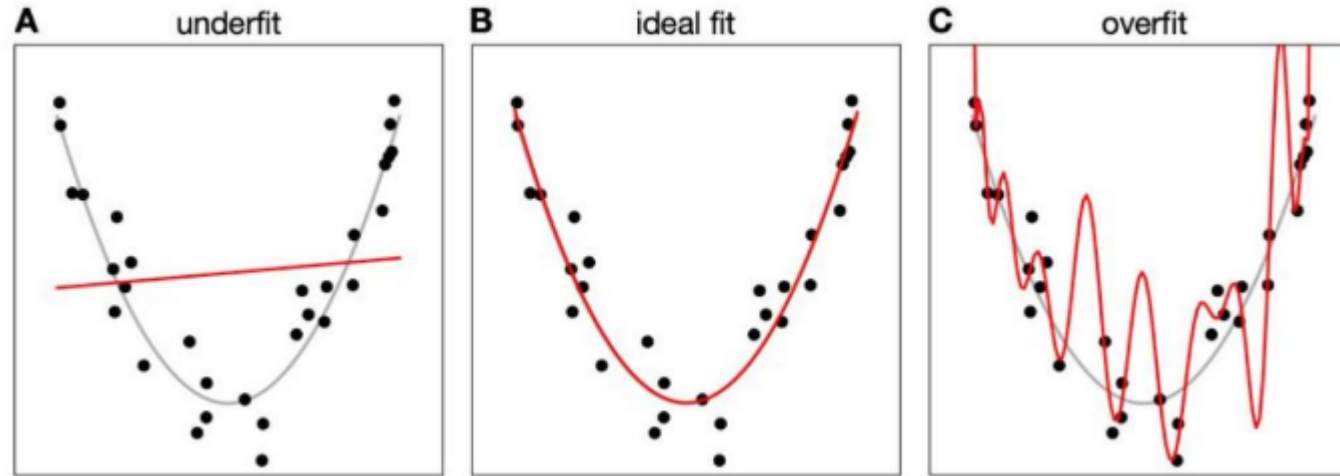
a library for making statistical graphics in Python.



Toolbox for imbalanced dataset in machine learning.



