



UNIVERSITÀ DEGLI STUDI DI MILANO
FACOLTÀ DI SCIENZE E TECNOLOGIE

Application of AI for ciphertext identification

Supervisor	Prof. Stelvio CIMATO
Co-Supervisor	Prof. Bernhard ESSLINGER
Author	Stefano SALA
Matricola	12490A
Academic year	2024-2025

dedicato a mamma, papà, Francesco e Alice ...

Astratto

In un mondo in cui l'intelligenza artificiale è sempre più in sviluppo, sia nella teoria che nella pratica, questa tesi svolge il compito di introdurre al progetto NCID di CrypTool, il quale utilizza il Machine Learning per il rilevamento del cifrario dato un testo cifrato.

L'obiettivo della tesi è introdurre i framework Keras e PyTorch, con particolare attenzione alla traduzione dell'attuale implementazione del modello di allenamento e valutazione da Keras a PyTorch. Le differenze tra i due framework vengono analizzate nel dettaglio, insieme alle rispettive caratteristiche e strategie di implementazione.

Infine, la tesi si propone di confrontare i risultati ottenuti dalla nuova implementazione con quelli della versione originale, al fine di analizzarne le differenze e trarre delle conclusioni significative.

Abstract

In a world where artificial intelligence continues to evolve both theoretically and practically, this thesis presents the NCID project by CrypTool, which uses machine learning techniques to identify the type of cipher used in a given ciphertext.

This thesis focuses on the use of the Keras and PyTorch frameworks, with particular attention to the translation of the current model training and evaluation implementation from Keras to PyTorch. The differences between the two frameworks are examined in detail, together with an analysis of their characteristics and implementation strategies.

Finally, the thesis aims to compare the results obtained from the new implementation with those of the original version, in order to analyze the differences and draw meaningful conclusions.

Contents

Listings	7
List of Figures	9
1 Introduction	10
1.1 NCID app	10
1.2 Cryptography aspects	11
1.3 Developer aspects	12
2 Theory and Foundations	13
2.1 Machine Learning basics	13
2.2 Machine Learning architectures	15
2.2.1 FFNN	16
2.2.2 LSTM	17
2.3 Training process	18
2.4 Statistical features	21
2.5 Keras vs PyTorch	22
3 Implementation	24
3.1 Gutenberg library	25
3.2 Preprocess texts and calculate statistics	25
3.3 Fit and validate model	25
3.3.1 Imports	26
3.3.2 FFNN PyTorch class	26
3.3.3 FFNN train function	28
3.3.4 FFNN prediction function	37
3.3.5 LSTM PyTorch class	38
3.3.6 LSTM train function	40
3.3.7 LSTM prediction function	41
3.3.8 Other changes	41
3.4 Evaluation code of the model	44
3.4.1 Benchmark function	44
3.4.2 Model loading	44
3.4.3 Input validation	46
4 Evaluation and results	47
4.1 First runs	48

4.2	DGX runs	49
4.2.1	Results from Keras	51
4.2.2	Results from PyTorch	53
4.3	Comparison	55
5	Conclusion	57
6	Acknowledgments	58
7	Bibliography	59
8	Appendix: Full Python code	61
8.1	Full listing of <code>train.py</code>	61
8.2	Full listing of <code>eval.py</code>	79

Listings

2.1	Keras fit implementation	23
3.1	Imports for PyTorch	26
3.2	Definition of the TorchFFNN class	26
3.3	Example of input	27
3.4	First linear layer (Linear input → hidden)	27
3.5	ReLU activation function	27
3.6	Example of output	27
3.7	Adam optimizer and Loss function	28
3.8	Loop structure, code lines from train.py, see appendix from 155 to 161	30
3.9	Batches before being splitted	31
3.10	Batches after being splitted	31
3.11	Splitting the data between trainin and validation	31
3.12	Wrong implementation	32
3.13	Example of right implementation	32
3.14	Right implementation, see lines from 175 to 194 of the appendix	33
3.15	Evaluation 8.1	34
3.16	Early stopping check	37
3.17	Custom predict function for FFNN	38
3.18	Input LSTM class	39
3.19	Input LSTM class	39
3.20	Forward method LSTM class	40
3.21	LSTM tensor conversion	40
3.22	FFNN tensor conversion	41
3.23	Printing the model summary	41
3.24	Printing the model summary of FFNN	41
3.25	Printing the model summary of LSTM	42
3.26	Saving the model with correct extension	43
3.27	Saving the model with correct extension	43
3.28	Newer input validation	43
3.29	New evaluation for FFNN and LSTM architecture	44
3.30	Setting model parameters	45
3.31	Changes to Ensemble architecture	45
3.32	Newer input validation	46
4.1	Docker run command	47

4.2	Redundant normalization	48
4.3	Shell script for PyTorch	50
	./includes/train.py	61
	./includes/eval.py	79

List of Figures

1.1	Screenshot of NCID app	11
2.1	Structure of a feedforward neural network with two hidden layers (blue).	16
2.2	Internal structure of an LSTM unit	18
3.1	Training process [6]	24

1 Introduction

During my studies, I developed a strong interest in cryptography, and in particular in the breaking of encrypted texts and all the mathematical logic behind it. For this reason, I decided to become part of the CrypTool (CT) team. CrypTool¹ is an international open-source project which develops free e-learning software for illustrating cryptographic and cryptanalytic concepts. One of their outcomes is CrypTool-Online (CTO), a web-based tool that offers knowledge about cryptology and free tools to self-educate and to teach about cryptology. Part of CTO is the Neural Cipher Identifier (NCID) app² which can be found on GitHub [1].

The task assigned to me was to implement modifications to the code base that is used in the **Neural Cipher Identifier (NCID)** app. Specifically, the modifications consisted of implementing two different Machine Learning architectures among the five available, namely FFNN and LSTM (see section 2.2), migrating from the current Keras framework to the new PyTorch framework.

Therefore, in order to better understand the process carried out, we first need to introduce some aspects: What is NCID? How is it related to cryptography? What tools did I use to implement these modifications?

1.1 NCID app

The NCID project started as a master thesis supervised by the University of Applied Sciences Upper Austria, Hagenberg, and the CrypTool project. The NCID project is based on the 2021 paper by Nils Kopal “*Of Ciphers and Neurons – Detecting the Type of Ciphers Using Artificial Neural Networks*” [2]. Of course, this project has been further developed and improved over the years, and in the theory chapter we will see who implemented the various ciphers and architectures of the project (see section 2.2).

The NCID app, which is shown in fig. 1.1, allows the user to input an encrypted text string of at least 30 characters, and from that identify the cipher used for the encryption. For that, NCID uses several neural networks from which one or more can be selected. By doing so, it removes a part that can be considered complicated for a user who is approaching the world of cryptography for the first time.

¹CrypTool: <https://www.cryptool.org/en/>

²NCID: <https://www.cryptool.org/en/cto/ncid/>

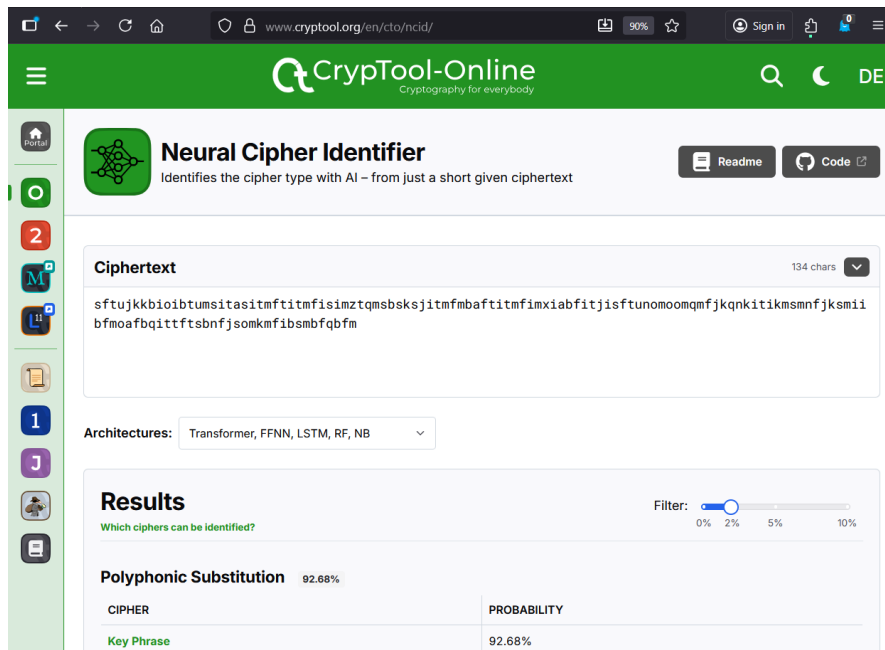


Figure 1.1: Screenshot of NCID app

In fact, in cryptanalysis, the identification of the cipher is the first step required to be able to decrypt an encrypted text. If you do not know which algorithm (cipher) was used to encrypt the text, you can hardly choose the correct attack technique or analysis method. For this reason, decryption often begins with statistical or structural observations on the ciphertext, which serve precisely to hypothesize which cipher has been used. In this regard, we will introduce in the theory chapter what are the statistical features used that allow the model to provide us with an accurate result, that is a correct result in predicting the type of cipher.

1.2 Cryptography aspects

ACA refers to the American Cryptogram Association, an organization founded in 1930 in the United States, dedicated to enthusiasts of classical cryptography and cryptogram puzzles. It is a community of amateur and professional cryptographers that for nearly a century has published the bulletin *The Cryptogram*, where classical ciphers (substitution, transposition, Vigenère, Playfair, etc.) are proposed and solved, using their own notation and standards.

The NCID project initially intended to solve a challenge based on identifying ACA cipher types³. Starting with a paper by Kopal [2] only the 56 ACA ciphers were originally planned. Later, the project was extended to recognize the five rotor ciphers — Enigma, M-209, Purple, SIGABA and Typex — also thanks to the work of Dalton and Stamp [3].

³See <https://mysterytwister.org/challenges?search=Cipher+ID> for more background.

The NCID app obviously does not aim at detecting modern ciphers, which require enormous computational power to break. First of all the practical problem: modern key spaces are gigantic. For example, a 128-bit key requires up to 2^{128} attempts for a brute-force attack. In addition, modern cryptography is not just simple encryption: it uses modes of operation, authentication, salting, password-iteration schemes, and more. Even if a 'partial' vulnerability was found, the combination of mechanisms often protects the overall communication.

1.3 Developer aspects

For the technical implementation of this project, I used Python, and more specifically I worked closely with the Keras and PyTorch frameworks. Keras is a high-level framework for building ML models, while PyTorch is lower-level and allows the programmer a much broader modification of the data flow (see section 2.5).

All the code provided to me comes from the project's GitHub repository⁴, which I was able to fork and then modify and test on my local machine. In particular, we will discuss two files: `train.py` and `eval.py` (see appendices on pages 61 and 79). The first contains all the functions necessary for training the model, while the second includes everything needed for its evaluation. This aspect will then be further explored in the implementation chapter.

⁴See [1]

2 Theory and Foundations

This chapter focuses on outlining the foundational concepts necessary to address the topics covered in this thesis. We begin by examining the fundamentals of machine learning (section 2.1), after which we will shift the focus to the architectures that I modified in the project and explain their evolutionary process within it (section 2.2). We will then look at how the model training process works, highlighting some key concepts of Machine Learning and about the NCID project (section 2.3). We will continue by discussing what statistical features are (section 2.4), and finally, we will discuss the differences between the two frameworks adopted, Keras and PyTorch (section 2.5).

2.1 Machine Learning basics

”A computer program is said to learn from experience **E** with respect to some class of tasks **T** and performance measure **P**, if its performance at tasks in **T**, as measured by **P**, improves with experience **E**.”

This is the definition that Tom Mitchell gives of Machine Learning in his book [4]. As we will notice throughout our work, this definition perfectly applies to our case. In fact, our *Task* consists in identifying the cipher, the *Performance* is measured through the accuracy of the trained model in identifying the correct cipher, and the *Experience* is determined by the dataset of pairs (x, y) , where x is our input (that is, the strings encrypted with one of the 61 ciphers), and y are the labels, namely the corresponding correct cipher for each x .

But how does my model map the inputs x to our label y ? To accomplish this task, a machine learning model can be trained in different ways: supervised, unsupervised, or reinforcement learning.

Supervised learning In supervised learning, the training set (named D) is composed of pairs (input, label).

$$D_{\text{supervised}} = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$$

The goal is to learn a function f that maps the input x to its label y :

$$f : X \rightarrow Y \quad \text{such that} \quad f(x) \approx y$$

Once trained, the model must be able to predict y for new, unseen inputs x .

Unsupervised learning In this case, the dataset contains only inputs, and thus there are no labels y :

$$D_{\text{unsupervised}} = \{x_1, x_2, \dots, x_n\}$$

The goal is to discover hidden structures, patterns, or relationships in the data, and in this case the model learns a function g that transforms or groups the data:

$$g : X \rightarrow Z \quad \text{where } Z \text{ is a space of representations, clusters, or rules}$$

Reinforcement learning Reinforcement learning is a type of machine learning in which an agent learns to make decisions by interacting with an environment, aiming to maximize a cumulative reward. It works without explicit labels for each action, but the agent learns through trial and error, observing the consequences of its actions.

An example could be a robot trying to exit a labyrinth, where the environment is the labyrinth, and the actions are moving up, down, left or right. The robot explores the maze randomly (exploration phase) and records what happens after each action. Over time, it learns that some moves lead to higher rewards than others (for example +10 if it exit and -1 for every step done), so it starts to exploit (exploitation phase) what it has learned.

Clearly, our project uses supervised learning, because as anticipated we provide our model with encrypted texts as input, and we also provide the corresponding labels. After sufficiently training the model, we test it with previously unseen encrypted sentences: the better the model performs at this task, the more accurate we consider it to be.

In particular, ours is a multiclass classification problem, that is, a supervised learning task in which each input $x \in X$ belongs to at least one class among a set of $k > 2$ possible classes. Formally:

$$f : X \rightarrow \{c_1, c_2, \dots, c_k\}$$

where $k \geq 3$ and $f(x)$ returns the correct class c_i for each x .

Now let us focus on how the inputs are mapped, that is, how do we implement our function f ?

2.2 Machine Learning architectures

”The choice of architecture inherently constrains the representational capacity of a model, determining not only what it can learn, but how efficiently it can transform inputs into useful outputs” [5]

Just as different vehicles serve distinct functions, a car for daily commuting, a tractor for agricultural work, and a truck for freight transport, different machine learning architectures are designed to address specific computational tasks. Certain models excel at image classification, while others specialize in sequential data prediction or clustering analysis. Some architectures demonstrate versatility across multiple domains, while others are finely tuned for specialized applications. This diversity in design reflects the need for tailored solutions in artificial intelligence, where the choice of architecture fundamentally shapes a model’s capabilities and performance. This is precisely where our function lies: each architecture adopts a different way of processing and evaluating the inputs, yet they all share the common goal of mapping them into a meaningful output.

But which architectures does the NCID project employ? And who first introduced them?

The idea originates from the paper by Kopal [6], which employs the **Artificial Neural Network** (ANN) for cipher identification. In the paper Kopal uses the book *Make your own neural network* [7] to provide a good introduction to ANNs: neurons are interconnected through input and output links that carry signals, each associated with a specific weight. A neuron includes an activation function a , which determines its output based on the incoming inputs. Concretely, the neuron combines all incoming signals with their corresponding weights, adds a bias term b , and then applies the activation function to this aggregated result to produce the final output (see section 2.3).

A typical design in artificial neural networks is to arrange neurons into distinct layers. The input data are first provided to an input layer composed of n neurons. This layer is connected to one or more hidden layers, which process the information further. The final hidden layer is then linked to the output layer. In this structure, every neuron in one layer is connected to every neuron in the subsequent layer, this behavior is called fully connected layer. The learning, in general, is performed by adapting the weights of the connections between the neurons.

Subsequently, with the paper *A Massive Machine-Learning Approach for Classical Cipher Type Detection Using Feature Engineering*, architectures that computed features (see section 2.4) were introduced through feature engineering, namely Feedforward Neural Networks, Decision Trees, and Naive Bayes networks. But what does feature engineering actually mean?

In traditional machine learning, feature engineering plays a central role. It consists of manually designing and selecting features that best represent the data, often relying on domain expertise and statistical analysis. While this approach can yield effective results, it is limited by the knowledge and creativity of the researcher, and may not capture complex or hidden patterns in the data.

In contrast, feature learning, automatically extracts representations from raw data. Neural networks, for example, learn hierarchical features directly during the training process, reducing the need for manual intervention. This allows models to discover intricate structures that would be difficult to design by hand. However, feature learning requires larger datasets and higher computational resources compared to traditional feature engineering.

In fact, the paper *Detection of Classical Cipher Types with Feature-Learning Approaches* by Dalton and Stamp served as inspiration for the addition of rotor ciphers in the project, which was introduced by Maik Bastian.

It should be noted that Maik Bastian also had to incorporate an SVM classifier, since the main architectures, although able to distinguish between ACA ciphers and rotor ciphers, were not capable of effectively discriminating among the different types of rotor ciphers. The SVM was therefore specifically trained for this task and integrated as an additional module: whenever the primary model detected a rotor cipher, the final classification was performed by the SVM, resulting in higher accuracy in distinguishing between the five rotor ciphers considered.

To sum up, for cipher identification, NCID uses a combination of five different architectures: FFNN, LSTM, Transformer, RF and NBN. In the work I carried out, I had to translate from Keras to PyTorch the Feedforward Neural Network and Long Short-Term Memory, and we will now examine them in more detail.

2.2.1 FFNN

It is a subclass of ANN, in which the information flows only forward, from the input layer to the output layer, without cycles or connections that send the signal backward. It is the most classical and simple ANN model, sometimes also called MLP (Multilayer Perceptron). Figure 2.1 shows the structure of an FFNN. The aim of the training is to guide $f(x)$ to obtain the probability of belonging to a class (our label, y).

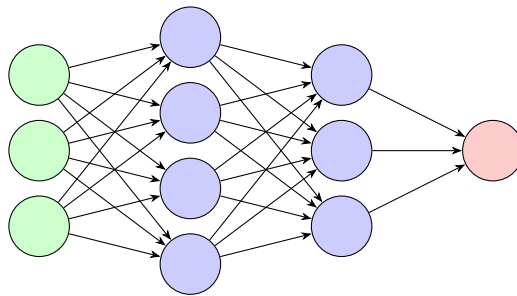


Figure 2.1: Structure of a feedforward neural network with two hidden layers (blue).

We can build deep networks with more than just two hidden layers; in fact, modern architectures often include over a hundred layers, each containing thousands of hidden layers. The term *width* refers to the number of hidden units within a single layer, while *depth* describes the number of hidden layers in the network. The total number of hidden units across all layers serves as an

indicator of the network's overall capacity.

If you are interested in exploring the topic further, I found useful consulting the paper by Goodfellow et al. [5]

2.2.2 LSTM

Long-Short Term Memory (LSTMs) is a recurrent neural network architecture designed by Sepp Hochreiter and Jürgen Schmidhuber in 1997 [8]. In contrast to FFNN, recurrent neural networks incorporate loops, enabling information from later stages of processing to be sent back to earlier stages. The LSTM has proven to be highly effective across a wide range of applications, including unconstrained handwriting recognition, speech recognition, handwriting generation, and many others [5].

To gain a clearer understanding of how the LSTM operates, its internal structure is illustrated in fig. 2.2. This graphic was designed by myself. Similar graphics can be found in Shi Yan's blog post "Understanding LSTM and its Diagrams" (2016) [9].

The LSTM architecture is built around a single core component called the memory unit, or LSTM cell. This unit comprises four neural networks, each containing an input layer and an output layer. Together, these networks form what are called the three gates: the Forget Gate, the Input Gate, and the Output Gate.

The Forget Gate determines how much of the long-term memory should be retained. This information is computed by combining the current input x_t with the previous hidden state h_{t-1} , which represents the short-term memory. A bias is then added, and the result is passed through a sigmoid activation function. The output of the Forget Gate is subsequently multiplied by the previous cell state c_{t-1} , which represents the long-term memory.

The Input Gate computes the candidate long-term memory to be added, along with the proportion of it that should be stored. This process consists of combining the current input x_t with the previous hidden state h_{t-1} , passing the result once through a sigmoid activation and once through a tanh activation. The two outputs are then multiplied element-wise, and the resulting vector is added to the cell state.

Lastly, the Output Gate updates the short-term memory by passing the combination of the current input x_t and the previous hidden state h_{t-1} through a sigmoid activation. This output is then multiplied element-wise with the tanh of the updated cell state, and the result forms the new hidden state h_t .