Underfitting and Overfitting

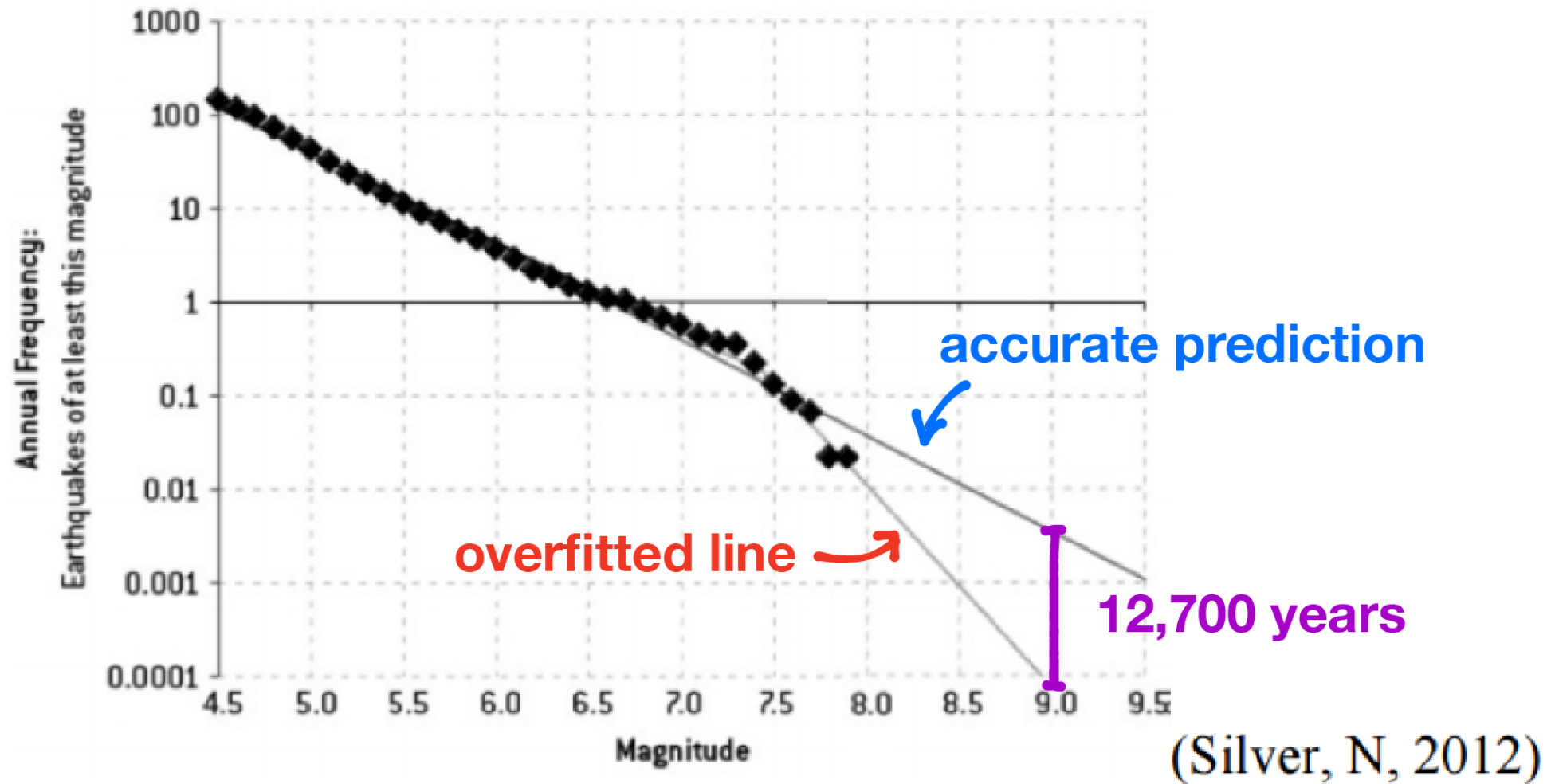


Complexity

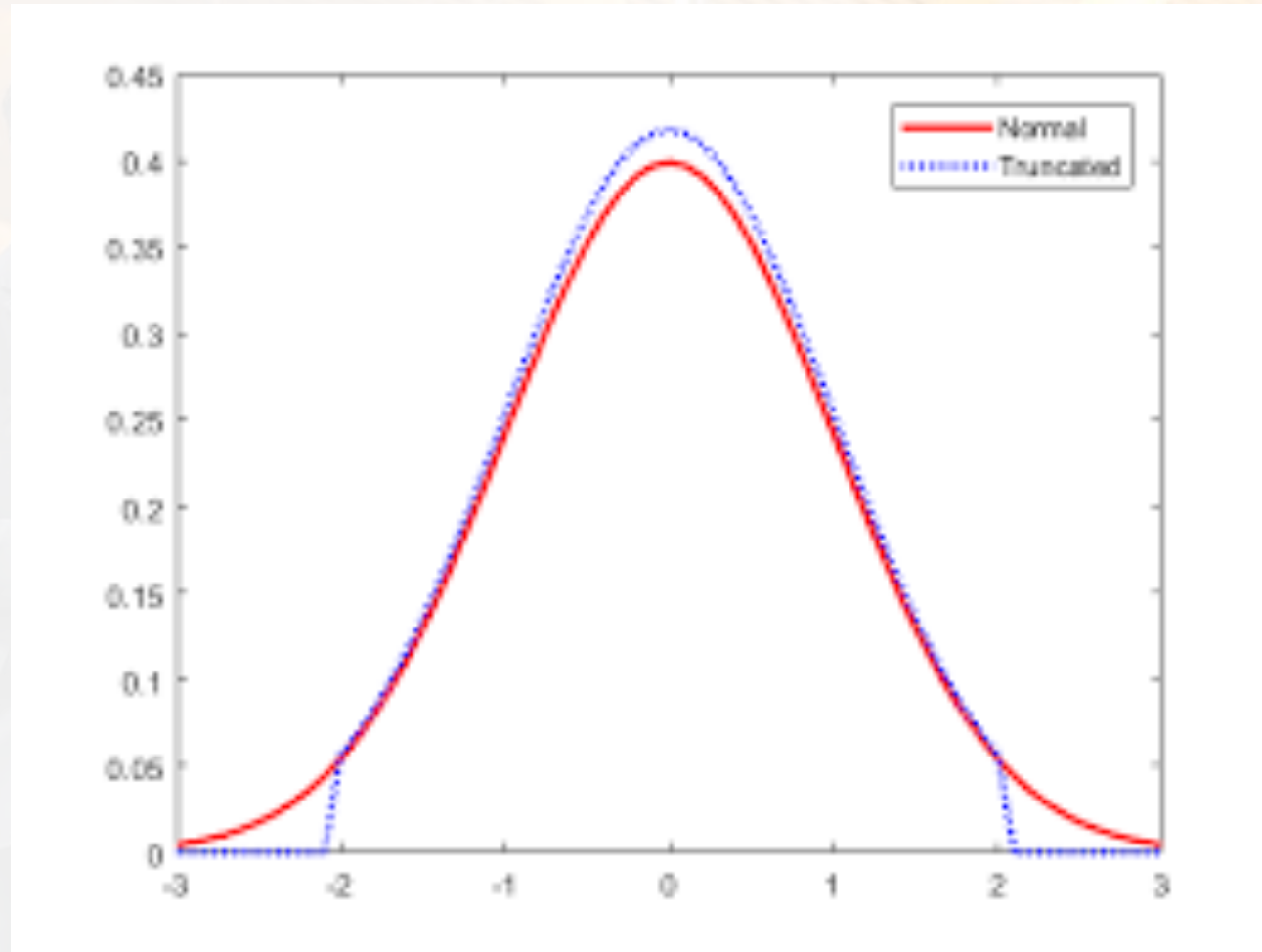


Overfitting

FIGURE 5-7C: TŌHOKU, JAPAN EARTHQUAKE FREQUENCIES
CHARACTERISTIC FIT

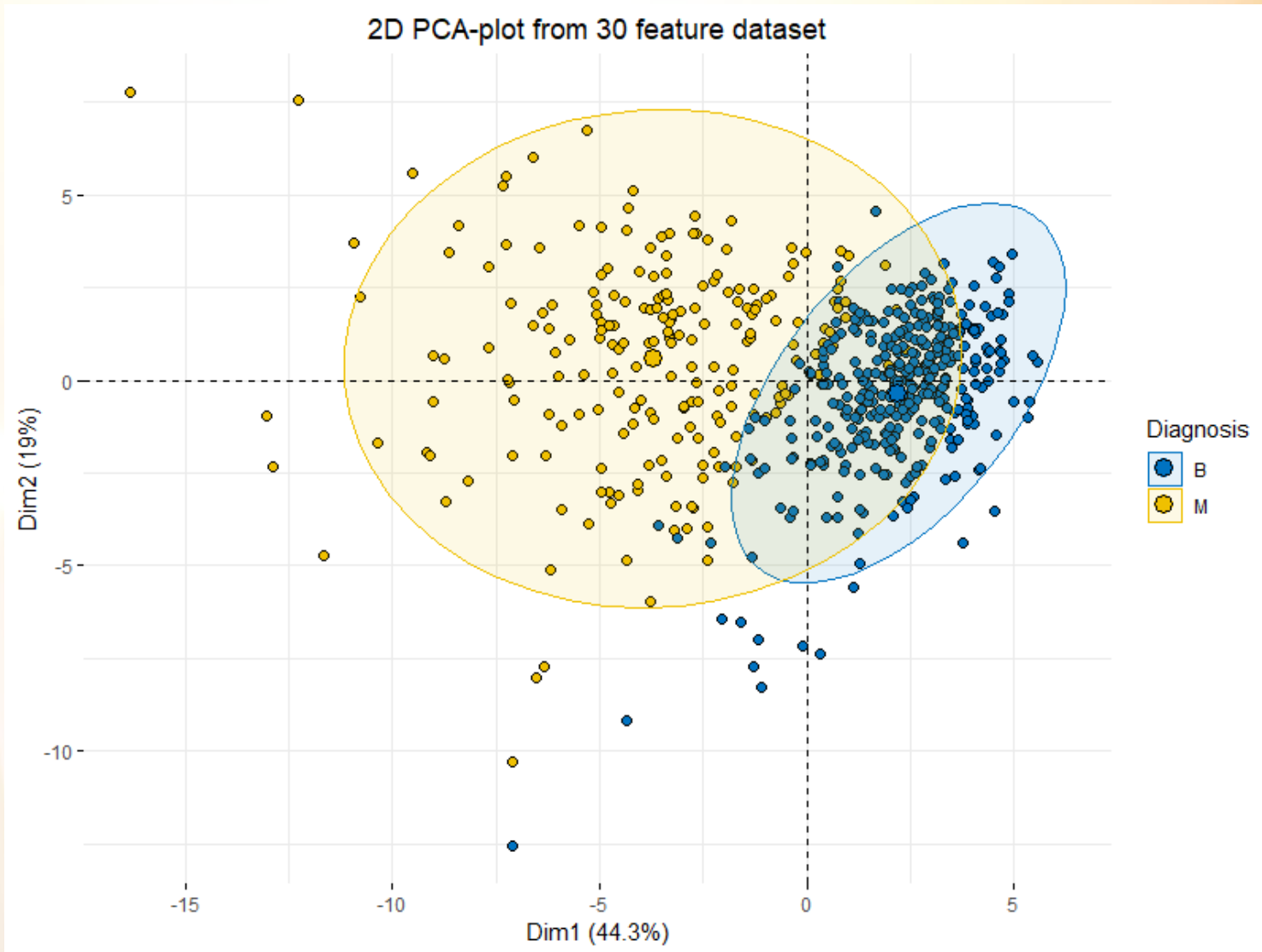


Outliers



object	temperature
object 1	71 °F
object 2	70 °F
object 3	73 °F
object 4	70 °F
object 5	70 °F
object 6	69 °F
object 7	70 °F
object 8	72 °F
object 9	71 °F
object 10	300 °F - outlier
object 11	71 °F
object 12	69 °F

Principal Component Analysis



Demonstration 2

```
df = pd.read_csv('creditcard.csv')
df.head()
```

Out[1]:

	Time	V1	V2	V3	V4	V5	V6	V7	V8	V9	...	V21	V22	V23	V24	
0	0.0	-1.359807	-0.072781	2.536347	1.378155	-0.338321	0.462388	0.239599	0.098698	0.363787	...	-0.018307	0.277838	-0.110474	0.066928	0.128
1	0.0	1.191857	0.266151	0.166480	0.448154	0.060018	-0.082361	-0.078803	0.085102	-0.255425	...	-0.225775	-0.638672	0.101288	-0.339846	0.167
2	1.0	-1.358354	-1.340163	1.773209	0.379780	-0.503198	1.800499	0.791461	0.247676	-1.514654	...	0.247998	0.771679	0.909412	-0.689281	-0.327
3	1.0	-0.966272	-0.185226	1.792993	-0.863291	-0.010309	1.247203	0.237609	0.377436	-1.387024	...	-0.108300	0.005274	-0.190321	-1.175575	0.647
4	2.0	-1.158233	0.877737	1.548718	0.403034	-0.407193	0.095921	0.592941	-0.270533	0.817739	...	-0.009431	0.798278	-0.137458	0.141267	-0.206

Scale time and amount

Scaling

Time and **Amount** should be scaled as the other columns.

In [2]: *# Since most of our data has already been scaled we should scale Amount and Time to be between -1 and 1 too*

```
from sklearn.preprocessing import RobustScaler

rob_scaler = RobustScaler()

df['scaled_amount'] = rob_scaler.fit_transform(df['Amount'].values.reshape(-1,1))
df['scaled_time'] = rob_scaler.fit_transform(df['Time'].values.reshape(-1,1))

df.drop(['Time','Amount'], axis=1, inplace=True)
scaled_amount = df['scaled_amount']
scaled_time = df['scaled_time']

df.drop(['scaled_amount', 'scaled_time'], axis=1, inplace=True)
df.insert(0, 'scaled_amount', scaled_amount)
df.insert(1, 'scaled_time', scaled_time)

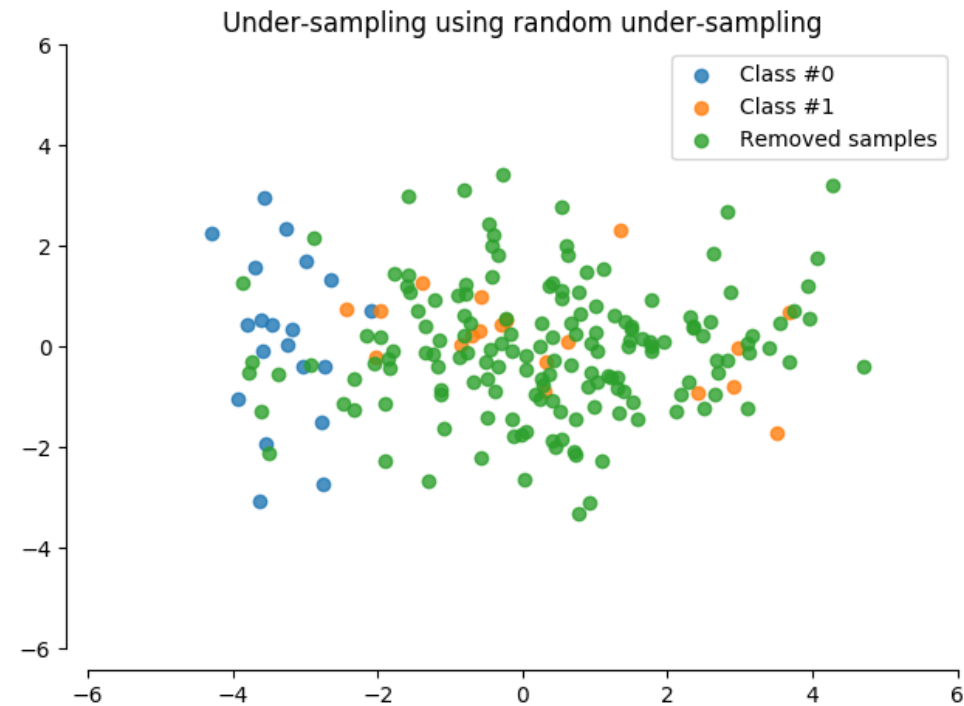
# Amount and Time are Scaled!

df.head()
```

Out[2]:

	scaled_amount	scaled_time	V1	V2	V3	V4	V5	V6	V7	V8	...	V20	V21	V22	
0	1.783274	-0.994983	-1.359807	-0.072781	2.536347	1.378155	-0.338321	0.462388	0.239599	0.098698	...	0.251412	-0.018307	0.277838	-0.11
1	-0.269825	-0.994983	1.191857	0.266151	0.166480	0.448154	0.060018	-0.082361	-0.078803	0.085102	...	-0.069083	-0.225775	-0.638672	0.10

Random under-sampling



In this phase of the project we will implement "Random Under Sampling" which basically consists of removing data in order to have a more **balanced dataset** and thus avoiding our models overfitting.

```
df = df.sample(frac=1)

# amount of fraud classes 492 rows.
fraud_df = df.loc[df['Class'] == 1]
non_fraud_df = df.loc[df['Class'] == 0][:492]

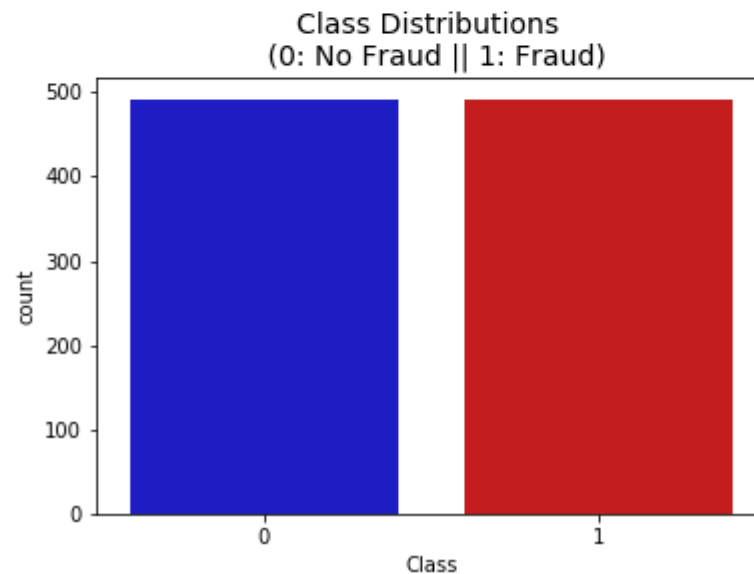
normal_distributed_df = pd.concat([fraud_df, non_fraud_df])

# Shuffle dataframe rows
new_df = normal_distributed_df.sample(frac=1, random_state=42)

colors = ["#0101DF", "#DF0101"]

sns.countplot('Class', data=new_df, palette=colors)
plt.title('Class Distributions \n (0: No Fraud || 1: Fraud)', fontsize=14)
```

Out[3]: Text(0.5, 1.0, 'Class Distributions \n (0: No Fraud || 1: Fraud)')



Correlation Matrix

Used to show which features heavily influence whether a transaction is a fraud

Now that we have our dataframe correctly balanced, we can go further with our **analysis** and **data preprocessing**.

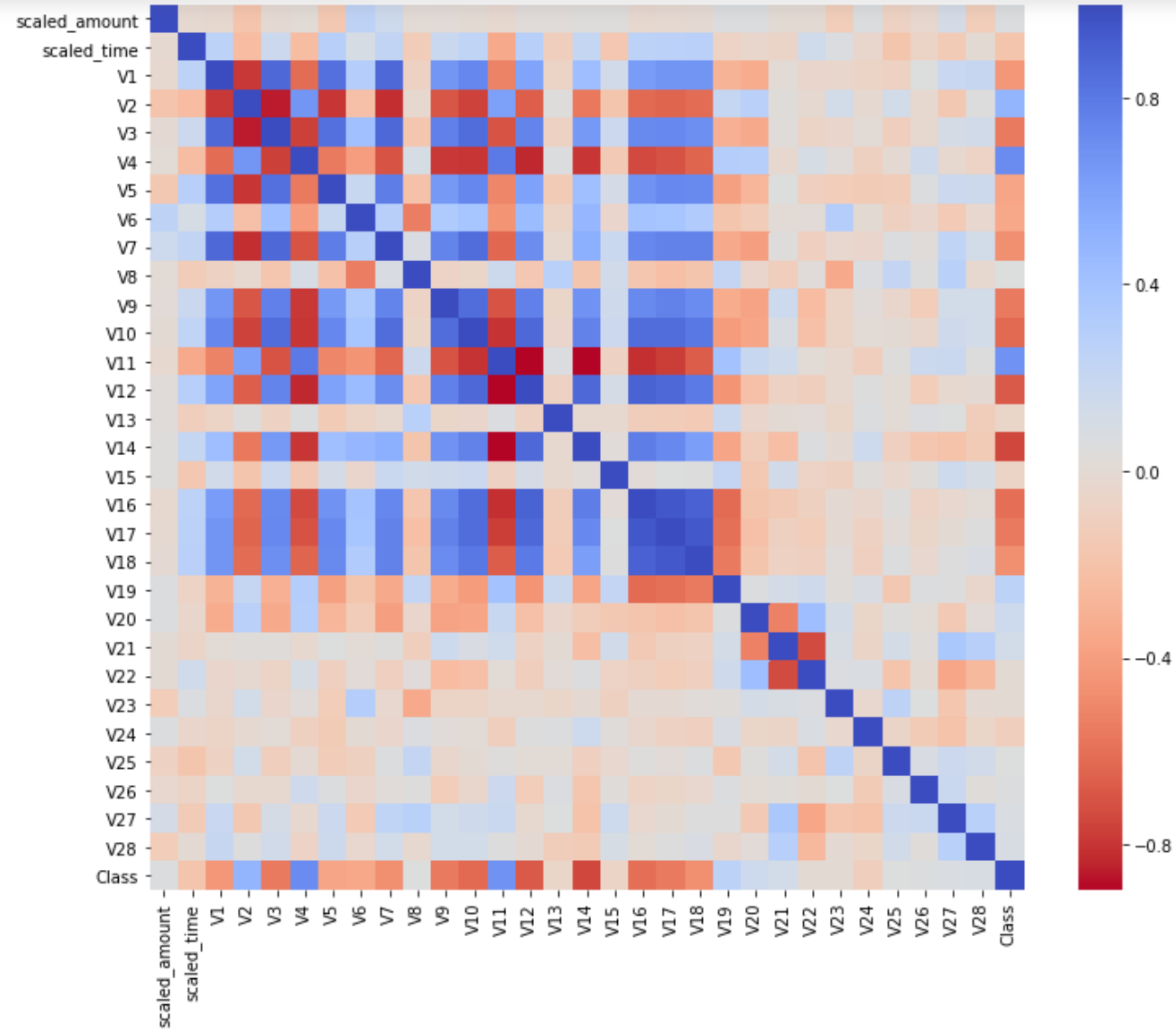
Correlation Matrices

A correlation matrix is used to show which features heavily influence whether a specific transaction is a fraud.

```
# Make sure we use the subsample in our correlation

f, (ax1) = plt.subplots(figsize=(12,10))

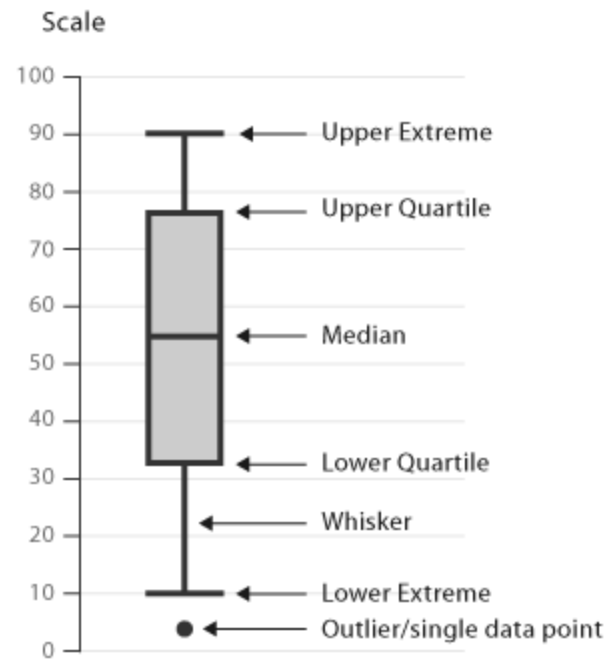
sub_sample_corr = new_df.corr()
sns.heatmap(sub_sample_corr, cmap='coolwarm_r', annot_kws={'size':20}, ax=ax1)
ax1.set_title('SubSample Correlation Matrix \n', fontsize=14)
plt.show()
```



Summary and Explanation:

- **Positive Correlations:** V2, V4, V11, and V19 are positively correlated. The higher these values are, the more likely the end result will be a fraud transaction.
- **Negative Correlations:** V17, V14, V12 and V10 are negatively correlated. The lower these values are, the more likely the end result will be a fraud transaction.

Anomaly detection



Our main aim in this section is to remove "extreme outliers" from features that have a high correlation with our classes. This will have a positive impact on the accuracy of our models.

```
f, axes = plt.subplots(ncols=4, figsize=(20,4))

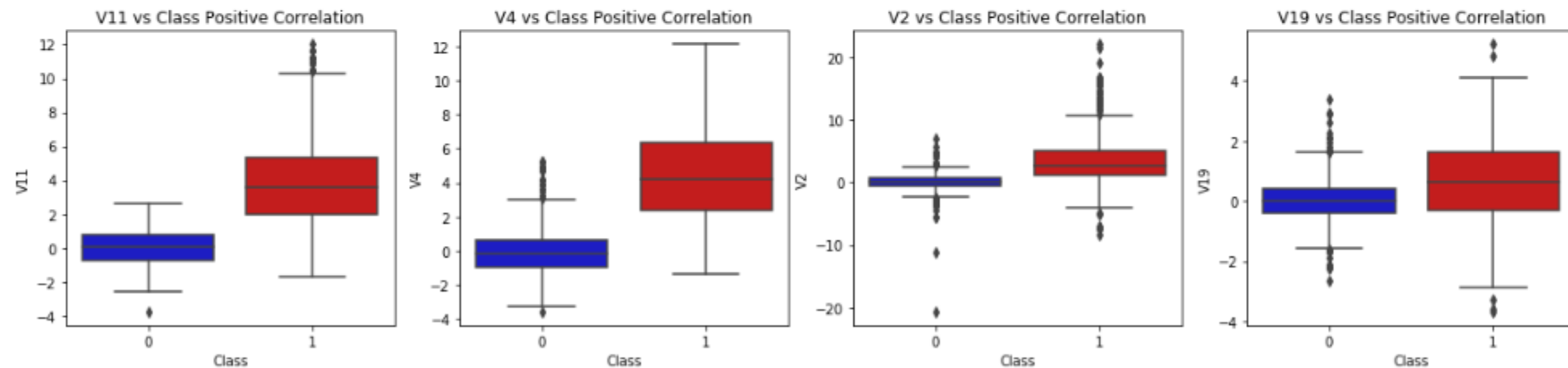
# Positive correlations (The higher the feature the probability increases that it will be a fraud transaction)
sns.boxplot(x="Class", y="V11", data=new_df, palette=colors, ax=axes[0])
axes[0].set_title('V11 vs Class Positive Correlation')

sns.boxplot(x="Class", y="V4", data=new_df, palette=colors, ax=axes[1])
axes[1].set_title('V4 vs Class Positive Correlation')

sns.boxplot(x="Class", y="V2", data=new_df, palette=colors, ax=axes[2])
axes[2].set_title('V2 vs Class Positive Correlation')

sns.boxplot(x="Class", y="V19", data=new_df, palette=colors, ax=axes[3])
axes[3].set_title('V19 vs Class Positive Correlation')

plt.show()
```



```
# Removing outliers V10 Feature
v10_fraud = new_df['V10'].loc[new_df['Class'] == 1].values
q25, q75 = np.percentile(v10_fraud, 25), np.percentile(v10_fraud, 75)
v10_iqr = q75 - q25

v10_cut_off = v10_iqr * 1.5
v10_lower, v10_upper = q25 - v10_cut_off, q75 + v10_cut_off
print('V10 Lower: {}'.format(v10_lower))
print('V10 Upper: {}'.format(v10_upper))
outliers = [x for x in v10_fraud if x < v10_lower or x > v10_upper]
print('V10 outliers: {}'.format(outliers))
print('Feature V10 Outliers for Fraud Cases: {}'.format(len(outliers)))
new_df = new_df.drop(new_df[(new_df['V10'] > v10_upper) | (new_df['V10'] < v10_lower)].index)
print('Number of Instances after outliers removal: {}'.format(len(new_df)))
```