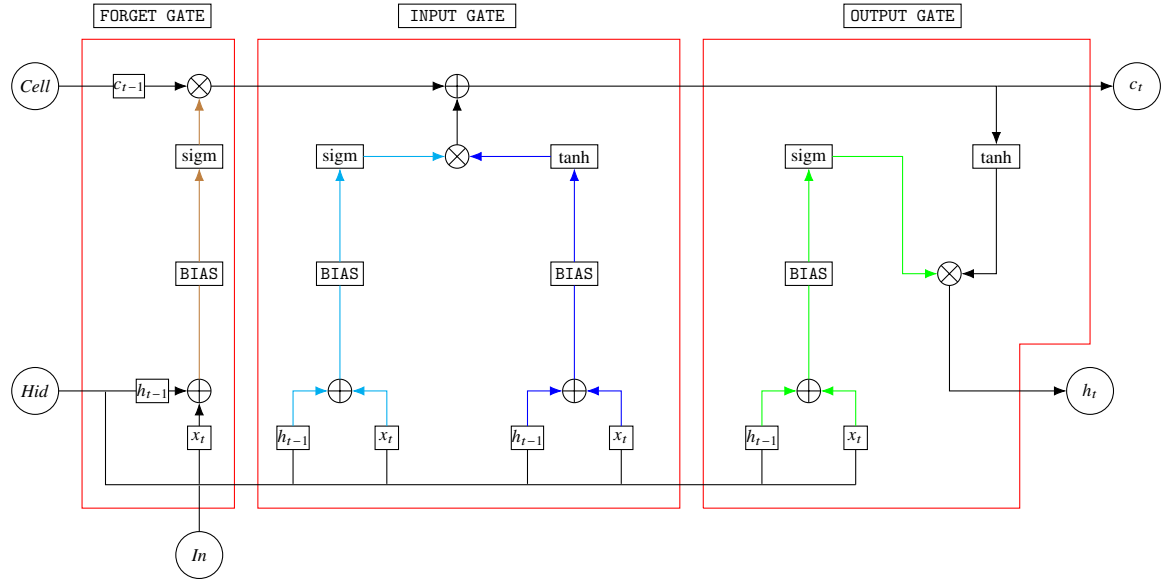


Figure 2.2: Internal structure of an LSTM unit

2.3 Training process

I will take this section to explain some fundamental concepts about ML and the project that will be useful multiple times throughout the work.

The path followed by a training process usually includes data setup, training and validation, and finally a final evaluation. In chapter 3, we will see in detail how this procedure is structured in NCID.

Let's now introduce some common Machine Learning terminology:

Weights The weights in a neural network are numerical parameters that determine how much each input influences the output. You can think of them as knobs that “weigh” the importance of each input. During training, the model searches for weight values that minimize the error between the predictions and the actual data. Therefore, the weights are not fixed: they are learned from the data.

Bias The bias in a neural network is an additional value that allows the model to shift the output independently of the inputs. It is similar to the intercept in a line $y = mx + q$, where the bias corresponds to that q . Without a bias, the output would always be constrained to pass through the origin (zero), which can limit the model's ability to fit the data. With a bias, even if all inputs are zero, the output can still take a value different from zero.

Linear layer A linear layer is used to transform the inputs into outputs through a linear combination of the values:

$$y = xW^T + b$$

where x is the input vector, i.e., the data entering the layer and W is the weight matrix. In practice, it is the basic building block of many neural networks. It serves to change the representation of the data so that the model can later recognize patterns or make predictions.

Rectified Linear Unit The activation function **ReLU** instead applies a non-linear, element-wise transformation, defined as:

$$\text{ReLU}(x) = \max(0, x)$$

In other words, negative values are replaced with zero, while positive values remain unchanged. The introduction of non-linearity is essential, as it enables the network to learn complex functions that could not be represented by linear transformations alone.

Optimizer The Adam optimizer (Adaptive Moment Estimation) is a widely used weight update algorithm in neural networks. It is a method based on computing the moments¹ of the gradient: Adam maintains an exponential estimate of the first moments (the mean of the gradients) and the second moments (the mean of the squared gradients), and uses these values to dynamically adjust the learning rate² of each parameter. The combination of these two moments allows for more stable and faster convergence compared to classical methods such as standard gradient descent.

Simply put: when the network makes a prediction, we compare the result with the correct answer and compute an error (loss). The optimizer 'looks' at this error and decides how to adjust the network's weights to reduce it in the next iteration.

Softmax Softmax is a mathematical function widely used in neural networks, particularly in the output layer of multi-class classification models. Its purpose is to transform a vector of real values (i.e., the logits) into a probability vector, meaning numbers between 0 and 1 that sum to 1. Therefore, softmax is used to interpret a network's outputs as probabilities, making it easier to select the most likely class.

¹Momentum is an optimization technique that introduces a form of "inertia" in the weight update process, accumulating part of the past gradients' direction to make the descent more stable and faster.

²The learning rate is a hyperparameter that controls the step size during gradient descent, determining how much the model's parameters are updated in response to the estimated error at each iteration.

Loss function The loss function used in the NCID project is the **Cross-Entropy Loss**, typically employed for multiclass classification problems like ours. Cross-entropy measures the distance between the probability distribution predicted by the network and the true distribution of the labels. It performs the following steps:

1. Applies the softmax function to the logits (the raw output of the architectures) to transform them into probabilities.
2. Compares these probabilities with the true labels.
3. Returns the average log-loss³ across all samples in the batch. The result of this computation is a scalar tensor (a vector that contains a single value), which when printed outputs a single value representing the loss.

Tokens Tokens are the smallest units of text into which a neural network splits a sentence in order to process it. In practice, tokenization transforms raw text into a sequence of numerical elements that the model can interpret. As we will see, these will be used in the LSTM (and more generally in sequential models) because they allow the network to process sequences of elements in order, preserving and learning the temporal or contextual dependencies between tokens.

Padding Padding is a technique that consists of adding values (typically zeros) to sequences of variable length to make them all the same size. This technique is used in LSTMs because each batch of data must have the same length since matrix operations require rectangular tensors.

Hyperparameters Finally, we would like to specify the parameters used as command line inputs. These will be particularly useful during the evaluation phase, since they must remain the same in order to compare the Keras models with the PyTorch ones. A quick reminder: We will talk about samples, but these depend on the context and the architectures used, as well as on the stage of the pipeline in which they are processed. For example, in the preprocessing phase, samples consist of plaintext sentences; after encryption, they become ciphered sentences, and later they are transformed into numerical features or representations suitable for the model.

The parameter `--batch_size` indicates the batch size, that is, how many samples are used in a single weight update step during training. Larger values make training more stable but more memory-intensive, while smaller values make the model noisier but allow more frequent updates.

`--train_dataset_size` specifies the total number of samples used in each fitting phase. It is important that this number is divisible by the amount of ciphers chosen with `--ciphers`, in order to obtain a dataset that has the same number of inputs for each cipher type.

`--dataset_workers` determines how many parallel processes are used to read the samples. More workers speed up the loading, which is useful with large datasets.

³For a deeper understanding the CrossEntropyLoss section of the PyTorch guide is available here [10]

`--epochs` defines how many times the same dataset is presented to the model during training. Increasing the number of epochs improves learning, but beyond a certain point it may lead to overfitting⁴.

`--plaintext_input_directory` is the folder that contains the plaintexts used as a basis for creating ciphertexts. These plaintexts are used to generate the training data (see section 3.1).

`--model_name` defines the name of the final saved model file, which must have the `.h5` extension for Keras and `.pth` for PyTorch.

`--min_train_len` and `--max_train_len` set the minimum and maximum length, respectively, of plaintexts used for training. If the value is set to `-1`, no constraint is applied. The same logic applies to `--min_test_len` and `--max_test_len`, which refer instead to the data for the evaluation.

`--architecture` defines the model architecture to be used. Combinations are also possible, such as `[FFNN,NB]` or `[DT,ET,RF,SVM,kNN]`, as well as a special case for rotor ciphers (`SVM-Rotor`).

2.4 Statistical features

Features are numerical measures that can be computed from the ciphertext, highlighting certain regularities. These regularities arise from the fact that the encrypted message retains statistical correlations with the original language (English, Italian, etc.). Essentially, statistical features are tools to quantify how much a ciphertext deviates from random text and how much it “resembles” natural language text. These differences correspond to the known weaknesses of the ACA ciphers and allow them to be identified and attacked.

One of the most widely used features in cryptanalysis is the Index of Coincidence (IoC), which is the probability that two randomly chosen letters from a text are the same. In a natural language like English, where some letters are much more frequent than others, the IoC takes on a characteristic value of around 0.066. In completely random text, the IoC is much lower, approximately 0.038.

Using substitution or transposition ciphers, this value remains unchanged, because the letters are simply rearranged. Consequently, an IoC that remains close to that of natural language is a strong indication that the cipher belongs to one of these families.

In contrast, for ciphers like the Bacon cipher, the frequency distribution changes drastically. Here, the IoC no longer reflects that of the original language and may appear anomalous or much higher. Therefore, if the model observes an IoC value that is particularly high or low compared to the typical English range, it can immediately rule out a simple transposition or substitution. In this sense, the IoC becomes a discriminative feature: it allows distinguishing between cipher families

⁴Overfitting is a phenomenon that occurs when a machine learning model learns the training data too well, memorizing its noise and specific details instead of capturing general patterns.

that preserve the statistical structure of the language and others that alter it significantly. It is through this type of analysis that the model can narrow down the possible ciphers used.

The project, as cited in the paper by Kopal et al. [6], initially implemented 28 statistical features, divided into groups: distribution statistics, frequency statistics, binary features, and cipher-specific features. After a selection phase, not all 28 features proved useful: only 20 were actually used in the final models, because some did not improve accuracy or introduced noise, meaning that a feature adds no new information and is redundant relative to others already present, leading to incorrect classification. If you are interested in seeing all the features used in the project, they can be found in [6].

2.5 Keras vs PyTorch

In the Machine Learning landscape, selecting the right framework is crucial for achieving desired results. Keras [11] and PyTorch [12] are the most popular open-source frameworks (along with TensorFlow, used for large-scale machine learning and deep learning applications; Scikit-learn, used for traditional machine learning algorithms and data preprocessing; and NumPy, used for efficient numerical computations and array manipulation), both used to train and evaluate neural networks. Despite sharing a common goal, they differ in their design, making them suitable for different use cases. This section aims to underline the main differences between Keras (introduced in 2015 by François Chollet) and PyTorch (released in 2016 by Meta AI, formerly Facebook AI Research) by analyzing their principal functionalities and their respective strengths and weaknesses. In the implementation chapter, we examine the actual translation from the Keras version to the PyTorch one.

So, why translating a code from Keras to PyTorch?

There are several reasons that make the conversion of a model from Keras to PyTorch an essential modification. Firstly, there is the desire to experiment with a new framework or the need to adapt to the demands of the job market. In fact, Keras has declined in popularity in the last few years among working environments, since PyTorch offers more flexibility for coding training loops. The key differences between Keras and PyTorch reside in their architectures, how they define and train models, and the level of control they provide to the coder over those processes.

Keras presents itself as a modular library designed for fast training and experimentation with deep neural networks. Its architecture favors quick and easy prototyping, offering straightforward “building blocks”. It is considered a “plug-and-play” framework that allows the coder to quickly build, train, and evaluate models, simplifying the development process.

PyTorch, on the contrary, is a low-level framework. It offers a more flexible approach and an imperative programming style, where network behavior is defined directly in Python code.

But how do we train and evaluate a model with those frameworks?

In Keras, training and evaluating a model is fairly easy, since the whole training process is managed by the `fit` method. For example, in the Keras implementation the FFNN was implemented like this:

```

1 history = model.fit(statistics, labels,
2                     batch_size=args.batch_size,
3                     validation_data=(val_data, val_labels),
4                     epochs=args.epochs,
5                     callbacks=[early_stopping_callback,
6                               tensorboard_callback,
7                               custom_step_decay_lrate_callback,
8                               checkpoint_callback])

```

Listing 2.1: Keras fit implementation

In PyTorch, instead, the architecture is divided into two parts: the definition of the states and the definition of the path that the data has to follow during the training. Training and evaluation require manual writing of complete training and test loops. This must include the upload of data inside batches, the passage of data through the network, loss calculation, backpropagation, and weight updates. [13]

Without delving into much theory, in pseudo code a neural network training loop looks like this: [14]

Algorithm 1 Neural network training loop

```

for each epoch do
    for each batch do
        Load data
        Pass data through network
        Calculate losses
        Calculate gradients
        Adjust network weights through backpropagation
    end for
end for

```

At this point you are probably wondering: Does PyTorch really expect me to write all that code? The answer is yes, but as we said, it comes with benefits for the programmer who wants to experiment with Machine Learning.

In summary, the decision between Keras and PyTorch should be driven by your project goals and your familiarity with deep-learning concepts. If you need to prototype quickly or teach fundamental models, Keras is likely the better fit; if you require research-grade flexibility, and fine-grained control, PyTorch will better meet those demands. If you are eager to know what can PyTorch accomplish I suggest you to read the implementation section 3.3.3 about the validation steps.

3 Implementation

In this chapter, we examine the code modifications that were introduced, explaining both their purpose and the rationale behind each change. Figure 3.1 shows the training process of the cipher classification model. These are the fundamental steps performed by the code to create a new model, starting from the extraction of texts from the Gutenberg Library and ending with the evaluation of the trained model. In the following paragraphs, each of these steps will be analyzed in detail, with particular emphasis on the `fit` (i.e., the actual training) and validation stages, as they required the most significant modifications.

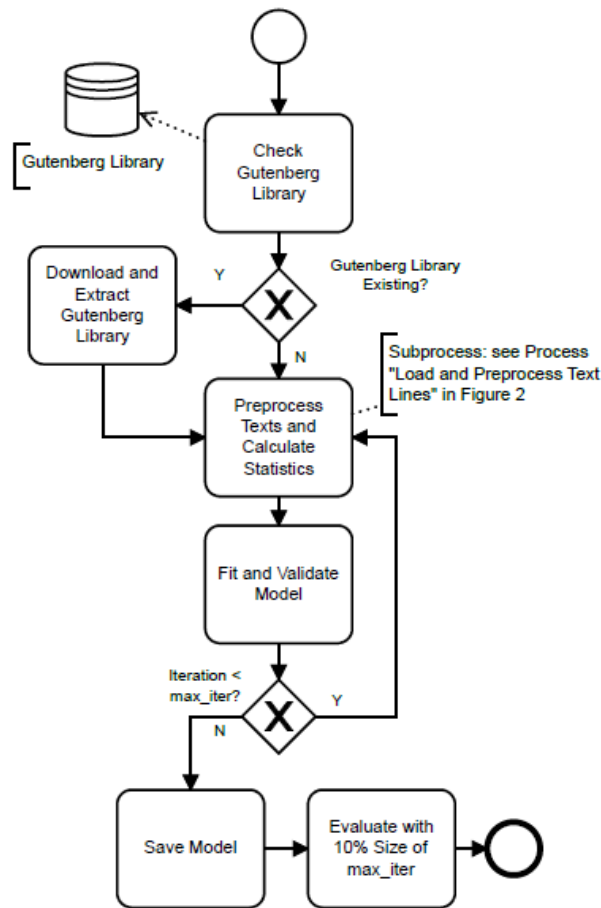


Figure 3.1: Training process [6]

3.1 Gutenberg library

The initial step involves extracting plaintext. For this purpose, we utilize the Gutenberg Library, an open-source collection of public domain texts, which provides a diverse set of literary works suitable for training and testing our models. These texts serve as the basis for generating the encrypted input used in our supervised learning tasks.

In the previous Keras-based implementation, the procedure for handling the Gutenberg plaintext dataset was the same as in the PyTorch version. The dataset was downloaded from the same public Google Drive link, automatically extracted into a temporary directory, and then relocated to the project's designated plaintext folder (`data/gutenberg_en`). After the extraction, auxiliary directories created by the download manager were removed, leaving only the necessary dataset in the target directory.

This design choice was not altered during the transition from Keras to PyTorch because the handling of plaintext data is a preprocessing step independent of the training framework. Both frameworks rely on the same input format: plaintext files are required to generate ciphertext samples through encryption with the implemented classical ciphers, which are then converted into statistical feature vectors or character sequences for model training. Consequently, the data pipeline, from downloading and extracting the Gutenberg texts to splitting them into training and testing subsets, remained unchanged.

3.2 Preprocess texts and calculate statistics

In this phase, the plaintexts are first encrypted using different cryptographic algorithms and then preprocessed to be transformed into a numerical form suitable for training machine learning models. The preprocessing consists of text normalization (for example, converting all letters to lowercase), the removal or replacement of unrecognized symbols, and the conversion of character sequences into numerical vectors. In summary, this phase makes it possible to translate the encrypted texts into standardized numerical representations, making them comparable to one another and enabling the models to learn the distinctive patterns of the different ciphers.

Like the previous plaintext extraction phase, this step did not require any changes, since the data are preprocessed in the same way for both Keras and PyTorch. The only precaution is to provide them, during training, in the appropriate format: tensors for PyTorch and NumPy arrays for Keras.

3.3 Fit and validate model

Let's now move on to the major changes.

Initially, I tried to understand how the GitHub repository was formed [1]. The file that we will examine now is called `train.py`, so every modification from now on (page 25 to 43) has been

done on that file. I was first focused on understanding how the training sessions worked, which I did by testing and running some basic training scripts on my local machine. These experiments were mainly aimed at grasping the overall workflow of the code.

Once I understood the overall functioning, I began working on the changes related to the training and validation of the model. In order to translate the code from Keras to PyTorch there are several logical steps to follow.

3.3.1 Imports

The first step consisted of importing the PyTorch library along with all the necessary modules required to run the code.

```
1 import torch
2 import torch.nn as nn
3 import torch.optim as optim
4 from torchinfo import summary
5 import numpy as np
6 from torch.utils.data import TensorDataset, DataLoader
```

Listing 3.1: Imports for PyTorch

These include essential components such as `torch.nn` for defining neural networks, `torch.optim` for optimization algorithms, and `torch.utils.data` for managing datasets and data loaders. Additionally, `torchinfo` was imported to provide a summary of the model's structure, and NumPy was used for numerical operations and compatibility with existing data structures.

We will now delve into the details of the individual architectures.

3.3.2 FFNN PyTorch class

As anticipated in the theory chapter I personally worked on and translated the FFNN and LSTM (see section 2.2), and in the following sections we will examine in detail the code modifications that were introduced in order to use them with the new framework.

Let's start from the FFNN class.

```
1 class FFNN(nn.Module):
2     def __init__(self, input_size, hidden_size, output_size,
3         ↪ num_hidden_layers):
4         super().__init__()
5
6         self.input_size = input_size
7         self.hidden_size = hidden_size
8         self.output_size = output_size
9         self.num_hidden_layers = num_hidden_layers
10
11         layers = [nn.Linear(input_size, hidden_size), nn.ReLU()]
12         for _ in range(num_hidden_layers - 1):
```

```

12         layers += [nn.Linear(hidden_size, hidden_size), nn.ReLU()]
13         layers.append(nn.Linear(hidden_size, output_size))
14         self.net = nn.Sequential(*layers)
15
16     def forward(self, x):
17         return self.net(x)

```

Listing 3.2: Definition of the TorchFFNN class

To gain an intuitive understanding of how the network works, let us consider a simplified example. Suppose we have an FFNN model with three input neurons, two hidden layers of five neurons, and an output layer with two possible classes (for instance, distinguishing whether a text belongs to cipher A or cipher B, so a binary classification).

As starting data, we take a small batch consisting of four vectors, each with three numerical features:

```

1 model = FFNN(input_size=3, hidden_size=5, output_size=2, num_hidden_layers=2)
2
3 Input:
4 tensor([[ 0.1000,  0.2000,  0.3000],
5         [ 1.0000, -1.0000,  0.5000],
6         [ 0.0000,  0.0000,  0.0000],
7         [ 2.0000,  1.0000, -1.0000]])

```

Listing 3.3: Example of input

At this stage, the input is multiplied by the weight matrix of the first layer and added to a bias term. The result is a linear transformation that projects the data into the space of the hidden layer:

```

1 tensor([[ 0.0317,  0.2515, -0.1857,  0.0923, -0.0142],
2         [ 0.6049, -0.7834,  0.1521, -0.4418,  0.2357],
3         [ 0.0005,  0.0003, -0.0007,  0.0002, -0.0001],
4         [ 0.3672,  0.8550, -0.6538,  0.3214, -0.1579]])

```

Listing 3.4: First linear layer (Linear input → hidden)

ReLU (see section 2.3) replaces negative values with zero:

```

1 tensor([[0.0317, 0.2515, 0.0000, 0.0923, 0.0000],
2         [0.6049, 0.0000, 0.1521, 0.0000, 0.2357],
3         [0.0005, 0.0003, 0.0000, 0.0002, 0.0000],
4         [0.3672, 0.8550, 0.0000, 0.3214, 0.0000]])

```

Listing 3.5: ReLU activation function

The linear transformation and the ReLU are then applied again until the final output state is obtained. The last layer maps the data into a vector that has as many entries as the number of classes:

```

1 tensor([[ -0.1613,  0.0144],
2         [ 0.0568, -0.1409],
3         [ -0.1585,  0.0267],

```

```
4 [-0.0652, -0.0819]])
```

Listing 3.6: Example of output

These values are not probabilities, but logits. In other words, they are the “raw” outputs produced by the last linear layer of your FFNN before any transformation that normalizes them between 0 and 1. Each row corresponds to a batch example, and each column corresponds to a class. These logits can be positive or negative and are not constrained to sum to 1. When you apply the softmax function to each row, you obtain a probability vector that sums to 1, which is useful for interpreting the model and for computing the cross-entropy loss (in PyTorch, cross-entropy loss can accept logits directly, see section 2.3).