After implementing outlier reduction our accuracy has been improved by **over 3%**! Some outliers can distort the accuracy of our models but remember, we have to avoid an extreme amount of information loss or else our model runs the risk of underfitting.

```
V10 Lower: -14.89885463232024
V10 Upper: 4.920334958342141
V10 outliers: [-20.949191554361104, -15.346098846877501, -17.141513641289198, -14.9246547735487, -16.7460441053944, -15.2318333
653018, -18.9132433348732, -15.124162814494698, -15.2399619587112, -16.6496281595399, -15.563791338730098, -15.563791338730098,
-23.2282548357516, -19.836148851696, -15.2399619587112, -15.1237521803455, -22.1870885620007, -16.6011969664137, -18.2711681738
888, -16.2556117491401, -22.1870885620007, -24.5882624372475, -16.3035376590131, -24.403184969972802, -22.1870885620007, -14.92
46547735487, -22.1870885620007]
Feature V10 Outliers for Fraud Cases: 27
Number of Instances after outliers removal: 945
```

Dimensionality Reduction and Clustering

of the information.

```
# New_df is from the random undersample data (fewer instances)
X = new_df.drop('Class', axis=1)
y = new_df['Class']

# T-SNE Implementation
t0 = time.time()
X_reduced_tsne = TSNE(n_components=2, random_state=42).fit_transform(X.values)
t1 = time.time()

# PCA Implementation
t0 = time.time()
X_reduced_pca = PCA(n_components=2, random_state=42).fit_transform(X.values)
t1 = time.time()

f, (ax1, ax2) = plt.subplots(1, 2, figsize=(24,6))
# labels = ['No Fraud', 'Fraud']
f.suptitle('Clusters using Dimensionality Reduction', fontsize=14)

blue_patch = mpatches.Patch(color='#0A0AFF', label='No Fraud')
red_patch = mpatches.Patch(color='#AF0000', label='Fraud')

# t-SNE scatter plot
ax1.scatter(X_reduced_tsne[:,0], X_reduced_tsne[:,1], c=(y == 0), cmap='coolwarm', label='No Fraud', linewidths=2)
ax1.scatter(X_reduced_tsne[:,0], X_reduced_tsne[:,1], c=(y == 1), cmap='coolwarm', label='Fraud', linewidths=2)
ax1.set_title('t-SNE', fontsize=14)

ax1.grid(True)

ax1.legend(handles=[blue_patch, red_patch])

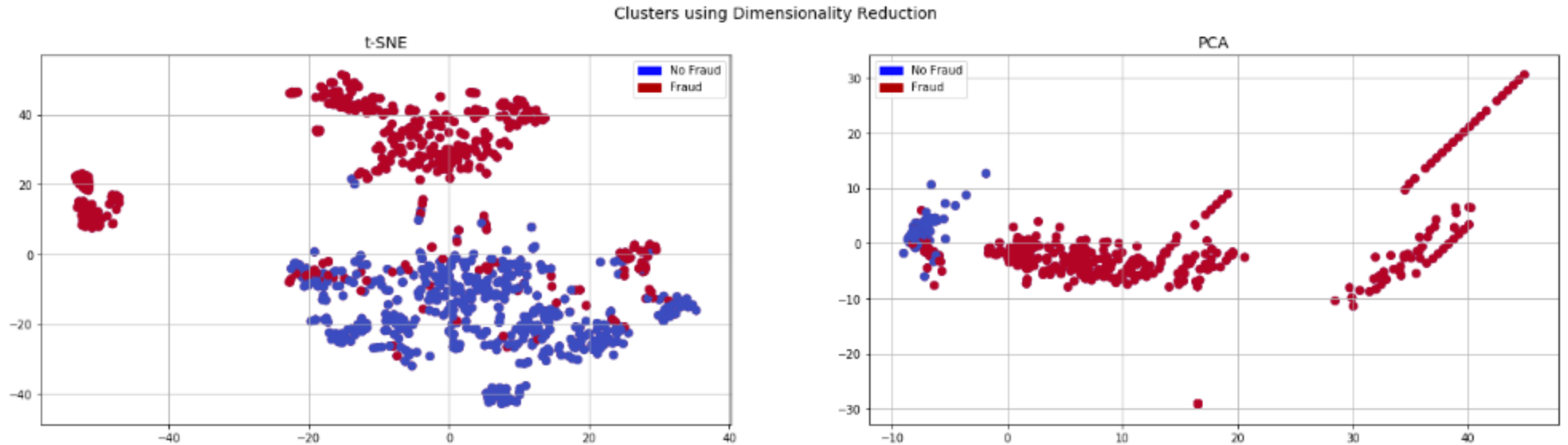
# PCA scatter plot
ax2.scatter(X_reduced_pca[:,0], X_reduced_pca[:,1], c=(y == 0), cmap='coolwarm', label='No Fraud', linewidths=2)
ax2.scatter(X_reduced_pca[:,0], X_reduced_pca[:,1], c=(y == 1), cmap='coolwarm', label='Fraud', linewidths=2)
ax2.set_title('PCA', fontsize=14)

ax2.grid(True)

ax2.legend(handles=[blue_patch, red_patch])

plt.show()
```

Dimensionality Reduction and Clustering

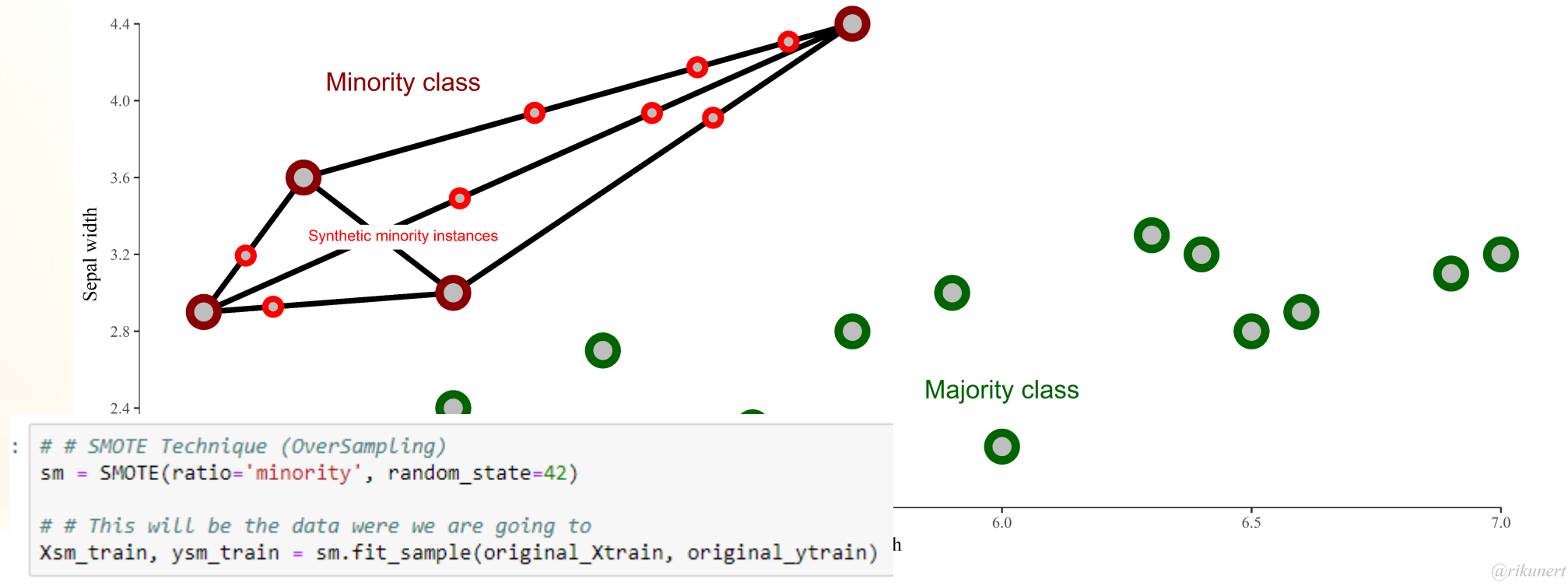


Summary:

- t-SNE algorithm can pretty accurately cluster the cases that were fraud and non-fraud in our dataset.
- Although the subsample is pretty small, the t-SNE algorithm is able to detect clusters pretty accurately in every scenario (I shuffle the dataset before running t-SNE)
- This gives us an indication that further predictive models will perform pretty well in separating fraud cases from non-fraud cases.

t-SNE takes a high-dimensional dataset and reduces it to a low-dimensional graph whilst still retaining a lot of the information.

SMOTE



Synthetic Minority Over-sampling Technique

Solving the Class Imbalance: SMOTE creates synthetic points from the minority class in order to reach an equal balance between the minority and majority class.

Location of the synthetic points: SMOTE picks the distance between the closest neighbors of the minority class, in between these distances it creates synthetic points.

Final Effect: More information is retained since we didn't have to delete any rows unlike in random undersampling.

Compile the model

The following example uses *accuracy*, the fraction of the transactions that are correctly classified.

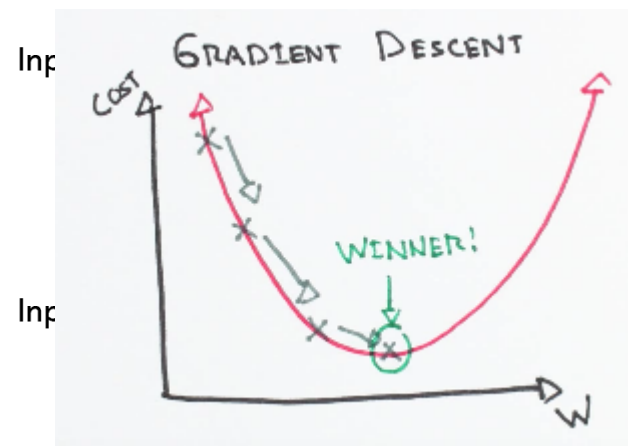
optimizers shape and mold your model into its most accurate possible form by futzing with the weights. The **loss** function is the guide to the terrain, telling the optimizer when it's moving in the right or wrong direction.

```
# import keras
from tensorflow.keras import backend as K
from tensorflow.keras.models import Sequential
from tensorflow.keras.layers import Activation
from tensorflow.keras.layers import Dense
from tensorflow.keras.optimizers import Adam
from tensorflow.keras.metrics import categorical_crossentropy
```

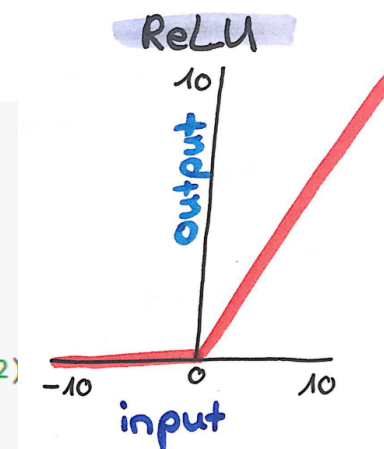
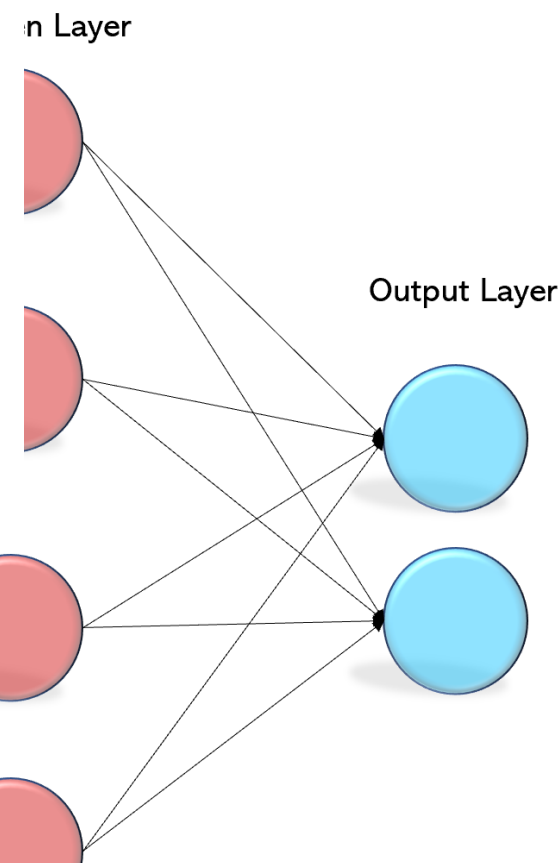
```
n_inputs = X_train.shape[1]
```

```
undersample_model = Sequential([
    Dense(n_inputs, input_shape=(n_inputs, ), activation='relu'),
    Model: "sequential"
```

Layer (type)	Output Shape	Param #
dense (Dense)	(None, 30)	930
dense_1 (Dense)	(None, 32)	992
dense_2 (Dense)	(None, 2)	66
Total params: 1,988		
Trainable params: 1,988		
undersample_fraud_predictions = undersample_model.predict_classes(original_Xtest, batch_size=200, verbose=0)		



```
ssentropy', metrics=['accuracy'])
size=25, epochs=20, shuffle=True, verbose=2)
atch_size=200, verbose=0)
```



Epoch 1/20
604/604 - 0s - loss: 0.7081 - accuracy: 0.6358 - val_loss: 0.4522 - val_accuracy: 0.7947
Epoch 2/20
604/604 - 0s - loss: 0.3993 - accuracy: 0.8046 - val_loss: 0.3239 - val_accuracy: 0.8675
Epoch 3/20
604/604 - 0s - loss: 0.3101 - accuracy: 0.8825 - val_loss: 0.2623 - val_accuracy: 0.9073
Epoch 4/20
604/604 - 0s - loss: 0.2573 - accuracy: 0.9255 - val_loss: 0.2272 - val_accuracy: 0.9139
Epoch 5/20
604/604 - 0s - loss: 0.2201 - accuracy: 0.9387 - val_loss: 0.2035 - val_accuracy: 0.9139
Epoch 6/20
604/604 - 0s - loss: 0.1927 - accuracy: 0.9437 - val_loss: 0.1856 - val_accuracy: 0.9139
Epoch 7/20
604/604 - 0s - loss: 0.1707 - accuracy: 0.9487 - val_loss: 0.1728 - val_accuracy: 0.9205
Epoch 8/20
604/604 - 0s - loss: 0.1549 - accuracy: 0.9503 - val_loss: 0.1638 - val_accuracy: 0.9205
Epoch 9/20
604/604 - 0s - loss: 0.1437 - accuracy: 0.9503 - val_loss: 0.1615 - val_accuracy: 0.9205
Epoch 10/20
604/604 - 0s - loss: 0.1344 - accuracy: 0.9503 - val_loss: 0.1553 - val_accuracy: 0.9139
Epoch 11/20
604/604 - 0s - loss: 0.1258 - accuracy: 0.9553 - val_loss: 0.1522 - val_accuracy: 0.9139
Epoch 12/20
604/604 - 0s - loss: 0.1209 - accuracy: 0.9487 - val_loss: 0.1488 - val_accuracy: 0.9139
Epoch 13/20
604/604 - 0s - loss: 0.1135 - accuracy: 0.9570 - val_loss: 0.1470 - val_accuracy: 0.9139
Epoch 14/20
604/604 - 0s - loss: 0.1078 - accuracy: 0.9586 - val_loss: 0.1470 - val_accuracy: 0.9139
Epoch 15/20
604/604 - 0s - loss: 0.1023 - accuracy: 0.9636 - val_loss: 0.1463 - val_accuracy: 0.9139
Epoch 16/20
604/604 - 0s - loss: 0.0981 - accuracy: 0.9636 - val_loss: 0.1463 - val_accuracy: 0.9205
Epoch 17/20
604/604 - 0s - loss: 0.0934 - accuracy: 0.9685 - val_loss: 0.1489 - val_accuracy: 0.9205
Epoch 18/20
604/604 - 0s - loss: 0.0901 - accuracy: 0.9719 - val_loss: 0.1474 - val_accuracy: 0.9205
Epoch 19/20
604/604 - 0s - loss: 0.0854 - accuracy: 0.9685 - val_loss: 0.1511 - val_accuracy: 0.9205
Epoch 20/20
604/604 - 0s - loss: 0.0821 - accuracy: 0.9752 - val_loss: 0.1476 - val_accuracy: 0.9272

Confusion Matrix

	Predicted: no	Predicted: Yes
Actual: no	True negative	False positive
Actual: yes	False negative	True positive

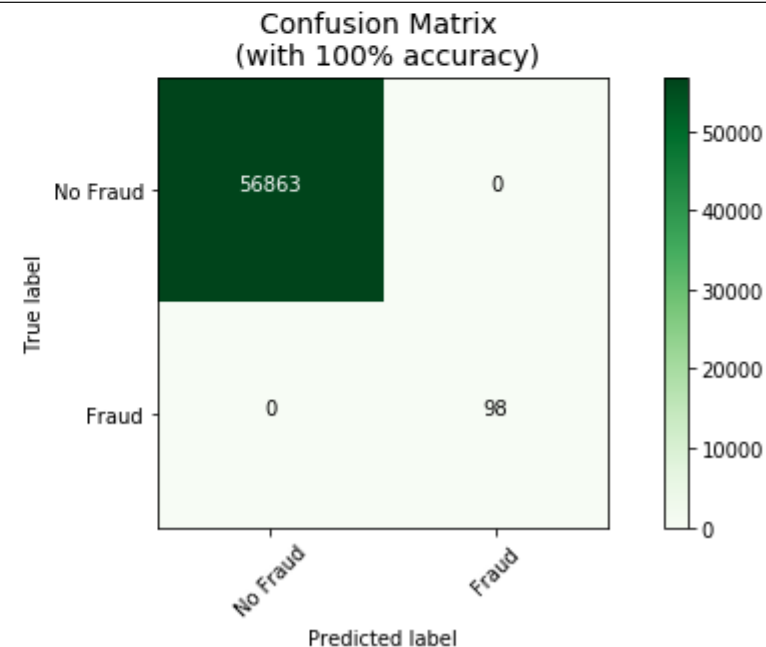
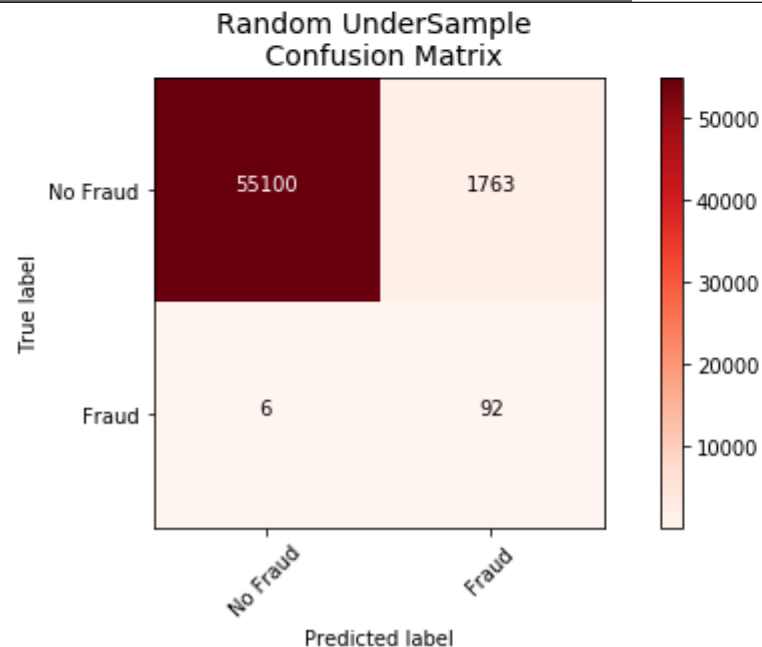
```
undersample_cm = confusion_matrix(original_ytest, undersample_fraud_predictions)
actual_cm = confusion_matrix(original_ytest, original_ytest)
labels = ['No Fraud', 'Fraud']

fig = plt.figure(figsize=(16,8))

fig.add_subplot(221)
plot_confusion_matrix(undersample_cm, labels, title="Random UnderSample \n Confusion Matrix", cmap=plt.cm.Red)

fig.add_subplot(222)
plot_confusion_matrix(actual_cm, labels, title="Confusion Matrix \n (with 100% accuracy)", cmap=plt.cm.Green)
```

	Predicted: no	Predicted: Yes
Actual: no	True negative	False positive
Actual: yes	False negative	True positive



Using SMOTE

Keras || OverSampling (SMOTE):

```
n_inputs = Xsm_train.shape[1]

oversample_model = Sequential([
    Dense(n_inputs, input_shape=(n_inputs, ), activation='relu'),
    Dense(32, activation='relu'),
    Dense(2, activation='softmax')
])

oversample_model.compile(Adam(lr=0.001), loss='sparse_categorical_crossentropy', metrics=['accuracy'])

oversample_model.fit(Xsm_train, ysm_train, validation_split=0.2, batch_size=300, epochs=20, shuffle=True, verbose=2)

oversample_predictions = oversample_model.predict(original_Xtest, batch_size=200, verbose=0)

oversample_fraud_predictions = oversample_model.predict_classes(original_Xtest, batch_size=200, verbose=0)
```