This simple example demonstrates how an initial set of numerical data, through a sequence of linear and non-linear transformations, is projected into a space that allows separation between different types of ciphers. In practice, the same logic is extended to more complex features and a larger number of classes, enabling NCID to recognize a wide range of ciphers.

An important aspect of the network implementation concerns the construction of the sequence of layers. In the code, the various layers of the network are initially stored in a Python list called `layers`. This list contains, in order: the linear transformation from the input to the hidden layer, the ReLU activation function, any additional pairs of linear layer and ReLU, and finally the output layer.

To transform this list into a proper PyTorch model, the `nn.Sequential` class is used, which allows multiple modules (`nn.Module`) to be concatenated into a single structure. In this way, the object `self.net` represents the entire neural network as an ordered sequence of operations.

The `*layers` syntax serves to “unpack” the list and pass its individual elements as arguments to `nn.Sequential`.

For example, if `layers = [A, B, C]`, then the statement `nn.Sequential(*layers)` is equivalent to `nn.Sequential(A, B, C)`.

3.3.3 FFNN train function

In this section, we will dive into the analysis of the lines of code that allow the neural network to learn, starting from setting the optimizer and the loss function, up to performing an evaluation.

Setting optimizer and loss function Let’s go through a detailed examination:

```
1 def train_torch_ffnn(model, args, train_ds):
2     device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
3     model.to(device)
4
5     optimizer = optim.Adam(
6         model.parameters(),
7         lr=config.learning_rate,
```

```

8         betas=(config.beta_1, config.beta_2),
9         eps=config.epsilon,
10        amsgrad=config.amsgrad
11    )
12    criterion = nn.CrossEntropyLoss()

```

Listing 3.7: Adam optimizer and Loss function

After defining the optimizer and the loss function (see section 2.3), the code sets the model to training mode using the command `model.train()`.

This instruction is crucial because it tells PyTorch that the model should behave as it does during training: in particular, certain layers such as *Dropout* (see the following paragraph) and *Batch Normalization* (see chapter 4) behave differently depending on whether the model is in training or validation/testing mode.

Dropout for example randomly deactivates, with a given probability (for example 50%), certain neurons at each forward pass during training to reduce overfitting. This prevents the network from relying too heavily on specific connections.

Without the `model.train()` call, the model would use the parameters fixed for evaluation mode, which could compromise the training process.

Immediately before the training loops, several control variables are also initialized:

- `best_val_acc = 0`: stores the highest accuracy achieved on the validation set so far, necessary for implementing *early stopping* (we will see it later in this section).
- `patience_counter = 0` and `patience_limit = 250`: used to count how many consecutive iterations do not improve performance. If the limit is exceeded, training is stopped to avoid wasting computational resources.
- `train_iter = 0`: keeps track of the total number of samples processed during training.
- `train_epoch = 0`: counts the number of completed epochs.
- `start_time = time.time()`: allows measuring the total duration of the training process.
- `val_data_created = False` together with `x_val` and `y_val`: ensures that the validation dataset is created only once, at the first iteration, and then reused without generating random splits at each cycle.

These assignments therefore serve to initialize the training state, so that the model can be trained, monitored, and evaluated correctly.

Training loop The code then enters the actual training phase (see listing 3.8). The outer loop for `epoch in range(args.epochs)` iterates over the number of epochs (see section 2.3). Inside, the loop `while train_ds.iteration < args.max_iter` ensures that training continues until the maximum number of iterations is reached.

The command `training_batches = next(train_ds)` retrieves a new batch of data from the generator `train_ds`. Each batch is organized as pairs of `(statistics, labels)`, where `statistics` contains the numerical features extracted from the ciphered texts i.e., the values that will be fed as input to the model while `labels` represents the corresponding class labels.

The statement `statistics, labels = training_batch.items()` separates the data from the labels. With `stats_np = statistics.numpy()` and `labels_np = labels.numpy()`, the PyTorch tensors are converted into NumPy arrays, which is useful for later splitting the dataset into training and validation sets using `train_test_split`.

For clarification, consider a numerical example: suppose the generator `train_ds` provides a batch with 4 samples, each described by 3 statistical features, along with their respective class labels. We would then have:

$$\text{statistics} = \begin{bmatrix} 0.2 & 0.5 & 0.7 \\ 0.1 & 0.4 & 0.3 \\ 0.9 & 0.8 & 0.6 \\ 0.0 & 0.2 & 0.1 \end{bmatrix}, \quad \text{labels} = \begin{bmatrix} 1 \\ 0 \\ 2 \\ 1 \end{bmatrix}$$

After conversion with `.numpy()`, we obtain the arrays

$$\text{stats_np} = \begin{bmatrix} 0.2 & 0.5 & 0.7 \\ 0.1 & 0.4 & 0.3 \\ 0.9 & 0.8 & 0.6 \\ 0.0 & 0.2 & 0.1 \end{bmatrix}, \quad \text{labels_np} = \begin{bmatrix} 1 \\ 0 \\ 2 \\ 1 \end{bmatrix}$$

Although they appear identical, these arrays can now be used either to create the PyTorch tensors required for training or to be split into a *training set* and a *validation set*.

In this way, the model receives as input the statistics and the corresponding labels, which allows the loss function to be computed and the weights to be updated via backpropagation.

```
for epoch in range(args.epochs):
    while train_ds.iteration < args.max_iter:
        training_batches = next(train_ds)
        for training_batch in training_batches:
            statistics, labels = training_batch.items()
            stats_np = statistics.numpy()
            labels_np = labels.numpy()
```

Listing 3.8: Loop structure, code lines from `train.py`, see appendix from 155 to 161

Splitting dataset In the following lines of `train.py`, the division of data into training and validation sets is handled. At the first pass, the variable `val_data_created` is set to `False`, so the function `train_test_split(stats_np, labels_np, test_size=0.3)` is executed, which splits the data into two parts: 70% for training (`x_train_np, y_train_np`) and 30% for validation (`x_val_np, y_val_np`).

Subsequently, the validation data are converted into PyTorch tensors and moved to the selected device (CPU or GPU).

After this initial split, the variable `val_data_created` is set to `True`, so in subsequent iterations the split is not repeated and the validation data remain fixed, while the set `x_train_np, y_train_np` is updated with the next batches provided by the dataset.

Suppose that the initial batch consists of four samples, each with three features, along with their corresponding labels:

```
1 stats_np = [[0.2, 0.5, 0.7],
2             [0.1, 0.4, 0.3],
3             [0.9, 0.8, 0.6],
4             [0.0, 0.2, 0.1]]
5
6 labels_np = [1, 0, 2, 1]
```

Listing 3.9: Batches before being splitted

By splitting with `test_size=0.3`, a possible outcome could be:

```
1 x_train_np = [[0.2, 0.5, 0.7],
2             [0.1, 0.4, 0.3],
3             [0.0, 0.2, 0.1]]
4
5 y_train_np = [1, 0, 1]
6
7 x_val_np = [[0.9, 0.8, 0.6]]
8
9 y_val_np = [2]
```

Listing 3.10: Batches after being splitted

In this example, three samples are used for training and one for validation. After conversion to PyTorch tensors, `x_val` and `y_val` will remain unchanged until the end of training, ensuring that the model's performance is always evaluated on the same reference data.

```
1 if not val_data_created:
2     x_train_np, x_val_np, y_train_np, y_val_np = train_test_split(stats_np,
3         ↪ labels_np, test_size=0.3)
4     x_val = torch.tensor(x_val_np, dtype=torch.float32).to(device)
5     y_val = torch.tensor(y_val_np, dtype=torch.long).to(device)
6     val_data_created = True
7 else:
8     x_train_np = stats_np
```

```
8 y_train_np = labels_np
```

Listing 3.11: Splitting the data between trainin and validation

Optimization of the gradients Now we move on to a very important part of the training, the optimization of the gradients that is the process of updating the weights of a neural network during training. The step we are going to analyze was initially implemented differently and produced poorly performing models, but thanks to the help of Maik Bastian, I was able to solve it. Below is a snippet of the code I had first implemented in listing 3.12:

```
1
2 optimizer.zero_grad()
3 outputs = model(x_train)
4 loss = criterion(outputs, y_train)
5 loss.backward()
6 optimizer.step()
```

Listing 3.12: Wrong implementation

The issue with these lines is that, regardless of the dataset size, the optimization is performed only once, whereas in the Keras implementation it was applied for every batch. The problem of calling the optimizer only once for the entire dataset (i.e., after processing all the samples of that batch) instead of after each batch is that the model performs too few weight updates. For example, if you have a dataset of 10k samples and process it as a single batch, the optimizer is updated only once every 10,000 samples. In practice, the network executes just one averaged backpropagation over the whole block, causing the model to learn much more slowly. This problem is known as a *larger effective batch size*.

To achieve behavior similar to Keras (mini-batches and more frequent updates), you should manually split x_{train} and y_{train} into mini-batches of size `args.batch_size`, as illustrated in listing 3.14.

In listing 3.13 is shown an example of the flow of data with the right implementation where we have a binary classification task. I need to emphasize to the reader that this example is not meant to follow precise calculations, but rather to understand the correct way to split the data. For this reason, the input data are represented by randomly chosen integers, as are the outputs of the loss function:

```
1
2 x_train_np = [
3     [1, 2],
4     [2, 1],
5     [3, 4],
6     [4, 3],
7     [5, 6],
8     [6, 5]
9 ]
10
11 y_train_np = [0, 0, 1, 1, 0, 1]
12
13 # Batch 1
```



```

14 x_batch1 = [[2, 1], [4, 3]]
15 y_batch1 = [0, 1]
16
17 # Batch 2
18 x_batch2 = [[5, 6], [1, 2]]
19 y_batch2 = [0, 0]
20
21 # Batch 3
22 x_batch3 = [[3, 4], [6, 5]]
23 y_batch3 = [1, 1]
24
25 outputs_batch1 = [[1.0, 0.5], [0.2, 1.2]]
26 outputs_batch2 = [[0.8, 0.4], [1.1, 0.3]]
27 outputs_batch3 = [[0.3, 0.7], [0.1, 1.0]]
28
29 loss_batch1 = 0.6
30 loss_batch2 = 0.5
31 loss_batch3 = 0.4
32
33 batch_losses = [loss_batch1, loss_batch2, loss_batch3]
34
35 epoch_loss = sum(batch_losses) / len(batch_losses)
36 print("Average Loss:", epoch_loss) # Output: 0.5

```

Listing 3.13: Example of right implementation

With the wrong implementation (see listing 3.12, `optimizer.step()` would update all the weights only once, using the gradients computed over all 6 samples together. Thus, to achieve this goal, it is not necessary to change the previous steps where optimization was performed, but simply to add a for loop that does it for each `x_batch` and `y_batch` contained in the `train_loader`. The lines before the for loop are used to prepare the data as PyTorch tensors and to create a PyTorch dataset that associates each input with its corresponding label:

```

1
2 x_train = torch.tensor(x_train_np, dtype=torch.float32)
3 y_train = torch.tensor(y_train_np, dtype=torch.long)
4
5 train_dataset = TensorDataset(x_train, y_train)
6 train_loader = DataLoader(train_dataset, batch_size=args.batch_size, shuffle=
    ↪ True)
7
8 batch_losses = []
9
10 for x_batch, y_batch in train_loader:
11     x_batch = x_batch.to(device)
12     y_batch = y_batch.to(device)
13
14     optimizer.zero_grad()
15     outputs = model(x_batch)
16     loss = criterion(outputs, y_batch)
17     loss.backward()
18     optimizer.step()
19

```

```

20     batch_losses.append(loss.item())
21     train_iter += len(y_batch)
22
23 epoch_loss = sum(batch_losses) / len(batch_losses)

```

Listing 3.14: Right implementation, see lines from 175 to 194 of the appendix

Evaluation The next step of our training loop consists of computing predictions using the model that we are training, see listing 3.15

```

199
200 model.eval()
201 with torch.no_grad():
202     val_outputs = model(x_val)
203     val_loss = criterion(val_outputs, y_val)
204     val_pred = torch.argmax(val_outputs, dim=1)
205     val_acc = (val_pred == y_val).float().mean().item()
206
207     top3 = torch.topk(val_outputs, k=3, dim=1).indices
208     y_val_exp = y_val.unsqueeze(1).expand_as(top3)
209     val_k3 = (top3 == y_val_exp).any(dim=1).float().mean().item()

```

Listing 3.15: Evaluation 8.1

The command `model.eval` in PyTorch is used to set the model into evaluation mode. By switching to `model.eval`, we inform PyTorch that training is no longer being performed, but rather validation or testing, and that the learned parameters should be used while disabling stochastic components. By stochastic components, we refer to mechanisms that introduce randomness during training, such as Dropout. In evaluation mode, Dropout is disabled and all neurons remain active. The context manager `torch.no_grad` in PyTorch specifies that, within its block of code, gradients should neither be computed nor stored. During training, PyTorch keeps track of all operations performed on tensors by building a computational graph, since at the end gradients are required through the `loss.backward()` operation. This process, however, consumes memory and involves additional computations. During validation or inference (that is, when using `model.eval`), gradients are not needed, as only the model predictions are of interest.

We can now examine how predictions are computed. For this purpose, we will consider an example with `train_dataset_size = 10` (see section 2.3), a total of 7 features, and only 5 cipher classes. The command `model(x_val)` executes the forward pass, where `x_val` is the input batch, a tensor containing the calculated feature values (see section 2.4) of the validation set for each sample given.

In this example, the input tensor has the shape `x_val.shape = (3, 7)`. It should be emphasized, however, that in the actual project a total of 724 feature values are employed. To see where in the code this is defined, refer to the appendix at line 504. As you can notice, the number 724 is assigned to the variable `input_layer_size`. This is because each neuron in the input layer corresponds to one feature.

An example of `x_val` and the corresponding `val_outputs` after the forward pass is shown below:

```
1 x_val =
2 tensor([[0.12, 0.05, 0.33, 0.21, 0.08, 0.17, 0.02],
3         [0.07, 0.18, 0.22, 0.10, 0.14, 0.05, 0.04],
4         [0.30, 0.11, 0.09, 0.25, 0.16, 0.08, 0.12]])
5 x_val.shape = torch.Size([3, 7])
```

The result of the forward pass is a tensor called `val_outputs` with shape (3, 5), since in this example 5 ciphers are used, while in the full project there are 61 of which 56 were implemented in the paper "A Massive Machine-Learning Approach For Classical Cipher Type Detection Using Feature Engineering" [15], and the 5 rotor ciphers were later implemented thanks to the contribution of Dalton and Stamp's paper [3].

```
1 val_outputs =
2 tensor([[ 0.80,  1.90, -0.30,  0.20,  0.50],
3         [ 2.20, -0.10,  0.70,  1.50, -0.40],
4         [-0.20,  0.30,  2.00,  0.10,  1.20]])
5 val_outputs.shape = torch.Size([3, 5])
```

This tensor will then be employed to compute both the predictions and the loss. The loss is calculated by applying the cross-entropy loss function. The result of this computation is a scalar tensor, `val_loss`, which when printed outputs a single value representing the loss.

As for the predictions, the `val_outputs` tensor is used again, this time combined with the `argmax` function to obtain the index of the maximum value along dimension 1:

```
1 val_pred = tensor([1, 0, 2])
```

Once the predictions are obtained, they can be compared with the true labels in order to compute the accuracy, which is stored in the variable `val_acc`. This value indicates how well the model is classifying the validation data. The process is as follows:

1. `(val_pred == y_val)` compares the predictions with the true labels (`y_val`), returning a vector of booleans (True if correct, False if incorrect).
2. `.float().mean()` converts the booleans to numerical values (1.0 for correct, 0.0 for incorrect) and computes the mean.
3. `.item()` extracts the value as a standard Python number.

For example, if `y_val = tensor([1, 3, 2])`, then `val_acc` will be 0.67, since only 2 out of 3 samples were predicted correctly. Thus, the model achieves an accuracy of 67%.

To compute the top-3 accuracy, which measures the percentage of times the correct class appears among the three highest predictions (rather than only the top-1), the `topk` function selects the three highest logits for each sample, producing a tensor with shape `(batch_size, 3)`, which contains the indices of the three most probable classes:

```

1 top3 =
2 [[1, 0, 4],
3  [0, 3, 2],
4  [2, 4, 1]]

```

After that, the command `y_val_exp = y_val.unsqueeze(1).expand_as(top3)` is used to make the dimensions of the true labels (`y_val`) and the multiple predictions comparable, since while `y_val` has shape (3,), `top3` has shape (3, 3). The expansion produces the following tensor:

```

1 y_val_exp =
2 tensor([[1, 1, 1],
3         [3, 3, 3],
4         [2, 2, 2]])

```

Finally, similarly to `val_acc`, the predictions are compared with the true labels using:

```

1 val_k3 = (top3 == y_val_exp).any(dim=1).float().mean().item()

```

which results in:

```

1 tensor([[ True, False, False],
2         [False,  True, False],
3         [ True, False, False]])

```

Therefore, in this example `val_k3 = 1`, corresponding to 100%.

In this example, one can observe the advantage of PyTorch over Keras. For instance, the line `val_loss = criterion(val_outputs, y_val)` shows how PyTorch allows for full flexibility in defining the loss function (here named `criterion`), not only by choosing among the standard implementations (such as `CrossEntropyLoss`) but also by creating entirely customized loss functions and directly integrating them into the training loop. This feature makes it possible to adapt the model to highly specific and complex scenarios, where a traditional error metric would not be sufficient. In Keras, by contrast, the choice of the loss function is typically limited to a predefined set of options (such as `Poisson`, `binary_crossentropy` for binary classification, or `mean_squared_error` for regression problems) or otherwise constrained by the high-level APIs.

Early stopping check The last addition was simply an early stopping check. In Keras, these types of checks are already available through the `callbacks` module, such as `EarlyStopping`, which can monitor a specified metric and stop training when it ceases to improve. In PyTorch, however, we need to implement them manually. This particular check is used to prevent overfitting by stopping the training process once the model's performance on the validation set stops improving, thus saving computational resources.

```

1 # --- Early stopping check ---
2 if val_acc > best_val_acc:
3     best_val_acc = val_acc
4     patience_counter = 0
5 else:
6     patience_counter += 1
7     if patience_counter >= patience_limit:
8         print("Early stopping triggered.")
9         elapsed = time.time() - start_time
10        t = time.gmtime(elapsed)
11        print(f"Finished training in {t.tm_yday - 1} days {t.tm_hour} hours {
            ↳ t.tm_min} minutes {t.tm_sec} seconds with {train_iter}
            ↳ iterations.")
12        class DummyEarlyStopping:
13            stop_training = True
14        return DummyEarlyStopping(), train_iter, f"Early stopped at epoch {
            ↳ epoch+1}"

```

Listing 3.16: Early stopping check

Although functional, we preferred not to use this check during the runs on the DGX machine (see chapter 4), in order to assess the model's true limits. By allowing the training to continue without early stopping, we could observe the maximum performance the model can achieve and better understand how it behaves when trained for extended periods.

And with that we end the training loop.

3.3.4 FFNN prediction function

After training and saving a model, we need to evaluate whether it accurately recognizes our ciphers. To this end, I implemented the prediction function, which evaluates the fully trained model on a dedicated test dataset, without updating the weights anymore.

Similar to the snippet observed in the evaluation section (see section 3.3.3), this function also employs `model.eval()` and with `torch.no_grad()` to disable dropout and gradient computation, ensuring stable predictions without additional memory costs. Indeed, this function also makes use of Cross Entropy Loss (see section 2.3), `torch.argmax`, `torch.topk()`, etc.

However, we are now in a different stage, namely the final testing phase which is computed by the FFNN prediction function (see lines from 355 to 393 of appendix). In the previous case, after validation, the model was switched back to `train()` mode to continue learning, while in this function the model remains in `eval()` mode for the entire test process.

Moreover, an iterative test dataset (`test_ds`) is used, which provides batches until `args.max_iter` is reached, instead of a fixed validation set created by splitting the training data. In this context, an essential aspect is that, in addition to loss and metrics, all predictions (`all_preds`) and all labels (`all_labels`) are stored and then returned by the function for subsequent calculations (e.g., confusion matrix and F1-score, which will be discussed later in chapter 4).

```

1 def predict_torch_ffnn(model, test_ds, args):
2     device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
3     model.eval()
4     model.to(device)
5     criterion = nn.CrossEntropyLoss()
6
7     all_preds = []
8     all_labels = []
9
10    with torch.no_grad():
11        while test_ds.iteration < args.max_iter:
12            testing_batches = next(test_ds)
13            for testing_batch in testing_batches:
14                statistics, labels = testing_batch.items()
15                stats_np = statistics.numpy()
16
17                x = torch.tensor(stats_np, dtype=torch.float32).to(device)
18                y = torch.tensor(labels.numpy(), dtype=torch.long).to(device)
19
20                outputs = model(x)
21                loss = criterion(outputs, y)
22
23                pred_top1 = torch.argmax(outputs, dim=1)
24                acc = (pred_top1 == y).float().mean().item()
25
26                top3 = torch.topk(outputs, k=3, dim=1).indices
27                y_expanded = y.unsqueeze(1).expand_as(top3)
28                k3_acc = (top3 == y_expanded).any(dim=1).float().mean().item()
29                ↪ ()
30
31                print(f"Eval -> Loss: {loss.item():.4f}, Accuracy: {acc:.4f},
32                ↪ Top-3 Accuracy: {k3_acc:.4f}")
33
34                preds = torch.softmax(outputs, dim=1).cpu().numpy()
35                all_preds.append(preds)
36                all_labels.append(labels.numpy())
37
38    all_preds = np.concatenate(all_preds, axis=0)
39    all_labels = np.concatenate(all_labels, axis=0)
40
41    return all_preds, all_labels

```

Listing 3.17: Custom predict function for FFNN

3.3.5 LSTM PyTorch class

It is now time to discuss our second architecture, namely the LSTM. The main difference between the two architectures lies in the class definition, that is, how the function maps inputs to outputs. In fact, the structure of the LSTM class is very different from that of the FFNN, while the training process, as we will see, is identical for both.

Let us start by examining the inputs received by the class:

The LSTM class defines a recurrent neural network model based on LSTM. The parameter `vocab_size` indicates the size of the vocabulary, i.e., the total number of distinct tokens (see section 2.3) present in the input data, and it is used to build the initial *embedding layer*. The embedding dimension is specified by `embed_dim`, which determines how rich the representation of each token will be: larger embeddings can capture more detailed information but also increase the model's complexity. The parameter `hidden_size` represents the number of hidden units in the LSTM and defines the network's capacity to store long-term information. `output_size` indicates the number of neurons in the final layer, corresponding to the number of classes in the classification problem or the desired output dimension. The parameter `num_layers` establishes the number of stacked LSTM layers, while `dropout` specifies the dropout probability between LSTM layers, useful to reduce overfitting when multiple layers are used.

```

1 class LSTM(nn.Module):
2     def __init__(self, vocab_size, embed_dim, hidden_size, output_size,
3         num_layers=1, dropout=0.0):
4         super().__init__()
5
6         self.vocab_size = vocab_size
7         self.embed_dim = embed_dim
8         self.hidden_size = hidden_size
9         self.output_size = output_size
10        self.num_layers = num_layers
11        self.dropout = dropout
12
13        self.embedding = nn.Embedding(
14            num_embeddings=vocab_size,
15            embedding_dim=embed_dim,
16            padding_idx=0
17        )

```

Listing 3.18: Input LSTM class

Within the class, the LSTM layer itself is defined using `nn.LSTM`. This PyTorch function is used to define the structure of our architecture, to which the layers and parameters such as dropout are passed. After the LSTM layer, a linear layer (`self.fc`) is defined, which takes the hidden state of the LSTM as input and produces the final output of dimension `output_size`, corresponding to the classes of the classification problem or the desired output dimension.

```

1         self.lstm = nn.LSTM(
2             input_size=embed_dim,
3             hidden_size=hidden_size,
4             num_layers=num_layers,
5             batch_first=True,
6             dropout=dropout if num_layers > 1 else 0.0
7         )
8         self.fc = nn.Linear(hidden_size, output_size)

```

Listing 3.19: Input LSTM class

The `forward` method defines the flow of data through the LSTM model. Initially, the dimension of the input tensor `x` is checked. If `x` has three dimensions and the last dimension is equal to 1, it is removed using `x.squeeze(2)` to ensure that the input has the correct shape for the embedding layer.

Next, the input tokens are converted into dense vectors, that is a vector where every element contains a significant information, through the embedding layer defined as `self.embedding(x)`, producing the tensor `emb`. This tensor is then passed to the LSTM layer, `self.lstm(emb)`, which returns two main values: `output`, containing the hidden states for each time step, and `hidden`, representing the last hidden state of each LSTM layer.

To obtain a compact representation of the entire sequence, the last hidden state of the final layer is selected, `last_hidden = hidden[-1]`. This vector summarizes the sequential information learned by the LSTM. Finally, `last_hidden` is passed through the linear layer `self.fc` to obtain the logits.

```

1     def forward(self, x):
2         if x.dim() == 3 and x.size(2) == 1:
3             x = x.squeeze(2)
4
5         emb = self.embedding(x)
6         output, (hidden, _) = self.lstm(emb)
7         last_hidden = hidden[-1]
8         logits = self.fc(last_hidden)
9
10        return logits

```

Listing 3.20: Forward method LSTM class

3.3.6 LSTM train function

The LSTM training function in `train.py` is generic with respect to the type of model: for both FFNN and LSTM, the training loop handles taking data batches, computing the loss, backpropagating the gradient, and updating the weights via the optimizer. In other words, the training loop does not “know” whether it is working on an FFNN or an LSTM: it applies the same optimization procedure. The substantial difference lies in how each model transforms inputs into outputs, meaning the input to output mapping is architecturally different, but the weight update mechanism remains the same.

The only note to make is the following. During the conversion of data into Torch tensors, the LSTM performs the following conversion:

```

1 stats_np = statistics.numpy().astype(int)

```

Listing 3.21: LSTM tensor conversion

In contrast, as we saw in section 3.3.3 the FFNN simply performs:


```
1 stats_np = statistics.numpy()
```

Listing 3.22: FFNN tensor conversion

The difference between the two cases is not arbitrary, but depends on the type of data each architecture expects as input. In `train.py`, when working with an FFNN, `statistics` are already continuous numerical features (floats), derived from the computation of statistics on the encrypted text, so it is sufficient to convert them into a float tensor using `.numpy()`. This is consistent with the linear operations and ReLU activations in the FFNN, which operate on `float32` values.

In the case of the LSTM, however, input sequences are often treated as discrete indices or integer values representing tokens or symbols. For this reason, a different conversion is applied. This way, the values are integers, which can then be transformed into embeddings or processed as sequences.

3.3.7 LSTM prediction function

3.3.8 Other changes

Additionally, I introduced two minor but necessary modifications to improve the compatibility between the legacy script and the newly implemented PyTorch-based architecture. These changes involve printing a model summary and saving the model with the appropriate file extension.

Model summary To visualize the architecture of the model, I used the `torchsummary` library. If the model does not implement its own `summary()` method, the code falls back to calling `summary()` with a predefined input size. This is particularly helpful when debugging or validating the model structure.

```
1 if architecture in ("FFNN", "CNN", "LSTM", "Transformer") and extend_model is
   ↳ None:
2     if hasattr(model, "summary"):
3         model.summary()
4     else:
5         summary(model, input_size=(1, 724))
```

Listing 3.23: Printing the model summary

Here the visualization of the summary for the FFNN architecture:

```
1 =====
2 Layer (type:depth-idx)           Output Shape           Param #
3 =====
4 FFNN                             [1, 61]                --
5 Sequential: 1-1                  [1, 61]                --
6   Linear: 2-1                   [1, 543]               393,675
7   ReLU: 2-2                     [1, 543]               --
8   Linear: 2-3                   [1, 543]               295,392
```

```

9      ReLU: 2-4                      [1, 543]          --
10     Linear: 2-5                     [1, 543]          295,392
11     ReLU: 2-6                       [1, 543]          --
12     Linear: 2-7                     [1, 61]           33,184
13 =====
14 Total params: 1,017,643
15 Trainable params: 1,017,643
16 Non-trainable params: 0
17 Total mult-adds (Units.MEGABYTES): 1.02
18 =====
19 Input size (MB): 0.00
20 Forward/backward pass size (MB): 0.01
21 Params size (MB): 4.07
22 Estimated Total Size (MB): 4.09
23 =====

```

Listing 3.24: Printing the model summary of FFNN

And here is the summary of LSTM:

```

1 =====
2 Layer (type:depth-idx)              Output Shape          Param #
3 =====
4      ↪
4 Embedding: 1-1                      [1, 1000, 64]         3,584
5 LSTM: 1-2                           [1, 1000, 500]        1,132,000
6 Linear: 1-3                         [1, 61]               30,561
7 =====
8 Total params: 1,166,145
9 Trainable params: 1,166,145
10 Non-trainable params: 0
11 Total mult-adds (G): 1.13
12 =====
13 Input size (MB): 0.00
14 Forward/backward pass size (MB): 4.51
15 Params size (MB): 4.66
16 Estimated Total Size (MB): 9.18
17 =====

```

Listing 3.25: Printing the model summary of LSTM

Saving model Originally, models were saved using the `.h5` extension. However, for the PyTorch-based models, it was necessary to switch to the `.pth` format. To avoid overwriting existing files, the function `save_model` (see appendix lines from 1218 to 1294) dynamically appends an incremental number to the file name. Moreover, if a specific name is provided via `--model_name` (see section 2.3).

Another important modification involved the `torch.save` function. Unlike Keras' `model.save`, which stores both the weights and the architecture (such as input size, output size, and hidden layer configuration), PyTorch's `torch.save` saves only the model's weights via the `state_dict()`. To reproduce the same behavior, it was necessary to manually include the structural parameters

when saving the model. This was achieved by using the FFNN class of PyTorch (see section 3.3.2) to pass all required values.

```

1     if args.model_name == 'm.h5':
2         i = 1
3         base_name = args.model_name.split('.')[0]
4         extension = '.pth' if architecture == "FFNN" else '.h5'
5         while os.path.exists(os.path.join(args.save_directory, base_name +
6             ↪ str(i) + extension)):
7             i += 1
8             model_name = base_name + str(i) + extension
9     else:
10        model_name = args.model_name
11        if architecture == "FFNN":
12            model_name = model_name.replace('.h5', '.pth')

```

Listing 3.26: Saving the model with correct extension

```

1     if architecture in ("FFNN", "LSTM"):
2         state_dict = {
3             'model_state_dict': model.state_dict(),
4             'hidden_size': model.hidden_size,
5             'output_size': model.output_size,
6         }
7
8         if architecture == "FFNN":
9             state_dict['input_size'] = model.input_size
10            state_dict['num_hidden_layers'] = model.num_hidden_layers
11        elif architecture == "LSTM":
12            state_dict['vocab_size'] = model.vocab_size
13            state_dict['embed_dim'] = model.embed_dim
14            state_dict['num_layers'] = model.num_layers
15            state_dict['dropout'] = model.dropout
16
17        torch.save(state_dict, model_path)

```

Listing 3.27: Saving the model with correct extension

Input validation Last but not least, the input validation logic was modified. Previously, the script accepted only filenames with the .h5 extension, consistent with Keras models. However, as previously discussed, PyTorch models are conventionally saved using the .pth extension. Therefore, the validation step was updated to enforce this requirement, ensuring that the file extension matches the framework being used.

```

1 if os.path.splitext(args.model_name)[1] not in ('.h5', '.pth'):
2     print('ERROR: The model must have extension ".h5" (for Keras) or ".pth" (
3         ↪ for PyTorch).', file=sys.stderr)
4     sys.exit(1)

```

Listing 3.28: Newer input validation

3.4 Evaluation code of the model

After porting `train.py` to PyTorch and successfully training models, I proceeded to adapt the `eval.py` script accordingly, see pages 61 and 79. Unlike `train.py`, I chose to retain the original function flow already present in `eval.py`, modifying only the conditional statements responsible for selecting the correct evaluation logic for the specified architecture.

3.4.1 Benchmark function

The benchmark function was adapted to evaluate PyTorch models using tensors. The evaluation is now performed as seen in listing 3.29:

```

1 if architecture == ("FFNN", "LSTM"):
2     if hasattr(model, "evaluate"): # Keras model
3         results.append(model.evaluate(statistics, labels, batch_size=args.
4             ↪ batch_size, verbose=1))
5     else: # PyTorch model
6         x = torch.tensor(statistics.numpy(), dtype=torch.float32)
7         y = torch.tensor(labels.numpy(), dtype=torch.long)
8         with torch.no_grad():
9             outputs = model(x)
10            loss = F.cross_entropy(outputs, y)
11            top1 = torch.argmax(outputs, dim=1)
12            acc = (top1 == y).float().mean()
13
14            top3 = torch.topk(outputs, k=3, dim=1).indices
15            y_expanded = y.unsqueeze(1).expand_as(top3)
16            k3_acc = (top3 == y_expanded).any(dim=1).float().mean()

```

Listing 3.29: New evaluation for FFNN and LSTM architecture

As can be observed, the evaluation method is identical to the one carried out by the `predict` function in `train.py` (see section 3.3.4). This is because the evaluations must be consistent for both cases, meaning that they should yield more or less the same result.

3.4.2 Model loading

The `load_model` function was updated to support PyTorch `.pth` models. Parameters previously saved using `torch.save` are now correctly reloaded with `torch.load`. listing 3.30 allow the checkpoint saved on disk with `torch.save()` to be restored, and depending on the chosen architecture (FFNN or LSTM), the appropriate class is initialized by passing the stored parameters (input dimensions, hidden size, number of layers, etc.). Subsequently, the model weights are loaded with `model.load_state_dict(...)` and the model is set to `eval()` mode in order to be used exclusively for prediction, without further training.

```

1 if architecture in ("FFNN", "LSTM") and model_path.endswith(".pth"):
2     if architecture == "FFNN":
3         from cipherTypeDetection.train import TorchFFNN
4     elif architecture == "LSTM":
5         from cipherTypeDetection.train import TorchLSTM
6
7     checkpoint = torch.load(model_path, map_location=torch.device("cpu"))
8
9     if architecture == "FFNN":
10         model = TorchFFNN(
11             input_size=checkpoint['input_size'],
12             hidden_size=checkpoint['hidden_size'],
13             output_size=checkpoint['output_size'],
14             num_hidden_layers=checkpoint['num_hidden_layers']
15         )
16     elif architecture == "LSTM":
17         model = LSTM(
18             vocab_size=checkpoint['vocab_size'],
19             embed_dim=checkpoint['embed_dim'],
20             hidden_size=checkpoint['hidden_size'],
21             output_size=checkpoint['output_size'],
22             num_layers=checkpoint['num_layers'],
23             dropout=checkpoint['dropout']
24         )
25
26     model.load_state_dict(checkpoint['model_state_dict'])
27     model.eval()
28
29     config.FEATURE_ENGINEERING = (architecture == "FFNN")
30     config.PAD_INPUT = (architecture == "LSTM")
31
32     return model

```

Listing 3.30: Setting model parameters

In addition, in this way, if an FFNN is loaded, `FEATURE_ENGINEERING=True` and `PAD_INPUT=False` are set (see section 2.4 and section 2.3).

Conversely, if an LSTM is loaded, `FEATURE_ENGINEERING=False` and `PAD_INPUT=True`, so that the model receives tokenized ciphertexts instead of numerical features.

In addition, I changed the previous logic that always evaluated models as ensembles. Now, the model is only wrapped in a `RotorDifferentiationEnsemble` when rotor ciphers are actually present. This modification was introduced to avoid unnecessary overhead when evaluating models that do not involve rotor ciphers, ensuring that the ensemble mechanism is only applied when it is actually required:

```

1 has_rotor_ciphers = any(c in config.ROTOR_CIPHER_TYPES for c in cipher_types)
2
3 if has_rotor_ciphers:
4     rotor_only_model_path = args.rotor_only_model
5     if not os.path.exists(rotor_only_model_path):
6         raise FileNotFoundError(f"Rotor-only model is required but not found
           ↳ at {rotor_only_model_path}")

```

```

7     with open(rotor_only_model_path, "rb") as f:
8         rotor_only_model = pickle.load(f)
9     return RotorDifferentiationEnsemble(architecture, model, rotor_only_model
    ↪ )
10
11 return model

```

Listing 3.31: Changes to Ensemble architecture

3.4.3 Input validation

Finally, the input validation in the main function was extended to support models saved with the .pth extension:

```

1 if os.path.splitext(args.model)[1] not in ('.h5', '.pth'):
2     print('ERROR: The model must have extension ".h5" (for Keras) or ".pth" (
    ↪ for PyTorch FFNN and LSTM).', file=sys.stderr)
3     sys.exit(1)

```

Listing 3.32: Newer input validation

4 Evaluation and results

As I have already discussed in section 3.4, the final part of this work consists of the evaluation phase, that is, testing the model with new data in order to obtain meaningful results such as accuracy. In section 3.4, the goal was to correctly implement the code that enables this phase, whereas in this chapter we will compare the obtained results, aiming to contrast the Keras and PyTorch implementations.

Before analyzing the content of this chapter, I need to clarify the purpose of these evaluations. The goal of this thesis is not to improve the performance of the app; therefore, as we will see, I will generally aim to replicate the results achieved by the previous Keras implementation. This is because I do not have the exhaustive knowledge required to enhance a learning model like NCID, but I have worked to make the two implementations as similar as possible.

After implementing our architectures, as discussed in the previous chapter, it is now time to test the models with large amounts of data. Until now, I had been testing all modifications locally with parameters that allowed for simple execution, mainly to highlight potential issues in the data flow. However, these runs had no significance in terms of the actual evaluation of the model and the study of its accuracy.

To perform more meaningful evaluations, I was able to use a machine at the Universität der Bundeswehr München, namely the NVIDIA DGX H100, with the assistance of Doris Behrendt and Maik Bastian. For details on the specifications of the machine used, you can consult the following website [16]. A small clarification: I did not use the full power of the DGX machine; I was only able to train the model on a single GPU.

In order to run the modifications I implemented, I had to create Dockerfiles which, once unpacked on the cluster, allowed the execution of the code from my personal GitHub repository. A Dockerfile is a simple text file that contains the instructions to build a Docker container, that is, an isolated environment in which to run an application in exactly the same way on any machine.

In listing 4.1 depicts the command line used to run the dockerfile previously built. I won't share the build command since it contains reserved informations.

```
1 docker run -v $(pwd)/output:/app/ncid/cipherTypeDetection/output -it --gpus
   ↪ all --name ncid-keras-benchmark ncid-keras-benchmark
```

Listing 4.1: Docker run command

That said, let's start analyzing the results obtained.

4.1 First runs

Result from Keras Initially, I used the Keras implementation to perform a small training run with the following input data (see section 2.3):

- `train_dataset_size = 1000`
- `batch_size = 256`
- `dataset_workers = 4`
- `min_train_len = 100`
- `max_train_len = 1000`
- `min_test_len = 100`
- `max_test_len = 1000`
- `max_iter = 10000`

The results training the FFNN were the following:

Accuracy: 0.1984 k3-accuracy: 0.3783

Here, **accuracy** indicates the probability that the cipher was correctly identified, and **k3-accuracy** represents the probability that the correct cipher was among the top 3 most probable predictions.

So now we have a clear goal: to get as close as possible to the results of the Keras implementation.

Result from PyTorch The code has undergone several modifications over the months; indeed, the first training run with the same inputs yielded these results.

Accuracy: 0.042523, k3-accuracy: 0.075315

Obviously, we were still far from a good implementation. For example, at the beginning I added an unnecessary data normalization step. This happened because I initially thought that the architecture class did not receive already normalized inputs. In fact, as we have seen in section 3.2, the data were already normalized, and therefore performing this operation a second time significantly worsened the model's accuracy. In listing 4.2 depicts the redundant normalization step that was added.

```
1 mean = stats_np.mean(axis=0)
2 std = stats_np.std(axis=0) + 1e-8
3 stats_np = (stats_np - mean) / std
```

Listing 4.2: Redundant normalization

So, after discussing it with Maik Bastian, I removed the redundant normalization. Then I obtained the following results:

Accuracy: 0.076720, k3-accuracy: 0.151515

As can be seen, both the single accuracy and the k3 accuracy improved, although only slightly. This is where the modification I mentioned in the implementation chapter, which Maik Bastian pointed out to me (see section 3.3.3), comes into play. After fixing the implementation as discussed earlier, I finally obtained results similar to those of Keras:

Accuracy: 0.165479, k3-accuracy: 0.364742

Target acquired!