Train on 363923 samples, validate on 90981 samples

Epoch 1/20

363923/363923 - 2s - loss: 0.0898 - accuracy: 0.9685 - val_loss: 0.0341 - val_accuracy: 0.9880

Epoch 2/20

363923/363923 - 2s - loss: 0.0187 - accuracy: 0.9953 - val_loss: 0.0118 - val_accuracy: 0.9997

Epoch 3/20

363923/363923 - 2s - loss: 0.0099 - accuracy: 0.9978 - val_loss: 0.0078 - val_accuracy: 0.9996

Epoch 4/20

363923/363923 - 2s - loss: 0.0065 - accuracy: 0.9985 - val_loss: 0.0038 - val_accuracy: 0.9999

Epoch 5/20

363923/363923 - 2s - loss: 0.0050 - accuracy: 0.9990 - val_loss: 0.0055 - val_accuracy: 1.0000

Epoch 6/20

363923/363923 - 2s - loss: 0.0039 - accuracy: 0.9992 - val_loss: 0.0019 - val_accuracy: 1.0000

Epoch 7/20

363923/363923 - 2s - loss: 0.0035 - accuracy: 0.9993 - val_loss: 0.0031 - val_accuracy: 0.9997

Epoch 8/20

363923/363923 - 2s - loss: 0.0028 - accuracy: 0.9994 - val_loss: 0.0015 - val_accuracy: 1.0000

Epoch 9/20

363923/363923 - 2s - loss: 0.0028 - accuracy: 0.9994 - val_loss: 0.0018 - val_accuracy: 1.0000

Epoch 10/20

363923/363923 - 2s - loss: 0.0022 - accuracy: 0.9995 - val_loss: 0.0015 - val_accuracy: 1.0000

Epoch 11/20

363923/363923 - 2s - loss: 0.0020 - accuracy: 0.9996 - val_loss: 0.0011 - val_accuracy: 1.0000

Epoch 12/20

363923/363923 - 2s - loss: 0.0021 - accuracy: 0.9996 - val_loss: 0.0033 - val_accuracy: 0.9993

Epoch 13/20

363923/363923 - 2s - loss: 0.0017 - accuracy: 0.9997 - val_loss: 7.2115e-04 - val_accuracy: 1.0000

Epoch 14/20

363923/363923 - 2s - loss: 0.0015 - accuracy: 0.9997 - val_loss: 6.4973e-04 - val_accuracy: 1.0000

Epoch 15/20

363923/363923 - 2s - loss: 0.0014 - accuracy: 0.9997 - val_loss: 4.9727e-04 - val_accuracy: 1.0000

Epoch 16/20

363923/363923 - 2s - loss: 0.0014 - accuracy: 0.9997 - val_loss: 3.8057e-04 - val_accuracy: 1.0000

Epoch 17/20

363923/363923 - 2s - loss: 0.0013 - accuracy: 0.9997 - val_loss: 0.0018 - val_accuracy: 0.9998

Epoch 18/20

363923/363923 - 3s - loss: 0.0013 - accuracy: 0.9997 - val_loss: 0.0011 - val_accuracy: 1.0000

Epoch 19/20

363923/363923 - 2s - loss: 9.5826e-04 - accuracy: 0.9998 - val_loss: 1.4088e-04 - val_accuracy: 1.0000

Epoch 20/20

363923/363923 - 2s - loss: 8.8616e-04 - accuracy: 0.9998 - val_loss: 4.3943e-04 - val_accuracy: 1.0000

```
oversample_smote = confusion_matrix(original_cm,
actual_cm = confusion_matrix(original_cm,
labels = ['No Fraud', 'Fraud'])
```

```
fig = plt.figure(figsize=(16,8))
```

```
fig.add_subplot(221)
plot_confusion_matrix(oversample_smote,
```

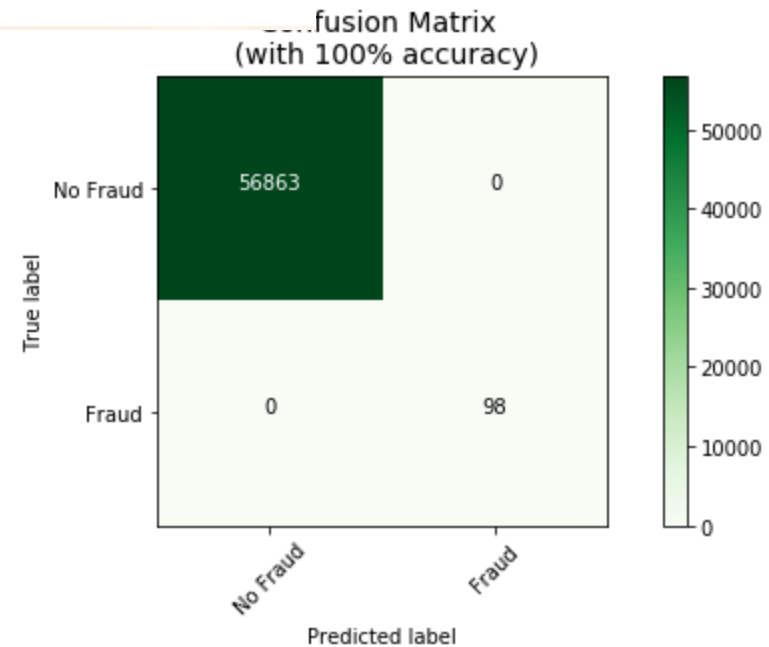
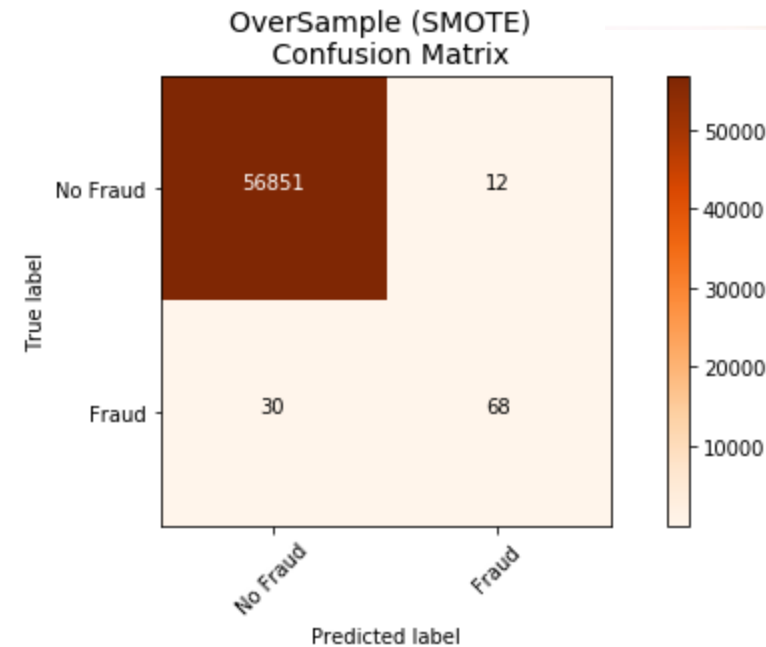
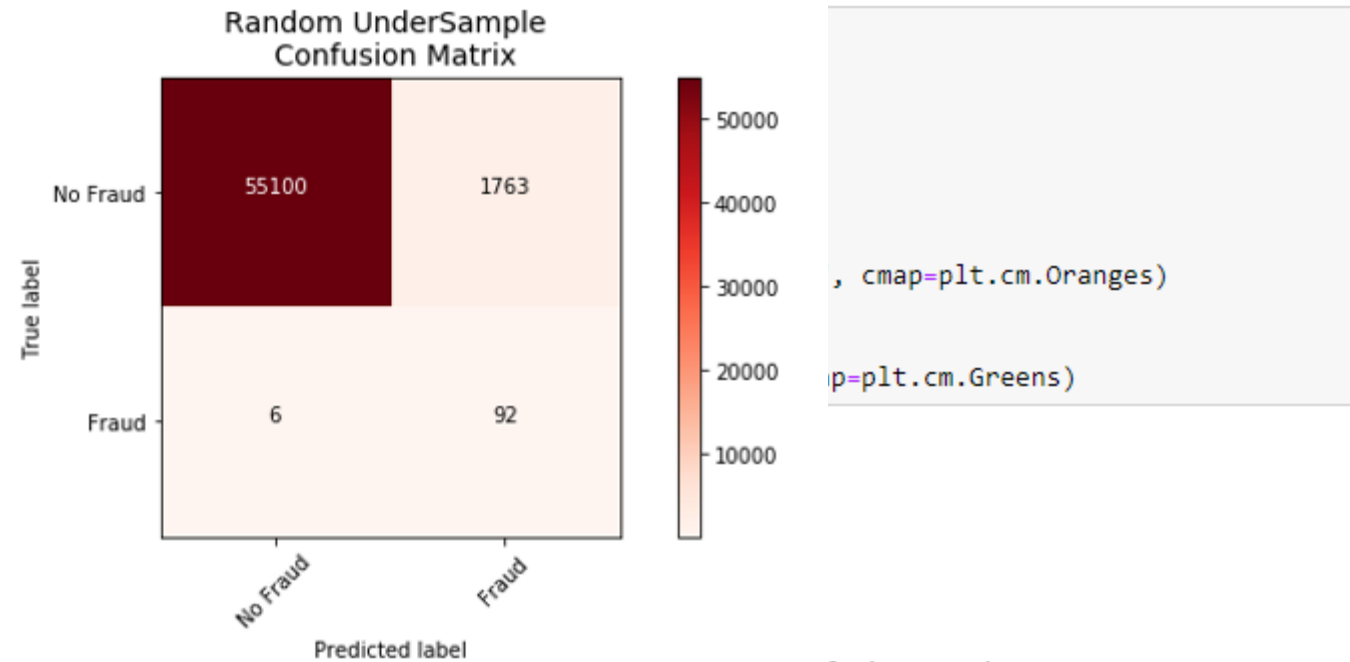
```
fig.add_subplot(222)
plot_confusion_matrix(actual_cm, labels=
```