4.2 DGX runs

It was then time to test our models on a large scale using the DGX machine. Below are the command line hyperparameters:

- `train_dataset_size = 976`
- `batch_size = 64`
- `dataset_workers = 16`
- `min_train_len = 100`
- `max_train_len = 1000`
- `min_test_len = 100`
- `max_test_len = 1000`
- `max_iter = 10000000`

To compare the two implementations, I prepared two shell scripts. The one shown in listing 4.3 is the script used on the DGX machine to test the PyTorch implementation. Of course, I also created a similar one for Keras that I won't add here:

```

1 BASE_DIR=output
2 mkdir -p "$BASE_DIR"
3
4 DATE=$(date +%Y%m%d")
5
6 COMMON_ARGS="--download_dataset=False \
7   --plaintext_input_directory=./data/gutenberg_en \
8   --rotor_input_directory=./data/rotor_ciphertexts \
9   --train_dataset_size=976 \
10  --dataset_workers=16 \
11  --batch_size=64 \
12  --max_iter=10000000 \
13  --min_train_len=100 \
14  --max_train_len=1000 \
15  --min_test_len=100 \
16  --max_test_len=1000 \
17  --epochs=1 \
18  --ciphers=all"
19
20 run_benchmark () {
21   ARCH=$1
22   for i in {1..3}; do
23     RUN_DIR="$BASE_DIR/${ARCH}_run_$i"
24     mkdir -p "$RUN_DIR"
25     echo "Launching $ARCH run $i..."
26     python train.py \
27       --architecture=$ARCH \
28       $COMMON_ARGS \
29       --save_directory="$RUN_DIR" \
30       --model_name="${ARCH}_run_$i.pth" \
31       > "$RUN_DIR/${ARCH}_var_10000000_run_${i}_${DATE}.txt" \
32       2> "$RUN_DIR/err_${ARCH}_var_10000000_run_${i}_${DATE}.txt"
33   done
34 }
35
36 # FFNN
37 run_benchmark FFNN
38
39 # LSTM
40 run_benchmark LSTM

```

Listing 4.3: Shell script for PyTorch

The goal is to test the FFNN architecture implementation three times and the LSTM implementation three times, for both frameworks. This process was carried out using the machine reserved for these tests, in order to ensure the most accurate results possible.

Several tests were performed to verify whether the implementation was correct, whether the models were being saved properly, and to address other potential issues. From this point on, I will report only the results of one result for both Keras and PyTorch. A quick note: the PyTorch results date back to the test performed on August 8, while the Keras results refer to the test conducted on October 6. The Keras results were repeated later because some anomalies had

been found in the model saving process, which prevented a correct evaluation.

4.2.1 Results from Keras

Once again, we will use the results obtained as a benchmark to assess whether we have successfully implemented our architectures. Moreover, it is fundamental introducing a the concept of *Precision*, *Recall* and *F1-score*. Precision indicates the percentage of correct positive predictions, that is, how often the model is right when it predicts a positive class. Recall measures the model's ability to identify all actual positive instances, that is, how many of the real positive classes have been correctly recognized. Precision and recall differ in focus: precision emphasizes the accuracy of positive predictions, while recall emphasizes the completeness of capturing all positive instances. The F1-score is a metric used to evaluate the performance of a classification model, especially in cases where the data are imbalanced, meaning that some classes have significantly more samples than others.

$$\text{Precision} = \frac{TP}{TP + FP} \quad (4.1)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (4.2)$$

where TP stands for true positive, FP for false positive and FN for false negative. All three variables (TP , FP and FN) are non-negative integer counts that represent the number of samples in each category. The F1 score is defined as the harmonic mean between precision and recall:

$$F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4.3)$$

This value ranges from 0 to 1, where 1 indicates perfect performance (both precision and recall are equals to 1), and 0 indicates the worst performance.

Now we will discuss the numerical results of the training sessions. To provide a direct reference, I include a list of the ciphers in their order of appearance, sorted alphabetically. This is the same order that appears when we obtain the results during the model evaluation.

1. amsco	2. autokey	3. baconian	4. bazeries	5. beaufort	6. bifid
7. cadenus	8. checkerboard	9. columnar_transposition	10. condi	11. cmbifid	12. digrafid
13. foursquare	14. fractionated_morse	15. grandpre	16. grille	17. gromark	18. gronsfeld
19. headlines	20. homophonic	21. key_phrase	22. monome_dinome	23. morbit	24. myszkowski
25. nicodemus	26. nihilist_substitution	27. nihilist_transposition	28. null	29. numbered_key	30. periodic_gromark
31. phillips	32. phillips_rc	33. plaintext	34. playfair	35. pollux	36. porta
37. portax	38. progressive_key	39. quagmire1	40. quagmire2	41. quagmire3	42. quagmire4
43. ragbaby	44. railfence	45. redefence	46. route_transposition	47. running_key	48. seriated_playfair
49. slidefair	50. swagman	51. tridigital	52. trifid	53. tri_square	54. two_square
55. variant	56. vigenere	57. enigma	58. m209	59. purple	60. sigaba
61. typex					

If you wish to skip directly to the comparison, see section [4.3](#).

FFNN results With the input hyperparameter specified in section 4.2, we obtained the following results for the Keras FFNN:

Total accuracy: 0.7427164374276443

Total k3: 0.9042634659180658

```
precision: [0.84247601 0.70099323 1.          0.99434685 0.23120063 0.76061568
0.87966872 0.95546632 0.93239781 0.88377353 0.80099023 0.77054351
0.93703211 0.99554579 0.94213309 0.9879969 0.89974178 0.7988629
0.99993887 0.8506699 0.99640551 0.99546737 0.92647763 0.63775634
0.85385154 0.99027014 0.73630802 1.          0.99939062 0.97517683
0.86251165 0.7487211 0.99750623 0.96016622 0.98180602 0.47141996
0.99668019 0.39556921 0.16037532 0.20766575 0.22230205 0.19764465
1.          0.41750766 0.43044917 0.99993903 0.49032965 0.84550003
0.48359369 0.99441375 0.77116174 0.8177442 0.95524203 0.89872368
0.2066785 0.20256524 0.31359946 0.46392964 0.34882021 0.29313037
0.32087629]
```

```
recall: [0.92615854 0.78323171 1.          0.99743902 0.23378049 0.86176829
0.66605189 0.9720122 0.97640244 0.73859756 0.75957317 0.83420732
0.81664634 0.99487805 0.99969512 0.93353659 0.97731707 0.97670732
0.99743902 0.98335366 0.9972561 0.7097561 0.94432927 0.4797561
0.96719512 0.99914634 0.92469512 1.          1.          0.89109756
0.73329268 0.89243902 1.          0.98621951 0.895 0.71560976
0.98853659 0.73817073 0.09067073 0.18335366 0.09432927 0.14121951
0.99963415 0.34054878 0.57908537 1.          0.76365854 0.8552439
0.75128049 0.97689024 0.87103659 0.74579268 0.95780488 0.85871951
0.11359756 0.22341463 0.16957317 0.6979878 0.20371951 0.2747561
0.12146341]
```

```
f1: [0.88233756 0.73983412 1.          0.99589054 0.2324834 0.80803865
0.7580995 0.96366824 0.95389289 0.80469009 0.7797321 0.80111258
0.87270713 0.99521181 0.97006094 0.95999498 0.9369264 0.87887849
0.99868738 0.91221223 0.99683062 0.82867618 0.93531828 0.54758673
0.90699603 0.99468844 0.81981836 1.          0.99969521 0.93124323
0.79267047 0.8142873 0.99875156 0.9730185 0.93639553 0.56839811
0.99259169 0.5151051 0.11584606 0.19475389 0.1324543 0.16473433
0.99981704 0.37512174 0.49382523 0.99996951 0.59720566 0.85034405
0.58842352 0.98557411 0.81806208 0.78011289 0.95652174 0.87826629
0.1466121 0.2124797 0.22012031 0.55738423 0.25721765 0.28364598
0.17622081]
```

LSTM results In this paragraph the Keras LSTM results are shown:

Total accuracy: 0.699571922430185

Total k3: 0.8750897243388391

```
precision: [0.87641498 0.65051067 0.99993903 0.97511339 0.23158169 0.62077346
0.73410459 0.84212268 0.6456672 0.32411558 0.62267438 0.79025262
0.92058659 0.97330374 0.98561677 0.99636811 0.99944673 0.85276796
1. 0.99115871 0.94371099 0.99415133 0.98101886 0.63922907
0.66583541 0.9971741 0.65088876 1. 1. 0.99129173
0.60972544 0.62101725 0.9995122 0.62416934 0.97185539 0.61542203
0.93694264 0.30687831 0.16879706 0.22280193 0.20307135 0.18400538
0.99939036 0.40662154 0.37781341 0.99993902 0.53409208 0.6180449
0.7126617 0.96941424 0.96679606 0.76808586 0.77763496 0.76023877
0.1938861 0.21226216 0.25221621 0.42009865 0.21202813 0.21203855
0.24711656]
```

```
recall: [0.91585366 0.65634146 1. 0.99628049 0.21658537 0.5295122
0.47924598 0.77603659 0.74054878 0.36591463 0.54487805 0.75536585
0.90335366 0.98926829 0.98609756 0.98695122 0.99134146 0.97121951
0.99695122 0.97067073 0.95890244 0.98463415 0.97695122 0.38829268
0.71634146 0.9897561 0.98689024 1. 0.99987805 0.9995122
0.76914634 0.51817073 0.9995122 0.86481707 0.99170732 0.84871951
0.9929878 0.51280488 0.05878049 0.14060976 0.12981707 0.08347561
0.99957317 0.43585366 0.37359756 0.99987805 0.6295122 0.5404878
0.88682927 0.95664634 0.98713415 0.78981707 0.70091463 0.83871951
0.33878049 0.24487805 0.19603659 0.65432927 0.15993902 0.07109756
0.18420732]
```

```
f1: [0.8957004 0.65341306 0.99996951 0.9855833 0.22383263 0.57152259
0.57990935 0.80773014 0.68986083 0.3437491 0.58118435 0.77241551
0.91188871 0.98122108 0.98585711 0.99163731 0.9953776 0.90814756
0.99847328 0.98080774 0.95124607 0.98936985 0.97898081 0.48311964
0.69016567 0.99345125 0.78442301 1. 0.99993902 0.99538499
0.68022002 0.56495147 0.9995122 0.72504665 0.981681 0.71348387
0.96415145 0.38397443 0.08719642 0.17241121 0.15838417 0.11484899
0.99948175 0.42073045 0.37569366 0.99990853 0.57788973 0.57667035
0.79026299 0.96298797 0.97685925 0.7787999 0.73728433 0.79755313
0.24662642 0.22740657 0.22060589 0.51168224 0.18233638 0.10648888
0.21107424]
```

4.2.2 Results from PyTorch

Now we present the PyTorch results, after which we will draw some conclusions.

FFNN results Starting with the FFNN:

Total accuracy: 0.7313708694494867

Total k3: 0.9020696031909635

```
precision: [0.77108553 0.64081574 1.          0.99379977 0.22407726 0.77174341
0.94204352 0.99584505 0.88031958 0.75089452 0.83762434 0.80873805
0.89755428 0.9911354  0.92796322 0.99900431 0.88510873 0.80968893
0.99987772 0.64407725 0.9967045  0.90825141 0.86191933 0.46486261
0.95426792 0.98740129 0.69856062 1.          0.99266889 0.97904397
0.76588775 0.86504006 0.99350689 0.96091244 1.          0.31089329
0.99572101 0.3690516  0.16369407 0.19653499 0.20846609 0.16659684
1.          0.35294118 0.43524855 1.          0.57956381 0.83109792
0.67200429 0.99152857 0.86308309 0.79863337 0.9533808  0.82372935
0.19591978 0.19693404 0.32494279 0.53625992 0.43583976 0.27847226
0.31505273]
```

```
recall: [0.9486084 0.8170166 1.          0.99786377 0.16711426 0.83782959
0.57701757 0.9508667  0.98858643 0.90942383 0.69897461 0.78973389
0.87805176 0.99633789 0.9977417  0.9185791  0.98132324 0.98339844
0.99816895 0.91595459 0.99682617 0.93048096 0.94561768 0.47601318
0.95391846 0.99975586 0.97454834 1.          1.          0.87255859
0.85620117 0.72491455 0.99926758 0.98730469 0.49353027 0.73370361
0.99420166 0.80847168 0.06610107 0.19525146 0.0838623  0.14562988
0.99981689 0.02783203 0.70648193 1.          0.70068359 0.85473633
0.53533936 0.97869873 0.74871826 0.80609131 0.95611572 0.94665527
0.1729126  0.18896484 0.09533691 0.65985107 0.24835205 0.29504395
0.10028076]
```

```
f1: [0.85068418 0.71826792 1.          0.99582762 0.19144845 0.80342981
0.71567321 0.97283627 0.93131702 0.82259089 0.76204418 0.799123
0.88769592 0.99372984 0.96158824 0.95710515 0.93073606 0.88812943
0.9990226  0.75632497 0.99676533 0.91923181 0.90183067 0.47037182
0.95409316 0.99354017 0.8137917  1.          0.99632096 0.9227393
0.80853026 0.78880255 0.99637891 0.9739298  0.66089089 0.43673025
0.99496075 0.50677175 0.09417391 0.19589112 0.11960827 0.15540937
0.99990844 0.05159538 0.53864768 1.          0.63439434 0.8427514
0.59593695 0.98507188 0.80184332 0.80234501 0.95474631 0.88092466
0.18369861 0.19286715 0.14742107 0.59167032 0.31640747 0.28651869
0.15213667]
```

LSTM results

Total accuracy: 0.7012079733647427

Total k3: 0.8742761016214723

```
precision: [0.85796652 0.65652963 0.99987794 0.99353659 0.19896735 0.68323991
0.56204133 0.85257582 0.8650025  0.27295937 0.64737512 0.83648348
0.85027119 0.98911391 0.99039579 0.99307159 0.99993889 0.90409117
0.99993883 0.99721172 0.91260593 0.99652397 0.98542363 0.71317752]
```

```

0.48079802 0.99211117 0.58891286 0.99987794 0.99993897 0.99871982
0.67685637 0.63413706 0.99987786 0.8504757 0.98306821 0.68676212
0.97486797 0.32358401 0.13672073 0.19156903 0.18345945 0.18534671
1. 0.35910812 0.35320563 0.99993895 0.52338435 0.66906094
0.76714031 0.92055668 0.99588022 0.71154232 0.72936094 0.81924149
0.20497916 0.19357802 0.23506053 0.4480726 0.22592295 0.201594
0.28153913]

```

```

recall: [0.94458008 0.66278076 1. 0.99450684 0.16699219 0.51330566
0.41035906 0.82525635 0.52874756 0.45025635 0.55847168 0.6428833
0.94726562 0.99267578 0.98815918 0.99731445 0.99865723 0.94415283
0.99768066 0.9822998 0.94647217 0.99737549 0.99029541 0.43768311
0.84283447 0.99786377 0.88311768 1. 1. 0.99993896
0.76220703 0.60998535 0.99926758 0.80749512 0.99578857 0.83435059
0.98016357 0.43621826 0.06585693 0.19360352 0.08435059 0.02593994
0.99981689 0.52099609 0.33087158 0.99969482 0.60107422 0.69360352
0.86663818 0.96893311 0.98852539 0.8661499 0.8449707 0.87939453
0.22814941 0.32159424 0.22399902 0.58459473 0.13745117 0.07873535
0.19873047]

```

```

f1: [0.89919238 0.65964038 0.99993897 0.99402147 0.18158288 0.58620569
0.47436993 0.83869367 0.65631274 0.3398756 0.59964611 0.72701546
0.89615151 0.99089164 0.98927622 0.9951885 0.99929765 0.92368782
0.99880847 0.9896996 0.92923058 0.99694955 0.98785351 0.54245622
0.61230462 0.99497916 0.7066149 0.99993897 0.99996948 0.99932902
0.71700063 0.62182678 0.99957262 0.8284283 0.98938751 0.75339635
0.9775086 0.3715526 0.08889438 0.1925809 0.11556633 0.04551052
0.99990844 0.42516312 0.34167402 0.99981687 0.55954545 0.68111121
0.81385951 0.94412561 0.99218918 0.78127065 0.78292097 0.84825292
0.21594454 0.24168062 0.22939651 0.50730932 0.17091682 0.11324233
0.23299581]

```

4.3 Comparison

The reported evaluations show very satisfactory results regarding the achievement of our goal, which, I would like to recall once again, was not to improve our previous implementation in terms of model accuracy, but rather to obtain results with PyTorch that were consistent with those achieved using Keras. At first glance, the two implementations behave in a very similar way, as can be seen in table 4.1 and table 4.2. It should also be reminded that both implementations were trained with identical hyperparameters, in order to ensure a fair evaluation.

Starting from the FFNN results, the Keras model achieved a total accuracy of 74.27% and a top-3 accuracy of 90.04%. The same architecture implemented in PyTorch reached a total accuracy of 73.13% and a top-3 accuracy of 90.02%. The total accuracy result is perfectly in line with our expectations, while the total accuracy is slightly lower. These discrepancies are not concerning, but they often result from differences between the two frameworks, which use slightly different

	Keras	PyTorch
FFNN	74.27%	73.13%
LSTM	69.95%	70.12%

Table 4.1: Table of total accuracies

	Keras	PyTorch
FFNN	90.04%	90.02%
LSTM	87.50%	87.42%

Table 4.2: Table of top 3 accuracies

weight initialization or floating-point rounding mechanisms, leading to minor variations in the results.

Similarly, for the LSTM architecture, the Keras implementation achieved a total accuracy of 69.95% and a top-3 accuracy of 87.50%, while the PyTorch version obtained a total accuracy of 70.12% and a top-3 accuracy of 87.42%. As with the FFNN, these small discrepancies are due to operational differences between the frameworks.

Let's now take a look to the evaluation metrics. A more detailed analysis of these highlights that both implementations perform consistently across most classes, with minor variations in Precision, Recall, and F1-score. Overall, the Keras model shows slightly higher values for these metrics.

The Precision values indicate that both models are generally reliable when predicting a positive class, although Keras tends to produce fewer false positives in several cipher categories. The Recall values are also comparable, suggesting that both models are capable of identifying the majority of the true instances. However, in a few specific classes, PyTorch shows lower Recall, meaning it occasionally fails to recognize some correct samples that Keras is able to detect.

The F1-score, which combines Precision and Recall into a single measure, reflects this balance. The Keras implementation achieves slightly higher F1-scores overall, confirming its better equilibrium between accurate and complete predictions.

5 Conclusion

As we have seen, the project aimed at the development and evaluation of FFNN and LSTM architectures, now familiar to you, for the automatic classification of classical ciphers, with particular attention to the comparison between implementations in Keras and PyTorch.

We started from the NCID application, analyzing what it is, how it works, its connections with cryptography, and the technologies it relies on. This provided the necessary foundation for understanding the work carried out. We began by reviewing the fundamentals of Machine Learning, including the concept of statistical features, and comparing the two frameworks Keras and PyTorch.

All of this served as preparation for the core part of the project: the implementation. In this phase, we analyzed the data flow within the NCID app and examined each stage: the extraction of plaintexts and their encryption, the preprocessing of texts and the calculation of statistical features, the training and validation of the models, and finally their evaluation. We highlighted the technical differences between the two frameworks and solved some technical issues. The most significant issues encountered were the incorrect implementation of data normalization and the improper optimization of gradients during model training, which, however, were resolved as discussed in chapter 3.

We concluded the work by comparing the results obtained from the Keras and PyTorch implementations. As already mentioned, the goal was not to outperform the previous implementation, but to replicate it accurately, thus providing a solid starting point for anyone working on this project. As shown in chapter 4, this goal was successfully achieved for both the FFNN and LSTM architectures.

Looking ahead, the work could be extended by implementing in PyTorch the remaining architectures, such as Transformer and Naive Bayes, as well as by making more in-depth modifications to the training and evaluation cycles. I hope that the project has provided a concrete contribution to the development of the NCID app, offering both significant experimental results and a useful foundation for future developments.

6 Acknowledgments

For the realization of this thesis, I would like to sincerely thank the entire CrypTool team for their support in the development of this work and for guiding me throughout the project. In particular, I would like to thank Dr. Doris Behrendt and Maik Bastian, who were always available to answer my questions and provided valuable guidance with their advice throughout the entire process. I would also like to thank Professor Bernhard Esslinger for giving me the opportunity to work on this project, which I truly appreciated and in which I fully immersed myself, and that has been fundamental in shaping my choice of master's degree.

Stefano Sala, October 13, 2025

7 Bibliography

- [1] Ernst Leierzopf. *NCID*. 2021. URL: <https://github.com/ernstleierzopf/ncid>.
- [2] Nils Kopal. “Of Ciphers and Neurons – Detecting the Type of Ciphers Using Artificial Neural Networks”. In: *Proceedings of the 3rd International Conference on Historical Cryptology (HistoCrypt 2020)*. Vol. 171. Linköping University Electronic Press, May 2020, pp. 77–86. DOI: [10.3384/ecp2020171011](https://doi.org/10.3384/ecp2020171011). URL: https://ep.liu.se/en/conference-article.aspx?series=ecp&issue=171&Article_No=11.
- [3] Brooke Dalton and Mark Stamp. *Classifying World War II Era Ciphers with Machine Learning*. 2023. arXiv: [2307.00501](https://arxiv.org/abs/2307.00501) [cs.LG]. URL: <https://arxiv.org/abs/2307.00501>.
- [4] Tom M. Mitchell. *Machine Learning*. Vol. 1. 9. McGraw-hill New York, 1997.
- [5] Ian Goodfellow, Yoshua Bengio, and Aaron Courville. *Deep Learning*. MIT Press, 2016. URL: <http://www.deeplearningbook.org>.
- [6] Ernst Leierzopf et al. “A Massive Machine-Learning Approach For Classical Cipher Type Detection Using Feature Engineering”. English. In: *Proceedings of the 4th International Conference on Historical Cryptology*. Aug. 2021, pp. 111–120. DOI: [10.3384/ecp183164](https://doi.org/10.3384/ecp183164).
- [7] Tariq Rashid. *Make Your Own Neural Network*. 1st. North Charleston, SC, USA: CreateSpace Independent Publishing Platform, 2016. ISBN: 1530826608.
- [8] Sepp Hochreiter and Jürgen Schmidhuber. “Long short-term memory”. In: *Neural computation* 9.8 (1997), pp. 1735–1780.
- [9] Amrit Singh Bist. *Comprehensive Understanding of LSTM Architecture: A Deep Dive into Long Short-Term Memory Networks*. Accessed: 2025-07-17. 2023. URL: <https://medium.com/@amritsinghbist/comprehensive-understanding-of-lstm-architecture-a-deep-dive-into-long-short-term-memory-networks-5405f97afc6b>.
- [10] PyTorch Contributors. *torch.nn.CrossEntropyLoss*. 2025. URL: <https://pytorch.org/docs/stable/generated/torch.nn.CrossEntropyLoss.html>.
- [11] Keras Team. *Keras API Reference*. 2024. URL: <https://keras.io/api/>.
- [12] PyTorch Team. *PyTorch: An Open Source Machine Learning Framework*. 2024. URL: <https://pytorch.org/>.
- [13] Luca Pietro Giovanni Antiga, Eli Stevens, and Thomas Viehmann. *Deep learning with PyTorch*. Simon and Schuster, 2020.
- [14] Pavan Patel. *From Keras to PyTorch*. 2020. URL: <https://medium.com/analytics-vidhya/from-keras-to-pytorch-722fa3b65cce>.

- [15] Ernst Leierzopf et al. “Detection of Classical Cipher Types with Feature-Learning Approaches”. In: *Data Mining*. Ed. by Yue Xu et al. Singapore: Springer Singapore, Dec. 2021, pp. 152–164. ISBN: 978-981-16-8531-6. DOI: [10.1007/978-981-16-8531-6_11](https://doi.org/10.1007/978-981-16-8531-6_11). URL: https://link.springer.com/chapter/10.1007/978-981-16-8531-6_11.
- [16] NVIDIA Corporation. *Introduction to NVIDIA DGX H100/H200*. 2025. URL: <https://docs.nvidia.com/dgx/dgxm100-user-guide/introduction-to-dgxm100.html>.

8 Appendix: Full Python code

Here is listed all the code I wrote during the internship.

8.1 Full listing of train.py

```
1 import multiprocessing
2 from pathlib import Path
3
4 import argparse
5 import sys
6 import time
7 import shutil
8 from sklearn.model_selection import train_test_split
9 import os
10 import math
11 import pickle
12 import functools
13
14 # PyTorch
15 import torch
16 import torch.nn as nn
17 import torch.optim as optim
18 from torchinfo import summary
19 import numpy as np
20 from torch.utils.data import TensorDataset, DataLoader
21
22 from sklearn.pipeline import Pipeline
23 from sklearn.preprocessing import StandardScaler
24 from sklearn.tree import DecisionTreeClassifier, plot_tree
25 from sklearn.ensemble import RandomForestClassifier, ExtraTreesClassifier
26 from sklearn.naive_bayes import MultinomialNB
27 from sklearn.svm import SVC
28 from sklearn.neighbors import KNeighborsClassifier
29 import matplotlib.pyplot as plt
30 from datetime import datetime
31
32 # This environ variable must be set before all tensorflow imports!
33 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
34 import tensorflow as tf
35 from tensorflow.keras.metrics import SparseTopKCategoricalAccuracy
36 from tensorflow.keras.optimizers import Adam # , Adamax
37 import tensorflow_datasets as tfds
38 sys.path.append("../")
39 from cipherTypeDetection.nullDistributionStrategy import NullDistributionStrategy
40 import cipherTypeDetection.config as config
41 from cipherTypeDetection.trainingBatch import TrainingBatch
42 from cipherTypeDetection.cipherStatisticsDataset import RotorCiphertextsDatasetParameters, PlaintextPathsDatasetParameters,
43     CipherStatisticsDataset
44 from cipherTypeDetection.predictionPerformanceMetrics import PredictionPerformanceMetrics
45 from cipherTypeDetection.miniBatchEarlyStoppingCallback import MiniBatchEarlyStopping
46 from cipherTypeDetection.transformer import TransformerBlock, TokenAndPositionEmbedding
47 from cipherTypeDetection.learningRateSchedulers import TimeBasedDecayLearningRateScheduler, CustomStepDecayLearningRateScheduler
48 # always flush after print as some architectures like RF need very long time before printing anything.
49 print = functools.partial(print, flush=True)
50 for device in tf.config.list_physical_devices('GPU'):
51     tf.config.experimental.set_memory_growth(device, True)
52
53 class FFNN(nn.Module):
54     def __init__(self, input_size, hidden_size, output_size, num_hidden_layers):
55         super().__init__()
56
57         # saves parameters so that they can be saved and loaded later
58         self.input_size = input_size
59         self.hidden_size = hidden_size
60         self.output_size = output_size
61         self.num_hidden_layers = num_hidden_layers
62
```

```

63     layers = [nn.Linear(input_size, hidden_size), nn.ReLU()]
64     for _ in range(num_hidden_layers - 1):
65         layers += [nn.Linear(hidden_size, hidden_size), nn.ReLU()]
66     layers.append(nn.Linear(hidden_size, output_size))
67     self.net = nn.Sequential(*layers)
68
69     def forward(self, x):
70         return self.net(x)
71
72
73 class LSTM(nn.Module):
74     def __init__(self, vocab_size, embed_dim, hidden_size, output_size, num_layers=1, dropout=0.0):
75         super().__init__()
76
77         # saves parameters so that they can be saved and loaded later
78         self.vocab_size = vocab_size
79         self.embed_dim = embed_dim
80         self.hidden_size = hidden_size
81         self.output_size = output_size
82         self.num_layers = num_layers
83         self.dropout = dropout
84
85         # Layers
86         self.embedding = nn.Embedding(
87             num_embeddings=vocab_size,
88             embedding_dim=embed_dim,
89             padding_idx=0
90         )
91         self.lstm = nn.LSTM(
92             input_size=embed_dim,
93             hidden_size=hidden_size,
94             num_layers=num_layers,
95             batch_first=True,
96             dropout=dropout if num_layers > 1 else 0.0
97         )
98         self.fc = nn.Linear(hidden_size, output_size)
99
100     # B: Batch size - number of sequences processed in parallel
101     # L: Sequence length - number of time steps (tokens) in each sequence
102     # D: Embedding dimension - size of each 'tokens' embedding vector
103     # H: Hidden size - number of features in the LSTM hidden state
104     # C: Number of classes - dimensionality of the output logits
105
106     def forward(self, x):
107         # x: LongTensor of shape [B, L] or [B, L, 1]
108         if x.dim() == 3 and x.size(2) == 1:
109             x = x.squeeze(2) # remove channel dimension → [B, L]
110
111         emb = self.embedding(x) # embeddings → [B, L, D]
112
113         # LSTM returns:
114         # - output: hidden state at each time step → [B, L, H]
115         # - hidden: final hidden state for each layer → [num_layers, B, H]
116         # not used as we only need the last hidden state, but can be useful for debugging
117         output, (hidden, _) = self.lstm(emb)
118
119         # hidden[-1] selects the final hidden state of the top (last) layer
120         # at the last time step → [B, H]
121         last_hidden = hidden[-1]
122
123         # apply the fully-connected layer to get logits → [B, C]
124         logits = self.fc(last_hidden)
125
126         return logits
127
128
129
130 def train_torch_ffnn(model, args, train_ds):
131     device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
132     model.to(device)
133
134     optimizer = optim.Adam(
135         model.parameters(),
136         lr=config.learning_rate,
137         betas=(config.beta_1, config.beta_2),
138         eps=config.epsilon,
139         amsgrad=config.amsgrad
140     )
141     criterion = nn.CrossEntropyLoss()
142     model.train()
143
144     best_val_acc = 0
145     patience_counter = 0
146     patience_limit = 250
147
148     train_iter = 0
149     train_epoch = 0
150     start_time = time.time()
151

```

```

152 val_data_created = False
153 x_val = y_val = None
154
155 for epoch in range(args.epochs):
156     while train_ds.iteration < args.max_iter:
157         training_batches = next(train_ds)
158         for training_batch in training_batches:
159             statistics, labels = training_batch.items()
160             stats_np = statistics.numpy()
161             labels_np = labels.numpy()
162
163             if not val_data_created:
164                 x_train_np, x_val_np, y_train_np, y_val_np = train_test_split(
165                     stats_np, labels_np, test_size=0.3
166                 )
167                 x_val = torch.tensor(x_val_np, dtype=torch.float32).to(device)
168                 y_val = torch.tensor(y_val_np, dtype=torch.long).to(device)
169                 val_data_created = True
170             else:
171                 x_train_np = stats_np
172                 y_train_np = labels_np
173
174             # Use DataLoader for creating minibatch
175             x_train = torch.tensor(x_train_np, dtype=torch.float32)
176             y_train = torch.tensor(y_train_np, dtype=torch.long)
177
178             train_dataset = TensorDataset(x_train, y_train)
179             train_loader = DataLoader(train_dataset, batch_size=args.batch_size, shuffle=True)
180
181             batch_losses = []
182
183             for x_batch, y_batch in train_loader:
184                 x_batch = x_batch.to(device)
185                 y_batch = y_batch.to(device)
186
187                 optimizer.zero_grad()
188                 outputs = model(x_batch)
189                 loss = criterion(outputs, y_batch)
190                 loss.backward()
191                 optimizer.step()
192
193                 batch_losses.append(loss.item())
194                 train_iter += len(y_batch)
195
196             epoch_loss = sum(batch_losses) / len(batch_losses)
197
198             # --- Validation step ---
199             model.eval()
200             with torch.no_grad():
201                 val_outputs = model(x_val)
202                 val_loss = criterion(val_outputs, y_val)
203                 val_pred = torch.argmax(val_outputs, dim=1)
204                 val_acc = (val_pred == y_val).float().mean().item()
205
206                 top3 = torch.topk(val_outputs, k=3, dim=1).indices
207                 y_val_exp = y_val.unsqueeze(1).expand_as(top3)
208                 val_k3 = (top3 == y_val_exp).any(dim=1).float().mean().item()
209
210             print(f"Epoch: {epoch+1}, Iteration: {train_iter}, "
211                  f"Train Loss: {epoch_loss:.4f}, Val Loss: {val_loss.item():.4f}, "
212                  f"Val Acc: {val_acc:.4f}, Val Top-3 Acc: {val_k3:.4f}")
213             model.train()
214
215             """
216             # --- Early stopping check ---
217             if val_acc > best_val_acc:
218                 best_val_acc = val_acc
219                 patience_counter = 0
220             else:
221                 patience_counter += 1
222                 if patience_counter >= patience_limit:
223                     print("Early stopping triggered.")
224                     elapsed = time.time() - start_time
225                     t = time.gmtime(elapsed)
226                     print(f"Finished training in {t.tm_yday - 1} days {t.tm_hour} hours {t.tm_min} minutes {t.tm_sec} seconds with {
227                         ↪ train_iter} iterations.")
228                     class DummyEarlyStopping: stop_training = True
229                     return DummyEarlyStopping(), train_iter, f"Early stopped at epoch {epoch+1}"
230
231             if train_iter >= args.max_iter:
232                 break
233             if train_iter >= args.max_iter:
234                 break
235             train_epoch += 1
236
237 elapsed = time.time() - start_time
238 t = time.gmtime(elapsed)
239 print(f"Finished training in {t.tm_yday - 1} days {t.tm_hour} hours {t.tm_min} minutes {t.tm_sec} seconds with {train_iter} iterations.
240 ↪")

```

```

240     class DummyEarlyStopping: stop_training = False
241     return DummyEarlyStopping(), train_iter, f"Trained for {train_epoch} epochs"
242
243
244 def train_torch_lstm(model, args, train_ds):
245     device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
246     model.to(device)
247
248     optimizer = optim.Adam(
249         model.parameters(),
250         lr=config.learning_rate,
251         betas=(config.beta_1, config.beta_2),
252         eps=config.epsilon,
253         amsgrad=config.amsgrad
254     )
255     criterion = nn.CrossEntropyLoss()
256     model.train()
257
258     best_val_acc = 0
259     patience_counter = 0
260     patience_limit = 250
261
262     train_iter = 0
263     train_epoch = 0
264     start_time = time.time()
265
266     val_data_created = False
267     x_val = y_val = None
268
269     for epoch in range(args.epochs):
270         while train_ds.iteration < args.max_iter:
271             training_batches = next(train_ds)
272             for training_batch in training_batches:
273                 statistics, labels = training_batch.items()
274                 stats_np = statistics.numpy().astype(int)
275                 labels_np = labels.numpy()
276
277                 if not val_data_created:
278                     x_train_np, x_val_np, y_train_np, y_val_np = train_test_split(stats_np, labels_np, test_size=0.3)
279                     x_val = torch.tensor(x_val_np, dtype=torch.long).to(device)
280                     y_val = torch.tensor(y_val_np, dtype=torch.long).to(device)
281                     val_data_created = True
282                 else:
283                     x_train_np = stats_np
284                     y_train_np = labels_np
285
286                 x_train = torch.tensor(x_train_np, dtype=torch.long)
287                 y_train = torch.tensor(y_train_np, dtype=torch.long)
288
289                 train_dataset = TensorDataset(x_train, y_train)
290                 train_loader = DataLoader(train_dataset, batch_size=args.batch_size, shuffle=True)
291
292                 batch_losses = []
293                 for x_batch, y_batch in train_loader:
294                     x_batch = x_batch.to(device)
295                     y_batch = y_batch.to(device)
296
297                     optimizer.zero_grad()
298                     outputs = model(x_batch)
299                     loss = criterion(outputs, y_batch)
300                     loss.backward()
301                     optimizer.step()
302
303                     batch_losses.append(loss.item())
304                     train_iter += len(y_batch)
305
306                 epoch_loss = sum(batch_losses) / len(batch_losses)
307
308                 # --- Validation step ---
309                 model.eval()
310                 with torch.no_grad():
311                     val_outputs = model(x_val)
312                     val_loss = criterion(val_outputs, y_val)
313                     val_pred = torch.argmax(val_outputs, dim=1)
314                     val_acc = (val_pred == y_val).float().mean().item()
315
316                     top3 = torch.topk(val_outputs, k=3, dim=1).indices
317                     y_val_exp = y_val.unsqueeze(1).expand_as(top3)
318                     val_k3 = (top3 == y_val_exp).any(dim=1).float().mean().item()
319
320                 print(f"Epoch: {epoch+1}, Iteration: {train_iter}, "
321                     f"Train Loss: {epoch_loss:.4f}, Val Loss: {val_loss.item():.4f}, "
322                     f"Val Acc: {val_acc:.4f}, Val Top-3 Acc: {val_k3:.4f}")
323                 model.train()
324
325                 """
326                 # --- Early stopping check ---
327                 if val_acc > best_val_acc:
328                     best_val_acc = val_acc

```



```

329         patience_counter = 0
330     else:
331         patience_counter += 1
332         if patience_counter >= patience_limit:
333             print("Early stopping triggered.")
334             elapsed = time.time() - start_time
335             t = time.gmtime(elapsed)
336             print(f"Finished training in {t.tm_yday - 1} days {t.tm_hour} hours {t.tm_min} minutes {t.tm_sec} seconds with {
337                 ← train_iter} iterations.")
338             class DummyEarlyStopping: stop_training = True
339             return DummyEarlyStopping(), train_iter, f"Early stopped at epoch {epoch+1}"
340
341         if train_iter >= args.max_iter:
342             break
343         if train_iter >= args.max_iter:
344             break
345         train_epoch += 1
346
347     elapsed = time.time() - start_time
348     t = time.gmtime(elapsed)
349     print(f"Finished {train_epoch} epochs in {t.tm_hour}h {t.tm_min}m {t.tm_sec}s")
350     class DummyEarlyStopping: stop_training = False
351     return DummyEarlyStopping(), train_iter, f"Trained for {train_epoch} epochs"
352
353
354
355 def predict_torch_ffnn(model, test_ds, args):
356     device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
357     model.eval()
358     model.to(device)
359     criterion = nn.CrossEntropyLoss()
360
361     all_preds = []
362     all_labels = []
363
364     with torch.no_grad():
365         while test_ds.iteration < args.max_iter:
366             testing_batches = next(test_ds)
367             for testing_batch in testing_batches:
368                 statistics, labels = testing_batch.items()
369                 stats_np = statistics.numpy()
370
371                 x = torch.tensor(stats_np, dtype=torch.float32).to(device)
372                 y = torch.tensor(labels.numpy(), dtype=torch.long).to(device)
373
374                 outputs = model(x)
375                 loss = criterion(outputs, y)
376
377                 pred_top1 = torch.argmax(outputs, dim=1)
378                 acc = (pred_top1 == y).float().mean().item()
379
380                 top3 = torch.topk(outputs, k=3, dim=1).indices
381                 y_expanded = y.unsqueeze(1).expand_as(top3)
382                 k3_acc = (top3 == y_expanded).any(dim=1).float().mean().item()
383
384                 print(f"Eval → Loss: {loss.item():.4f}, Accuracy: {acc:.4f}, Top-3 Accuracy: {k3_acc:.4f}")
385
386                 preds = torch.softmax(outputs, dim=1).cpu().numpy()
387                 all_preds.append(preds)
388                 all_labels.append(labels.numpy())
389
390     all_preds = np.concatenate(all_preds, axis=0)
391     all_labels = np.concatenate(all_labels, axis=0)
392
393     return all_preds, all_labels
394
395
396
397
398 def predict_torch_lstm(model, test_ds, args):
399     device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
400     model.eval()
401     model.to(device)
402     criterion = nn.CrossEntropyLoss()
403
404     all_preds = []
405     all_labels = []
406
407     with torch.no_grad():
408         while test_ds.iteration < args.max_iter:
409             testing_batches = next(test_ds)
410             for testing_batch in testing_batches:
411                 statistics, labels = testing_batch.items()
412
413                 stats_np = statistics.numpy().astype(int) # input tokenizzati
414                 x = torch.tensor(stats_np, dtype=torch.long).to(device)
415                 y = torch.tensor(labels.numpy(), dtype=torch.long).to(device)
416
417                 outputs = model(x)

```

```

418         loss = criterion(outputs, y)
419
420         preds = torch.softmax(outputs, dim=1).cpu().numpy()
421         all_preds.append(preds)
422         all_labels.append(labels.numpy())
423
424     all_preds = np.concatenate(all_preds, axis=0)
425     all_labels = np.concatenate(all_labels, axis=0)
426
427     return all_preds, all_labels
428
429
430
431 def str2bool(v):
432     return v.lower() in ("yes", "true", "t", "1")
433
434
435 def create_model_with_distribution_strategy(architecture, extend_model, output_layer_size, max_train_len):
436     """Creates models depending on the GPU count and on extend_model"""
437     print('Creating model...')
438
439     strategy = None
440     gpu_count = (len(tf.config.list_physical_devices('GPU')) +
441                  len(tf.config.list_physical_devices('XLA_GPU')))
442     if gpu_count > 1:
443         print("Multiple GPUs found.")
444         strategy = tf.distribute.MirroredStrategy()
445         print(f"Number of mirrored devices: {strategy.num_replicas_in_sync}.")
446         with strategy.scope():
447             if extend_model is not None:
448                 extend_model = tf.keras.models.load_model(extend_model, compile=False)
449                 model = create_model(architecture, extend_model, output_layer_size, max_train_len)
450             if architecture in ("FNN", "CNN", "LSTM", "Transformer") and extend_model is None:
451                 if hasattr(model, "summary"):
452                     model.summary()
453                 else:
454                     # for LSTM use a LongTensor dummy input of shape (1, max_train_len)
455                     if architecture == "LSTM":
456                         summary(model, input_size=(1, max_train_len), dtypes=[torch.long])
457                     else:
458                         summary(model, input_size=(1, 724))
459             else:
460                 print("Only one GPU found.")
461                 strategy = NullDistributionStrategy()
462                 if extend_model is not None:
463                     extend_model = tf.keras.models.load_model(extend_model, compile=False)
464                     model = create_model(architecture, extend_model, output_layer_size, max_train_len)
465                 if architecture in ("FNN", "CNN", "LSTM", "Transformer") and extend_model is None:
466                     if hasattr(model, "summary"):
467                         model.summary()
468                     else:
469                         # for LSTM use a LongTensor dummy input of shape (1, max_train_len)
470                         if architecture == "LSTM":
471                             summary(model, input_size=(1, max_train_len), dtypes=[torch.long])
472                         else:
473                             summary(model, input_size=(1, 724))
474
475
476     print('Model created.\n')
477     return model, strategy
478
479
480 def create_model(architecture, extend_model, output_layer_size, max_train_len):
481     """
482     Creates an un-trained model to use in the training process.
483
484     The kind of model that is returned, depends on the provided architecture and
485     the 'extend_model' flag.
486
487     Parameters
488     -----
489     architecture : str
490         The architecture of the model to create.
491     extend_model
492         When 'extend_model' is not None and architecture in ('FNN', 'CNN', 'LSTM'),
493         the 'extend_model' will be further trained.
494     output_layer_size : int
495         Defines the size of the output layer of the neural networks
496
497     """
498     optimizer = Adam(
499         learning_rate=config.learning_rate, beta_1=config.beta_1, beta_2=config.beta_2,
500         epsilon=config.epsilon, amsgrad=config.amsgrad)
501
502     # Depends on the number of features returned by 'calculate_statistics()' in
503     # 'featureCalculations.py'.
504     input_layer_size = 724
505     hidden_layer_size = int(2 * (input_layer_size / 3) + output_layer_size)
506