Confusion matrix, without normalization

```
[[56851  12]
 [  30   68]]
```

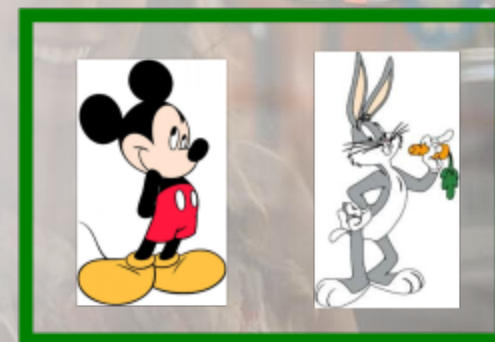
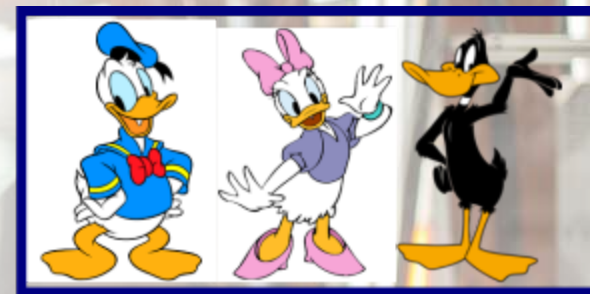
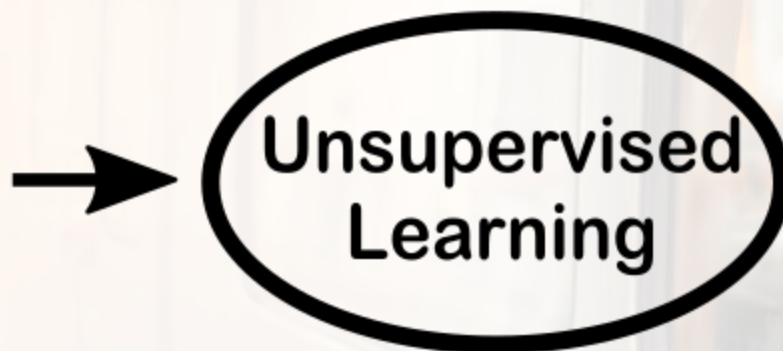
Confusion matrix, without normalization

```
[[56863   0]
 [   0   98]]
```





Unsupervised Learning



Iris Data Set



Iris setosa



Iris versicolor



Iris virginica



- 50 samples from each of three species of *Iris*.
 - Four features were measured from each sample:
 - the length and the width of the sepals and petals, in centimeters.
-
- the objective of K-means is simple:
 - group similar data points together and discover underlying patterns.
-
- To achieve this objective, K-means looks for a fixed number (k) of clusters in a dataset.

KMeans

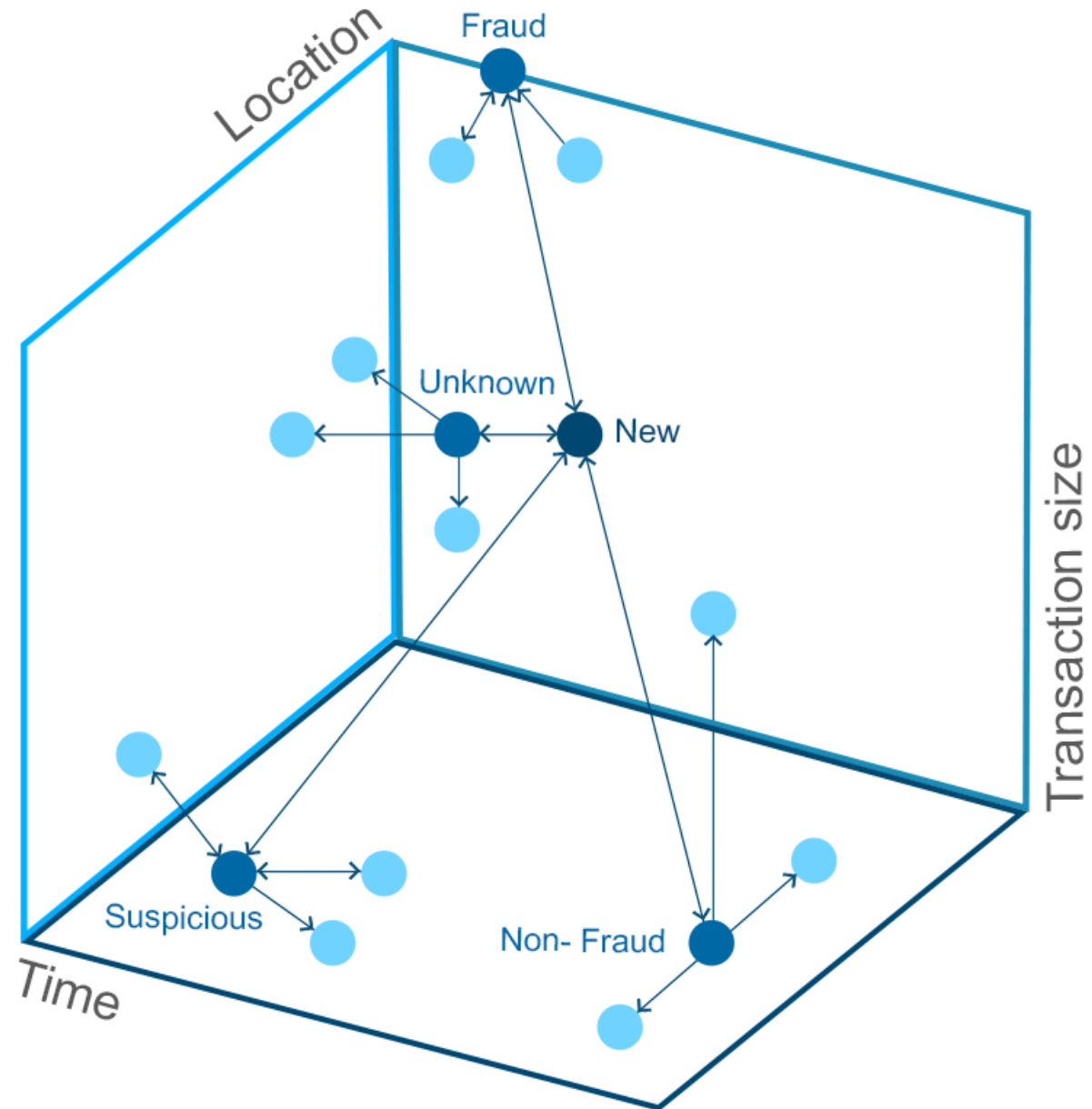
K-means clustering is a type of unsupervised learning, which is used when you have unlabeled data (i.e., data without defined categories or groups).

The goal of this algorithm is to find groups in the data, with the number of groups represented by the variable *K*.

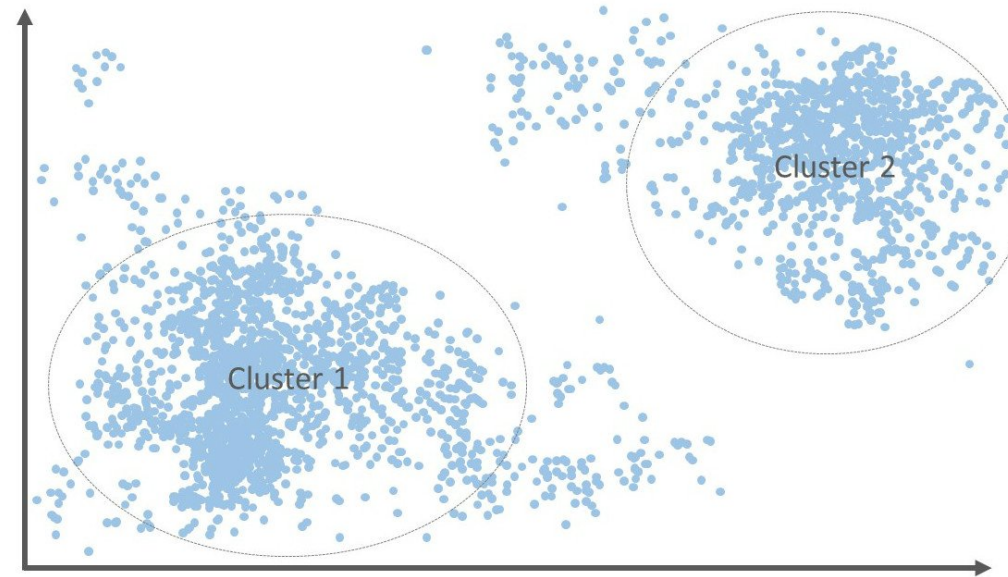
The algorithm works iteratively to assign each data point to one of *K* groups based on the features that are provided. Data points are clustered based on feature similarity.