```

```

507 # Create a model based on an existing one for further trainings
508 if extend_model is not None:
509     # remove the last layer
510     model = tf.keras.Sequential()
511     for layer in extend_model.layers[:-1]:
512         model.add(layer)
513     model.add(tf.keras.layers.Dense(output_layer_size, activation='softmax', name="output"))
514     model.compile(optimizer=optimizer, loss="sparse_categorical_crossentropy",
515                  metrics=["accuracy", SparseTopKCategoricalAccuracy(k=3, name="k3_accuracy")])
516     return model
517
518 # Create new model based on architecture
519 if architecture == "FFNN":
520     # Use PyTorch for FFNN
521     model = FFNN(
522         input_size=input_layer_size,
523         hidden_size=hidden_layer_size,
524         output_size=output_layer_size,
525         num_hidden_layers=config.hidden_layers
526     )
527     return model
528
529 elif architecture == "CNN":
530     config.FEATURE_ENGINEERING = False
531     config.PAD_INPUT = True
532     model = tf.keras.Sequential()
533     model.add(tf.keras.layers.Conv2D(
534         filters=config.filters, kernel_size=config.kernel_size,
535         input_shape=(max_train_len, 1), activation='relu'))
536     for _ in range(config.layers - 1):
537         model.add(tf.keras.layers.Conv2D(filters=config.filters, kernel_size=config.kernel_size, activation='relu'))
538     # model.add(tf.keras.layers.Dropout(0.2))
539     model.add(tf.keras.layers.MaxPooling2D(pool_size=2))
540     model.add(tf.keras.layers.Flatten())
541     model.add(tf.keras.layers.Dense(output_layer_size, activation='softmax'))
542     model.compile(optimizer=optimizer, loss="sparse_categorical_crossentropy",
543                  metrics=["accuracy", SparseTopKCategoricalAccuracy(k=3, name="k3_accuracy")])
544     return model
545
546 elif architecture == "LSTM":
547     config.FEATURE_ENGINEERING = False
548     config.PAD_INPUT = True
549     model = LSTM(
550         vocab_size=56,
551         embed_dim=64,
552         hidden_size=config.lstm_units,
553         output_size=output_layer_size,
554         num_layers=1,
555         dropout=0.0
556     )
557     return model
558
559 elif architecture == "DT":
560     return DecisionTreeClassifier(criterion=config.criterion, ccp_alpha=config.ccp_alpha)
561
562 elif architecture == "NB":
563     return MultinomialNB(alpha=config.alpha, fit_prior=config.fit_prior)
564
565 elif architecture == "RF":
566     return RandomForestClassifier(n_estimators=config.n_estimators, criterion=config.criterion,
567                                  bootstrap=config.bootstrap, n_jobs=30,
568                                  max_features=config.max_features, max_depth=30,
569                                  min_samples_split=config.min_samples_split,
570                                  min_samples_leaf=config.min_samples_leaf)
571
572 elif architecture == "ET":
573     return ExtraTreesClassifier(n_estimators=config.n_estimators, criterion=config.criterion,
574                                 bootstrap=config.bootstrap, n_jobs=30,
575                                 max_features=config.max_features, max_depth=30,
576                                 min_samples_split=config.min_samples_split,
577                                 min_samples_leaf=config.min_samples_leaf)
578
579 elif architecture == "Transformer":
580     config.FEATURE_ENGINEERING = False
581     config.PAD_INPUT = True
582     vocab_size = config.vocab_size
583     maxlen = max_train_len
584     embed_dim = config.embed_dim # Embedding size for each token
585     num_heads = config.num_heads # Number of attention heads
586     ff_dim = config.ff_dim # Hidden layer size in feed forward network inside transformer
587
588     inputs = tf.keras.layers.Input(shape=(maxlen,))
589     embedding_layer = TokenAndPositionEmbedding(maxlen, vocab_size, embed_dim)
590     x = embedding_layer(inputs)
591     transformer_block = TransformerBlock(embed_dim, num_heads, ff_dim)
592     x = transformer_block(x)
593     x = tf.keras.layers.GlobalAveragePooling1D()(x)
594     outputs = tf.keras.layers.Dense(output_layer_size, activation="softmax")(x)
595

```

```

596 model = tf.keras.Model(inputs=inputs, outputs=outputs)
597 model.compile(optimizer=optimizer, loss="sparse_categorical_crossentropy",
598               metrics=["accuracy", SparseTopKCategoricalAccuracy(k=3, name="k3_accuracy")])
599 return model
600
601 elif architecture == "SVM":
602     return SVC(probability=True, C=1, gamma=0.001, kernel="linear")
603
604 elif architecture == "SVM-Rotor":
605     pipe = Pipeline([
606         ('scale', StandardScaler()),
607         ('clf', SVC(probability=True, C=10, gamma=0.001, kernel="rbf"))])
608     return pipe
609
610 elif architecture == "kNN":
611     return KNeighborsClassifier(90, weights="distance", metric="euclidean")
612
613 elif architecture == "[FNN,NB]":
614     model_ffnn = tf.keras.Sequential()
615     model_ffnn.add(tf.keras.layers.Input(shape=(input_layer_size,)))
616     for _ in range(config.hidden_layers):
617         model_ffnn.add(tf.keras.layers.Dense(hidden_layer_size, activation='relu', use_bias=True))
618     model_ffnn.add(tf.keras.layers.Dense(output_layer_size, activation='softmax'))
619     model_ffnn.compile(optimizer=optimizer, loss="sparse_categorical_crossentropy",
620                       metrics=["accuracy", SparseTopKCategoricalAccuracy(k=3, name="k3_accuracy")])
621     model_nb = MultinomialNB(alpha=config.alpha, fit_prior=config.fit_prior)
622     return [model_ffnn, model_nb]
623
624 elif architecture == "[DT,ET,RF,SVM,kNN]":
625     dt = DecisionTreeClassifier(criterion=config.criterion, ccp_alpha=config.ccp_alpha)
626     et = ExtraTreesClassifier(n_estimators=config.n_estimators, criterion=config.criterion,
627                              bootstrap=config.bootstrap, n_jobs=30,
628                              max_features=config.max_features, max_depth=30,
629                              min_samples_split=config.min_samples_split,
630                              min_samples_leaf=config.min_samples_leaf)
631     rf = RandomForestClassifier(n_estimators=config.n_estimators, criterion=config.criterion,
632                               bootstrap=config.bootstrap, n_jobs=30,
633                               max_features=config.max_features, max_depth=30,
634                               min_samples_split=config.min_samples_split,
635                               min_samples_leaf=config.min_samples_leaf)
636     svm = SVC(probability=True, C=1, gamma=0.001, kernel="linear")
637     knn = KNeighborsClassifier(90, weights="distance", metric="euclidean")
638     return [dt, et, rf, svm, knn]
639
640 else:
641     raise Exception(f"Could not create model. Unknown architecture '{architecture}'.")
642
643 def parse_arguments():
644     parser = argparse.ArgumentParser(
645         description='CANN Ciphertext Detection Neuronal Network Training Script',
646         formatter_class=argparse.RawTextHelpFormatter)
647     parser.add_argument('--batch_size', default=128, type=int,
648                         help='Batch size for training.')
649     parser.add_argument('--train_dataset_size', default=16000, type=int,
650                         help='Dataset size per fit. This argument should be dividable \n'
651                              'by the amount of --ciphers.')
652     parser.add_argument('--dataset_workers', default=1, type=int,
653                         help='The number of parallel workers for reading the \ninput files.')
654     parser.add_argument('--epochs', default=1, type=int,
655                         help='Defines how many times the same data is used to fit the model.')
656     parser.add_argument('--plaintext_input_directory', default='../data/gutenberg_en', type=str,
657                         help='Input directory of the plaintexts for training the aca ciphers.')
658     parser.add_argument('--rotor_input_directory', default='../data/rotor_ciphertexts', type=str,
659                         help='Input directory of the rotor ciphertexts.')
660     parser.add_argument('--download_dataset', default=True, type=bool,
661                         help='Download the dataset automatically.')
662     parser.add_argument('--save_directory', default='../data/models/',
663                         help='Directory for saving generated models. \n'
664                              'When interrupting, the current model is \n'
665                              'saved as interrupted_...')
666     parser.add_argument('--model_name', default='m.h5', type=str,
667                         help='Name of the output model file. The file must \nhave the .h5 extension.')
668     parser.add_argument('--ciphers', default='all', type=str,
669                         help='A comma separated list of the ciphers to be created.\n'
670                              'Be careful to not use spaces or use \' \' to define the string.\n'
671                              'Possible values are:\n'
672                              '- mtc3 (contains the ciphers Monoalphabetic Substitution, Vigenere, \n'
673                              'Columnar Transposition, Plaifair and Hill)\n'
674                              '- aca (contains all currently implemented ciphers from \n'
675                              'https://www.cryptogram.org/resource-area/cipher-types/)\n'
676                              '- rotor (contains Enigma, M209, Purple, Sigaba and Typex ciphers)\n'
677                              '- all (contains aca and rotor ciphers)\n'
678                              '- all aca ciphers in lower case\n'
679                              '- simple_substitution\n'
680                              '- vigenere\n'
681                              '- columnar_transposition\n'
682                              '- playfair\n'
683

```

```

685         '- hill\n')
686     parser.add_argument('--keep_unknown_symbols', default=False, type=str2bool,
687                         help='Keep unknown symbols in the plaintexts. Known \n'
688                             'symbols are defined in the alphabet of the cipher.')
689     parser.add_argument('--max_iter', default=1000000, type=int,
690                         help='the maximal number of iterations before stopping training.')
691     parser.add_argument('--min_train_len', default=50, type=int,
692                         help='The minimum length of a plaintext to be encrypted in training. \n'
693                             'If this argument is set to -1 no lower limit is used.')
694     parser.add_argument('--min_test_len', default=50, type=int,
695                         help='The minimum length of a plaintext to be encrypted in testing. \n'
696                             'If this argument is set to -1 no lower limit is used.')
697     parser.add_argument('--max_train_len', default=-1, type=int,
698                         help='The maximum length of a plaintext to be encrypted in training. \n'
699                             'If this argument is set to -1 no upper limit is used.')
700     parser.add_argument('--max_test_len', default=-1, type=int,
701                         help='The maximum length of a plaintext to be encrypted in testing. \n'
702                             'If this argument is set to -1 no upper limit is used.')
703     parser.add_argument('--architecture', default='FFNN', type=str,
704                         choices=['FFNN', 'CNN', 'LSTM', 'DT', 'NB', 'RF', 'ET', 'Transformer',
705                                 'SVM', 'kNN', '[FFNN,NB]', '[DT,ET,RF,SVM,kNN]', 'SVM-Rotor'],
706                         help='The architecture to be used for training. \n'
707                             'Possible values are:\n'
708                             '- FFNN\n'
709                             '- CNN\n'
710                             '- LSTM\n'
711                             '- DT\n'
712                             '- NB\n'
713                             '- RF\n'
714                             '- ET\n'
715                             '- Transformer\n'
716                             '- SVM\n'
717                             '- kNN\n'
718                             '- [FFNN,NB]\n'
719                             '- [DT,ET,RF,SVM,kNN]\n'
720                             '- SVM-Rotor\n'
721                         )
722     parser.add_argument('--extend_model', default=None, type=str,
723                         help='Load a trained model from a file and use it as basis for the new training.')
724
725     return parser.parse_args()
726
727 def should_download_plaintext_datasets(args):
728     """Determines if the plaintext datasets should be loaded"""
729     return (args.download_dataset and
730             not os.path.exists(args.plaintext_input_directory) and
731             args.plaintext_input_directory == os.path.abspath('../data/gutenberg_en'))
732
733 def download_plaintext_datasets(args):
734     """Downloads plaintexts and saves them in the plaintext_input_directory"""
735     print("Downloading Datasets...")
736     checksums_dir = '../data/checksums/'
737     if not Path(checksums_dir).exists():
738         os.mkdir(checksums_dir)
739     tfds.download.add_checksums_dir(checksums_dir)
740
741     download_manager = tfds.download.DownloadManager(download_dir='../data/',
742                                                       extract_dir=args.plaintext_input_directory)
743     data_url = ('https://drive.google.com/uc?id=1bF5sSVjxTx3ADB-P5wxn87nxWnDRhK_V&export=download' +
744               '&confirm=t&uiid=afbc362d-9d52-472a-832b-c2af331a8d5b')
745     try:
746         download_manager.download_and_extract(data_url)
747     except Exception as e:
748         print("Download of datasets failed. If this issues persists, try downloading the dataset yourself "
749               "from: https://drive.google.com/file/d/1bF5sSVjxTx3ADB-P5wxn87nxWnDRhK_V/view."
750               "(For more information see the README.md of this project.)")
751         print("Underlying error:")
752         print(e)
753         sys.exit(1)
754
755     path = os.path.join(args.plaintext_input_directory,
756                        'ZIP_ucid_1bF5sSVjxTx-P5wxn87nxWn_V_export_downloadR9Cwhunev5CvJ-ic__'
757                        'HawxhTtGOISdcCrro4fxfEI8A')
758     os.path.basename(args.plaintext_input_directory))
759     dir_name = os.listdir(path)
760     for name in dir_name:
761         p = Path(os.path.join(path, name))
762         parent_dir = p.parents[2]
763         p.rename(parent_dir / p.name)
764     os.rmdir(path)
765     os.rmdir(os.path.dirname(path))
766     print("Datasets Downloaded.")
767
768 def load_rotor_ciphertext_datasets_from_disk(args, batch_size):
769     """
770     Load ciphertext input data for rotor ciphertexts from disk.
771
772     This method (in contrast to 'load_plaintext_datasets_from_disk') loads the data
773     immediately from disk. Currently this does not take too long, since the input files

```

```

774 and the number of ciphers that need ciphertexts as input for the feature extraction
775 are limited. (If this method should lazily load ciphertexts, 'RotorCiphertextsDataset'
776 needs to be adapted.)
777
778 Parameters
779 -----
780 args :
781     The arguments parsed by the 'ArgumentParser'. See 'parse_arguments()'.
782 batch_size : int
783     The number of samples and labels per batch.
784
785 Returns
786 -----
787 tuple[list, RotorCiphertextsDatasetParameters, list, RotorCiphertextsDatasetParameters]
788     A tuple with the training and testing ciphertexts as well as their
789     'RotorCiphertextsDatasetParameters' that are both provided to the
790     'CipherStatisticsDataset's for training and testing.
791 """
792
793 def validate_ciphertext_path(ciphertext_path, cipher_types):
794     """Check if the filename of the file at 'ciphertext_path' matches the
795     names inside 'cipher_types'."""
796     file_name = Path(ciphertext_path).stem.lower()
797     if not file_name in cipher_types:
798         raise Exception(f"Filename must equal one of the expected cipher types. "
799                         f"Expected cipher types are: {cipher_types}. Current "
800                         f"filename is '{file_name}'.")
801
802
803 # Filter cipher_types to exclude non-ciphertext ciphers
804
805 # Filter cipher_types to exclude non-ciphertext ciphers
806 rotor_cipher_types = [config.CIPHER_TYPES[i] for i in range(56, 61)]
807
808 # Load all ciphertexts in the train and dev folders in 'args.rotor_input_directory'
809 rotor_cipher_dir = args.rotor_input_directory
810
811 def find_ciphertext_paths_in_dir(folder_path):
812     """Loads all .txt files in the given folder and checks that their names match the
813     known cipher types."""
814     file_names = os.listdir(folder_path)
815     result = []
816     for name in file_names:
817         path = os.path.join(folder_path, name)
818         file_name, file_type = os.path.splitext(name)
819         if os.path.isfile(path) and file_name.lower() in rotor_cipher_types and file_type == ".txt":
820             validate_ciphertext_path(path, config.ROTOR_CIPHER_TYPES)
821
822     result.append(path)
823     return result
824
825 train_dir = os.path.join(rotor_cipher_dir, "train")
826 test_dir = os.path.join(rotor_cipher_dir, "dev")
827 train_rotor_ciphertext_paths = find_ciphertext_paths_in_dir(train_dir)
828 test_rotor_ciphertext_paths = find_ciphertext_paths_in_dir(test_dir)
829
830
831 # Return empty lists and parameters if no requested ciphers were found on disk
832 if len(train_rotor_ciphertext_paths) == 0 or len(test_rotor_ciphertext_paths) == 0:
833     empty_params = RotorCiphertextsDatasetParameters(config.ROTOR_CIPHER_TYPES,
834                                                       0,
835                                                       args.dataset_workers,
836                                                       args.min_train_len,
837                                                       args.max_train_len,
838                                                       generate_evaluation_data=False)
839     return ([], empty_params, [], empty_params)
840
841 # Create dataset parameters, which will be used for creating a 'CipherStatisticsDataset'.
842 # This class will provide an iterator in 'train_model' to convert the plaintext files
843 # (applying the provided options) into statistics (features) used for training.
844 train_rotor_ciphertexts_parameters = RotorCiphertextsDatasetParameters(config.ROTOR_CIPHER_TYPES,
845                                batch_size,
846                                args.dataset_workers,
847                                args.min_train_len,
848                                args.max_train_len,
849                                generate_evaluation_data=False)
850 test_rotor_ciphertexts_parameters = RotorCiphertextsDatasetParameters(config.ROTOR_CIPHER_TYPES,
851                                batch_size,
852                                args.dataset_workers,
853                                args.min_test_len,
854                                args.max_test_len,
855                                generate_evaluation_data=False)
856
857 # Return the tuples of training and testing rotor_ciphertexts as well as the parameter
858 # for initializing the 'CipherStatisticsDataset's.
859 return (train_rotor_ciphertext_paths, train_rotor_ciphertexts_parameters,
860         test_rotor_ciphertext_paths, test_rotor_ciphertexts_parameters)
861
862 def load_plaintext_datasets_from_disk(args, requested_cipher_types, batch_size):