Demonstration 3

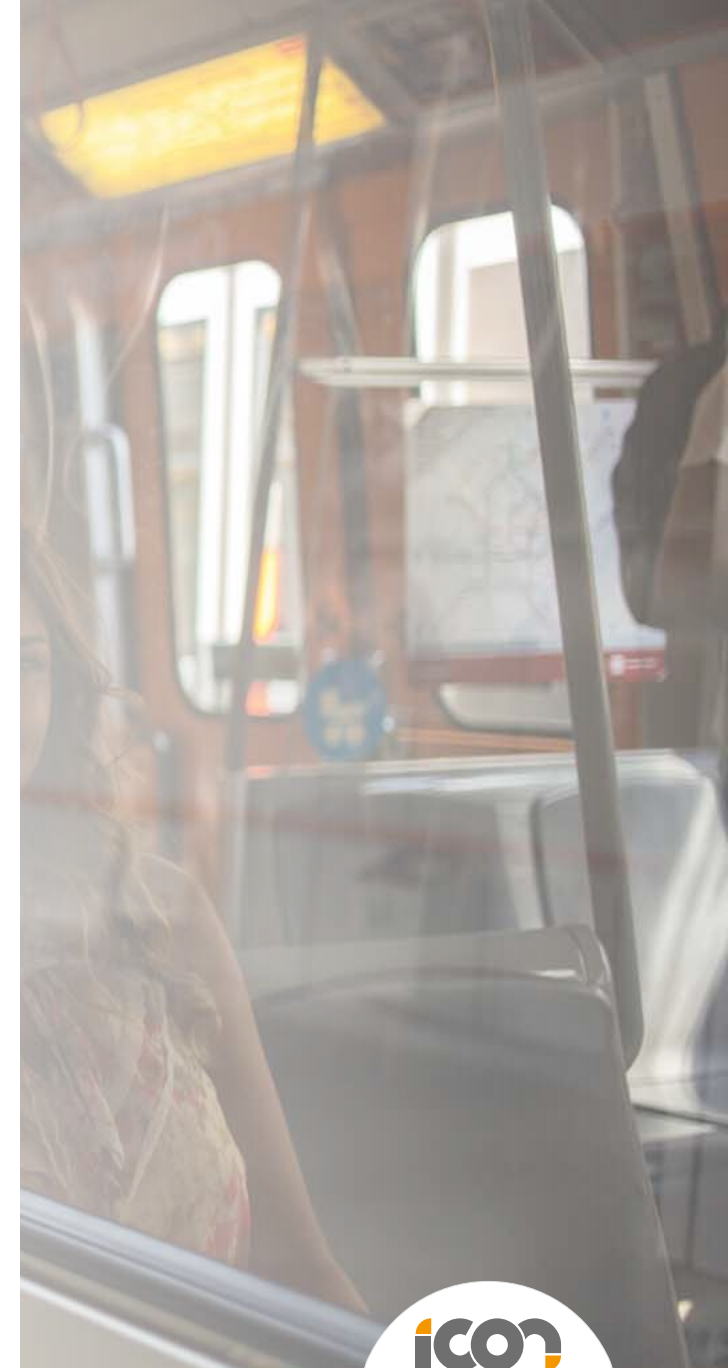
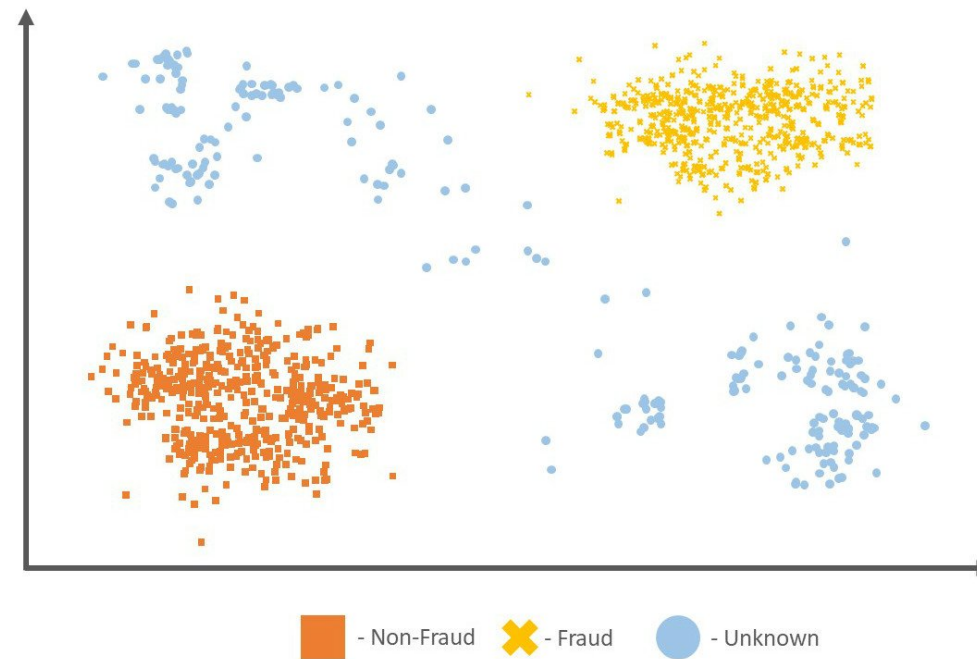
K-Nearest Neighbors Model



Unsupervised Machine Learning



Supervised Machine Learning



Machine Learning and Fraud Detection

- Payments
- Demonstration
- **Thoughts**

Two Questions Every Machine Learning Project Should Ask

Is the purpose of the project ethical?

Is the implementation of the project ethical?

Two Questions Every Machine Learning Project Should Ask

Is the purpose of the project ethical?

what are the additional benefits of the project?
who does it benefit?

Action Fraud in More than five million fall victim to financial lost £18,000 but scams, Lloyds Bank warns

Fraud complaints hit record high as banks' new anti-scam measures delayed

steals \$3K

The New York Times

Scam victims left unprotected:

More people tricked into transferring money to,

BBC

NEWS

Is the purpose of the project ethical?

herability,

and Find It in Older People

Swindlers prey on people over 70, who lose millions each year to their shifting tactics.

Bank payment scams claim 84,000 victims

nmers lived 'like Premier League

Fraud: The Public Threat ' after conning v

ITV REPORT

23 June 2019 at 11:59am

Scammers targeting flooding vict Woman loses £300,000

Two Questions Every Machine Learning Project Should Ask

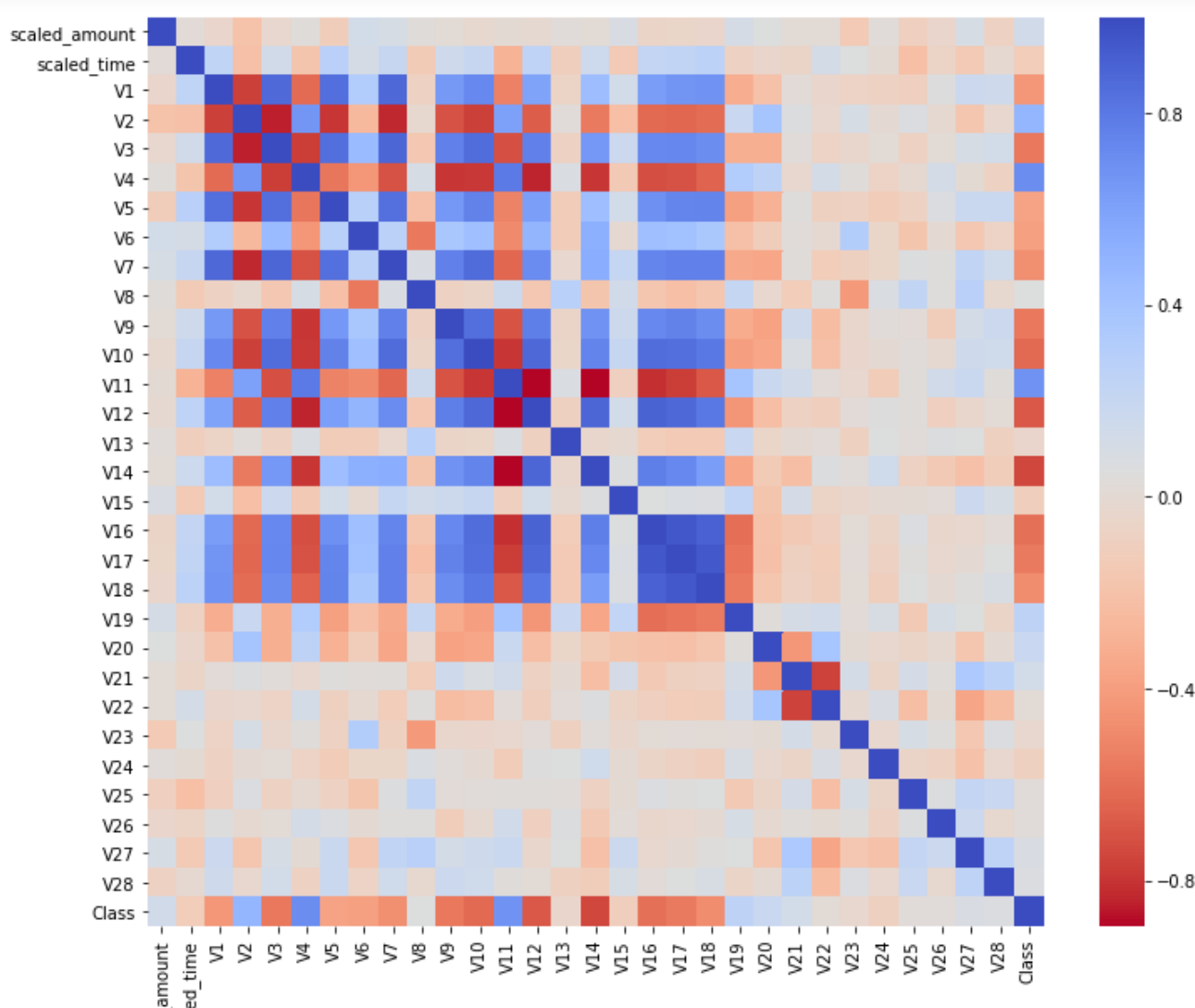
Is the implementation of the project ethical?

Does it implement unfair bias?

Disclose to stakeholders about their interactions with an AI

Governance:

- secure,
- reliable and robust, and
- appropriate processes are in place to ensure responsibility and accountability for those AI systems



Is the implementation of the project ethical?

- **Positive Correlations:** V2, V4, V11, and V19 are positively correlated. Notice how the higher these values are, the more likely the end result will be a fraud transaction.
- **Negative Correlations:** V17, V14, V12 and V10 are negatively correlated. Notice how the lower these values are, the more likely the end result will be a fraud transaction.

One last thing

Is it Intelligent?

Action Fraud investigation: Victim fall victim to financial

Fraud
anti-

More

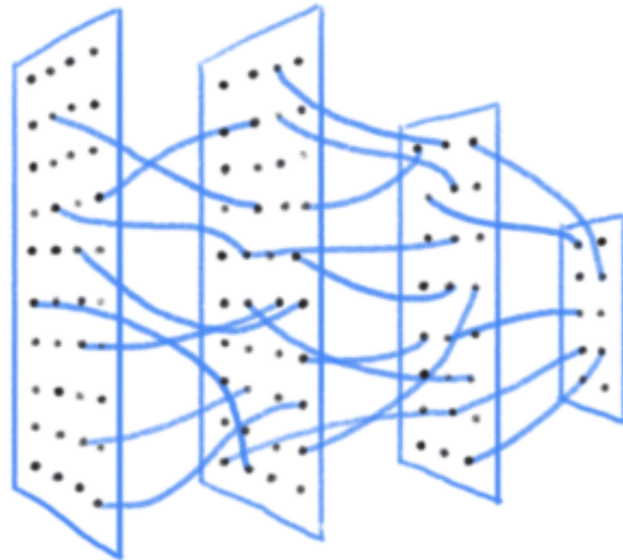
BBC

NEWS

Birmingham

Fraudulent
Transaction

Non-Fraudulent
Transaction



OUTPUT

Fraud
GOT
IT

Phone scammers lived like Premier League footballers' after conning victims out of £157,000

ITV REPORT

23 June 2019 at 11:59am

@CrosslandTamsin

targeting flooding victims Woman loses £300,000