```

```

863 """
864 Gets all plaintext paths found in 'args.plaintext_input_directory', and converts
865 them into training and testing list and parameters, used to create
866 'CipherStatisticsDataset's.
867
868 This method does not load the contents of the plaintext files. This is done
869 lazily by the 'CipherStatisticsDataset'.
870
871 Parameters
872 -----
873 args :
874     The arguments parsed by the 'ArgumentParser'. See 'parse_arguments()'.
875 requested_cipher_types : list
876     A list of cipher types to provide as parameters to the 'CipherStatisticsDataset'.
877     The list is filtered and only ACA ciphers are used as parameters, since the features
878     of rotor ciphers currently have to be extracted from ciphertext files.
879 batch_size : int
880     The number of samples and labels per batch.
881
882 Returns
883 -----
884 tuple[list, PlaintextPathsDatasetParameters, list, PlaintextPathsDatasetParameters]
885     A tuple with the training and testing plaintext paths as well as their
886     'PlaintextPathsDatasetParameters' that are both provided to the
887     'CipherStatisticsDataset's for training and testing.
888 """
889 # Filter cipher_types to exclude non-plaintext ciphers
890 aca_cipher_types = [config.CIPHER_TYPES[i] for i in range(56)]
891 cipher_types = [type for type in requested_cipher_types if type in aca_cipher_types]
892
893 # Get all paths to plaintext files in the 'plaintext_input_directory'
894 plaintext_files = []
895 dir_name = os.listdir(args.plaintext_input_directory)
896 for name in dir_name:
897     path = os.path.join(args.plaintext_input_directory, name)
898     if os.path.isfile(path):
899         plaintext_files.append(path)
900
901 # Use some plaintext for training and others for testing
902 train_plaintexts, test_plaintexts = train_test_split(plaintext_files, test_size=0.05,
903                                                     random_state=42, shuffle=True)
904
905 # Create dataset parameters, which will be used for creating a 'CipherStatisticsDataset'.
906 # This class will provide an iterator in 'train_model' to convert the plaintext files
907 # (applying the provided options) into statistics (features) used for training.
908 train_plaintext_parameters = PlaintextPathsDatasetParameters(cipher_types, batch_size,
909                                                             args.min_train_len, args.max_train_len,
910                                                             args.keep_unknown_symbols, args.dataset_workers)
911 test_plaintext_parameters = PlaintextPathsDatasetParameters(cipher_types, batch_size,
912                                                             args.min_test_len, args.max_test_len,
913                                                             args.keep_unknown_symbols, args.dataset_workers)
914
915 # Return the training and testing plaintexts as well as their parameters
916 return (train_plaintexts, train_plaintext_parameters, test_plaintexts, test_plaintext_parameters)
917
918 def load_datasets_from_disk(args, requested_cipher_types):
919     """
920     Loads training and testing data from the file system.
921
922     In case of the ACA ciphers the datasets are plaintext files that need to be
923     encrypted before the features can be extracted. In case of the rotor ciphers
924     there are already encrypted ciphertext files that can directly be used to
925     extract the features.
926     To simplify the training code, both kinds of input data are returned in
927     'CipherStatisticsDataset's that provide an iterator interface, returning
928     'TrainingBatch'es of the requested size on each 'next()' call.
929     Plaintext input is loaded lazily, while ciphertexts currently are loaded
930     immediately.
931
932     Parameters
933     -----
934     args :
935         The parsed commandline arguments. See also 'parse_arguments()'.
936     requested_cipher_types : list
937         A list of the requested cipher types. These are provided as parameters to
938         the returned 'CipherStatisticsDataset's as well as for selection of input
939         files.
940     max_rotor_lines : int
941         Limits the amount of input lines loaded from ciphertext files.
942
943     Returns
944     -----
945     tuple[CipherStatisticsDataset]
946         Training and testing 'CipherStatisticsDataset' that lazily calculate the
947         features for the input data on 'next()' calls.
948     """
949     print("Loading Datasets...")
950

```

```

952 # Filter cipher_types to exclude non-ciphertext ciphers
953 rotor_cipher_types = [config.CIPHER_TYPES[i] for i in range(56, 61)]
954 non_rotor_ciphers = [type for type in requested_cipher_types if type not in rotor_cipher_types]
955 rotor_cipher_types = [type for type in requested_cipher_types if type in rotor_cipher_types]
956
957 # Calculate batch size for rotor ciphers. If both aca and rotor ciphers are requested,
958 # the amount of samples of each rotor cipher per batch should be equal to the
959 # amount of samples of each aca cipher per loaded batch.
960 number_of_rotor_ciphers = len(rotor_cipher_types)
961 number_of_aca_ciphers = len(non_rotor_ciphers)
962 if number_of_aca_ciphers <= 0:
963     rotor_dataset_batch_size = args.train_dataset_size
964     aca_dataset_batch_size = 0
965 else:
966     amount_of_samples_per_cipher = args.train_dataset_size // (number_of_aca_ciphers + number_of_rotor_ciphers)
967     rotor_dataset_batch_size = amount_of_samples_per_cipher * number_of_rotor_ciphers
968     aca_dataset_batch_size = amount_of_samples_per_cipher * number_of_aca_ciphers
969
970 # Load the plaintext file paths and the rotor ciphertexts from disk.
971 (train_plaintexts,
972  train_plaintext_parameters,
973  test_plaintexts,
974  test_plaintext_parameters) = load_plaintext_datasets_from_disk(args,
975                                                                    requested_cipher_types,
976                                                                    aca_dataset_batch_size)
977 (train_rotor_ciphertext_paths,
978  train_rotor_ciphertexts_parameters,
979  test_rotor_ciphertext_paths,
980  test_rotor_ciphertexts_parameters) = load_rotor_ciphertext_datasets_from_disk(args,
981                                                                                rotor_dataset_batch_size)
982
983 # Convert the training and testing ciphertexts and plaintexts, as well as
984 # their parameters into 'CipherStatisticsDataset's.
985 train_ds = CipherStatisticsDataset(train_plaintexts, train_plaintext_parameters, train_rotor_ciphertext_paths,
986                                   train_rotor_ciphertexts_parameters)
987 test_ds = CipherStatisticsDataset(test_plaintexts, test_plaintext_parameters, test_rotor_ciphertext_paths,
988                                  test_rotor_ciphertexts_parameters)
989
990 if train_ds.key_lengths_count > 0 and args.train_dataset_size % train_ds.key_lengths_count != 0:
991     print("WARNING: the --train_dataset_size parameter must be dividable by the amount of --ciphers and the length configured "
992           "KEY_LENGTHS in config.py. The current key_lengths_count is %d" %
993         train_ds.key_lengths_count, file=sys.stderr)
994
995 print("Datasets loaded.\n")
996
997 # Return 'CipherStatisticsDataset's for training and testing.
998 return train_ds, test_ds
999
1000 def train_model(model, strategy, args, train_ds):
1001     """
1002     Trains the model with the given training dataset.
1003
1004     Depending on the value of 'args.architecture' a different approach is
1005     taken to train the model. Some architectures need to be trained in one
1006     iteration, while others can be trained with multiple input batches.
1007     While the training is in progress, status messages are logged to
1008     stdout to indicate the amount of seen input data as well as the current
1009     accuracy of the trained model.
1010
1011     Parameters
1012     -----
1013     model :
1014         The model that will be trained. Needs to match 'args.architecture'.
1015     strategy :
1016         A distribution strategy (of the 'Tensorflow' library) to distribute
1017         the 'fit' calls to multiple devices. Could also be a 'NullStrategy'
1018         if no GPU devices are found on the system.
1019     args :
1020         Commandline arguments entered by the user. See also: 'parse_arguments()'.
1021     train_ds :
1022         A 'CipherStatisticsDataset' providing the features to use for training.
1023
1024     Returns
1025     -----
1026     tuple
1027     """
1028     checkpoints_dir = Path('../data/checkpoints')
1029     def delete_previous_checkpoints():
1030         shutil.rmtree(checkpoints_dir)
1031
1032     def create_checkpoint_callback():
1033         """Provides a 'keras' 'ModelCheckpoint' used to periodically save a model in training"""
1034         if not checkpoints_dir.exists():
1035             os.mkdir(checkpoints_dir)
1036         checkpoint_file_path = os.path.join(checkpoints_dir,
1037                                             "epoch_{epoch:02d}-"
1038                                             "acc_{accuracy:.2f}.h5")
1039
1040     return tf.keras.callbacks.ModelCheckpoint(

```



```

1040         filepath=checkpoint_file_path ,
1041         save_weights_only=False ,
1042         monitor='val_accuracy' ,
1043         mode='max' ,
1044         save_best_only=False ,
1045         save_freq=100)
1046
1047     print('Training model...')
1048
1049     delete_previous_checkpoints()
1050
1051     # Create callbacks for tensorflow models
1052     tensorboard_callback = tf.keras.callbacks.TensorBoard(log_dir='../data/logs' ,
1053                                                         update_freq='epoch')
1054     early_stopping_callback = MiniBatchEarlyStopping(min_delta=1e-5,
1055                                                     patience=250,
1056                                                     monitor='accuracy' ,
1057                                                     mode='max' ,
1058                                                     restore_best_weights=True)
1059     custom_step_decay_lrate_callback = CustomStepDecayLearningRateScheduler(early_stopping_callback)
1060     checkpoint_callback = create_checkpoint_callback()
1061
1062     # Initialize variables
1063     architecture = args.architecture
1064     start_time = time.time()
1065     train_iter = 0
1066     train_epoch = 0
1067     val_data = None
1068     val_labels = None
1069     training_batches = None
1070     combined_batch = TrainingBatch("mixed", [], [])
1071     classes = list(range(len(config.CIPHER_TYPES)))
1072     should_create_validation_data = True
1073
1074     if args.architecture == "FFNN" and isinstance(model, FFNN):
1075         return train_torch_ffnn(model, args, train_ds)
1076
1077     elif args.architecture == "LSTM" and isinstance(model, LSTM):
1078         return train_torch_lstm(model, args, train_ds)
1079
1080
1081     # Perform main training loop while the iterations don't exceed the user provided max_iter
1082     while train_ds.iteration < args.max_iter:
1083         training_batches = next(train_ds)
1084
1085         # For architectures that only support one fit call: Sample all batches into one large batch.
1086         if architecture in ("DT", "RF", "ET", "SVM", "kNN", "SVM-Rotor", "[DT,ET,RF,SVM,kNN]"):
1087             for training_batch in training_batches:
1088                 combined_batch.extend(training_batch)
1089             if train_ds.iteration < args.max_iter:
1090                 print("Loaded %d ciphertexts." % train_ds.iteration)
1091                 continue
1092             train_ds.stop_outstanding_tasks()
1093             print("Loaded %d ciphertexts." % train_ds.iteration)
1094             training_batches = [combined_batch]
1095
1096         for index, training_batch in enumerate(training_batches):
1097             statistics, labels = training_batch.items()
1098             train_iter = train_ds.iteration - len(training_batch) * (len(training_batches) - index - 1)
1099
1100             # Create small validation dataset on first iteration
1101             if should_create_validation_data:
1102                 statistics, val_data, labels, val_labels = train_test_split(statistics.numpy(),
1103                                                                           labels.numpy(),
1104                                                                           test_size=0.3)
1105                 statistics = tf.convert_to_tensor(statistics)
1106                 val_data = tf.convert_to_tensor(val_data)
1107                 labels = tf.convert_to_tensor(labels)
1108                 val_labels = tf.convert_to_tensor(val_labels)
1109                 should_create_validation_data = False
1110                 train_iter -= len(training_batch) * 0.3
1111
1112             # scikit-learn architectures:
1113             if architecture in ("DT", "RF", "ET", "SVM", "kNN", "SVM-Rotor"):
1114                 train_iter = len(labels) * 0.7
1115                 print(f"Start training the {architecture}.")
1116                 if architecture == "kNN":
1117                     history = model.fit(statistics, labels)
1118                 elif architecture == "SVM-Rotor":
1119                     history = model.fit(list(statistics), list(labels))
1120                     # print(f"RFE-support: \n{list(model.support_)}\n")
1121                     # print(f"RFE-rank: \n{list(model.ranking_)}\n")
1122                     # print()
1123                 else:
1124                     history = model.fit(statistics, labels)
1125                 if architecture == "DT":
1126                     plt.gcf().set_size_inches(25, 25 / math.sqrt(2))
1127                     print("Plotting tree.")
1128                     plot_tree(model, max_depth=3, fontsize=6, filled=True)

```

```

1129         plt.savefig(args.model_name.split('.')[0] + '_decision_tree.svg',
1130                     dpi=200, bbox_inches='tight', pad_inches=0.1)
1131
1132     # Naive Bayes training
1133     elif architecture == "NB":
1134         history = model.partial_fit(statistics,
1135                                     labels,
1136                                     classes=classes)
1137
1138     # Ensemble: [FFNN,NB]
1139     elif architecture == "[FFNN,NB]":
1140         with strategy.scope():
1141             history = model[0].fit(statistics,
1142                                    labels,
1143                                    batch_size=args.batch_size,
1144                                    validation_data=(val_data, val_labels),
1145                                    epochs=args.epochs,
1146                                    callbacks=[early_stopping_callback,
1147                                              tensorboard_callback,
1148                                              custom_step_decay_lrate_callback,
1149                                              checkpoint_callback])
1150         history = model[1].partial_fit(statistics,
1151                                       labels,
1152                                       classes=classes)
1153
1154     # Ensemble: [DT,ET,RF,SVM,kNN]
1155     elif architecture == "[DT,ET,RF,SVM,kNN]":
1156         print(f"Start training the {architecture}.")
1157         dt, et, rf, svm, knn = model
1158         for index, m in enumerate([dt, et, rf, svm]):
1159             m.fit(statistics, labels)
1160             print(f"Trained model {index + 1} of {len(model)}")
1161         knn.fit(statistics, labels)
1162         print(f"Trained model {len(model)} of {len(model)}")
1163
1164     else:
1165         with strategy.scope():
1166             history = model.fit(statistics, labels,
1167                                batch_size=args.batch_size,
1168                                validation_data=(val_data, val_labels),
1169                                epochs=args.epochs,
1170                                callbacks=[early_stopping_callback,
1171                                          tensorboard_callback,
1172                                          custom_step_decay_lrate_callback,
1173                                          checkpoint_callback])
1174
1175     # print for Decision Tree, Naive Bayes and Random Forests
1176     if architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN", "SVM-Rotor"):
1177         val_score = model.score(val_data, val_labels)
1178         train_score = model.score(statistics, labels)
1179         print("train accuracy: %f, validation accuracy: %f" % (train_score, val_score))
1180
1181     if architecture == "[FFNN,NB]":
1182         val_score = model[1].score(val_data, val_labels)
1183         train_score = model[1].score(statistics, labels)
1184         print("train accuracy: %f, validation accuracy: %f" % (train_score, val_score))
1185
1186     if architecture == "[DT,ET,RF,SVM,kNN]":
1187         for m in model:
1188             val_score = m.score(val_data, val_labels)
1189             train_score = m.score(statistics, labels)
1190             print(f"{type(m).__name__}: train accuracy: {train_score}, "
1191                   f"validation accuracy: {val_score}")
1192
1193     if train_ds.epoch > 0:
1194         train_epoch = (train_ds.iteration
1195                        // ((train_iter + train_ds.batch_size * train_ds.dataset_workers)
1196                           // train_ds.epoch))
1197
1198     print("Epoch: %d, Iteration: %d" % (train_epoch, train_iter))
1199     if train_iter >= args.max_iter or early_stopping_callback.stop_training:
1200         break
1201
1202     if train_ds.iteration >= args.max_iter or early_stopping_callback.stop_training:
1203         train_ds.stop_outstanding_tasks()
1204         break
1205
1206     elapsed_training_time = datetime.fromtimestamp(time.time()) - datetime.fromtimestamp(start_time)
1207     training_stats = ("Finished training in %d days %d hours %d minutes %d seconds "
1208                      "with %d iterations and %d epochs.\n"
1209                      % (elapsed_training_time.days,
1210                        elapsed_training_time.seconds // 3600,
1211                        (elapsed_training_time.seconds // 60) % 60,
1212                        elapsed_training_time.seconds % 60,
1213                        train_iter,
1214                        train_epoch))
1215     print(training_stats)
1216     return early_stopping_callback, train_iter, training_stats
1217

```

```

1218 def save_model(model, args):
1219     """Writes the model and the commandline arguments to disk."""
1220     print('Saving model...')
1221     architecture = args.architecture
1222
1223     if not os.path.exists(args.save_directory):
1224         os.mkdir(args.save_directory)
1225
1226     # Gestione nome modello
1227     if args.model_name == 'm.h5':
1228         i = 1
1229         base_name = args.model_name.split('.')[0]
1230         extension = '.pth' if architecture == "FFNN" else '.h5'
1231         while os.path.exists(os.path.join(args.save_directory, base_name + str(i) + extension)):
1232             i += 1
1233         model_name = base_name + str(i) + extension
1234     else:
1235         model_name = args.model_name
1236         if architecture == "FFNN":
1237             model_name = model_name.replace('.h5', '.pth')
1238
1239     model_path = os.path.join(args.save_directory, model_name)
1240
1241     if architecture in ("FFNN", "LSTM"):
1242         state_dict = {
1243             'model_state_dict': model.state_dict(),
1244             'hidden_size': model.hidden_size,
1245             'output_size': model.output_size,
1246         }
1247
1248         if architecture == "FFNN":
1249             state_dict['input_size'] = model.input_size
1250             state_dict['num_hidden_layers'] = model.num_hidden_layers
1251         elif architecture == "LSTM":
1252             state_dict['vocab_size'] = model.vocab_size
1253             state_dict['embed_dim'] = model.embed_dim
1254             state_dict['num_layers'] = model.num_layers
1255             state_dict['dropout'] = model.dropout
1256
1257         torch.save(state_dict, model_path)
1258
1259     elif architecture in ("CNN", "Transformer"):
1260         model.save(model_path)
1261
1262     elif architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN", "SVM-Rotor"):
1263         with open(model_path, "wb") as f:
1264             pickle.dump(model, f)
1265
1266     elif architecture == "[FFNN,NB]":
1267         model[0].save('./data/models/' + model_path.split('.')[0] + "_ffnn.h5")
1268         with open('./data/models/' + model_path.split('.')[0] + "_nb.h5", "wb") as f:
1269             pickle.dump(model[1], f)
1270
1271     elif architecture == "[DT,ET,RF,SVM,kNN]":
1272         for index, name in enumerate(["dt", "et", "rf", "svm", "knn"]):
1273             with open('./data/models/' + model_path.split('.')[0] + f"_{name}.h5", "wb") as f:
1274                 pickle.dump(model[index], f)
1275
1276     """
1277
1278     # Saving parameters
1279     with open('./data/' + model_path.split('.')[0] + '_parameters.txt', 'w') as f:
1280         for arg in vars(args):
1281             f.write("{:23s}={:s}\n".format(arg, str(getattr(args, arg))))
1282
1283     # Managing logs
1284     if architecture in ("FFNN", "CNN", "LSTM", "Transformer"):
1285         logs_destination = './data/' + model_name.split('.')[0] + '_tensorboard_logs'
1286         try:
1287             if os.path.exists('./data/logs'):
1288                 if os.path.exists(logs_destination):
1289                     shutil.rmtree(logs_destination)
1290                     shutil.move('./data/logs', logs_destination)
1291             except Exception:
1292                 print(f"Could not move logs from './data/logs' to '{logs_destination}'.")
1293
1294     """
1295     print('Model saved.\n')
1296
1297
1298 def predict_test_data(test_ds, model, args, early_stopping_callback, train_iter):
1299     """
1300     Testing the predictions of the model.
1301
1302     The trained model is used to predict the data in 'test_ds' and the results
1303     are evaluated in regard to accuracy, precision, recall, etc. The calculated
1304     metrics are printed to stdout.
1305
1306     Parameters

```

```

1307 -----
1308 test_ds : CipherStatisticsDataset
1309     The dataset used for prediction.
1310 model :
1311     The trained model to evaluate.
1312 args :
1313     The commandline arguments provided by the user.
1314 early_stopping_callback :
1315     Indicates whether the training was stopped before 'args.max_iter' was
1316     reached. Used together with 'train_iter' and 'args.max_iter' to control
1317     the number of prediction iterations.
1318 train_iter : int
1319     The number of iterations used until the model converged. Used together with
1320     'early_stopping_callback' and 'args.max_iter' to control the number of
1321     prediction iterations.
1322
1323 Returns
1324 -----
1325 str
1326     The statistics of this prediction run.
1327 """
1328
1329 print('Predicting test data...\n')
1330
1331 architecture = args.architecture
1332 start_time = time.time()
1333 total_len_prediction = 0
1334 cntr = 0
1335 test_iter = 0
1336 test_epoch = 0
1337
1338 # Determine the number of iterations to use for evaluating the model
1339 prediction_dataset_factor = 10
1340 if early_stopping_callback.stop_training:
1341     while test_ds.dataset_workers * test_ds.batch_size > train_iter / prediction_dataset_factor and prediction_dataset_factor > 1:
1342         prediction_dataset_factor -= 1
1343     args.max_iter = int(train_iter / prediction_dataset_factor)
1344 else:
1345     while test_ds.dataset_workers * test_ds.batch_size > args.max_iter / prediction_dataset_factor and prediction_dataset_factor > 1:
1346         prediction_dataset_factor -= 1
1347     args.max_iter /= prediction_dataset_factor
1348
1349 # Initialize 'PredictionPerformanceMetrics' instances for all classifiers. These
1350 # are used to save and evaluate the batched prediction results of the models.
1351 prediction_metrics = {}
1352 if architecture == "[FFNN,NB]":
1353     prediction_metrics = {"FFNN": PredictionPerformanceMetrics(model_name="FFNN"),
1354                          "NB": PredictionPerformanceMetrics(model_name="NB")}
1355 elif architecture == "[DT,ET,RF,SVM,kNN]":
1356     prediction_metrics = {"DT": PredictionPerformanceMetrics(model_name="DT"),
1357                          "ET": PredictionPerformanceMetrics(model_name="ET"),
1358                          "RF": PredictionPerformanceMetrics(model_name="RF"),
1359                          "SVM": PredictionPerformanceMetrics(model_name="SVM"),
1360                          "kNN": PredictionPerformanceMetrics(model_name="kNN").}
1361 else:
1362     prediction_metrics = {architecture: PredictionPerformanceMetrics(model_name=architecture)}
1363
1364 combined_batch = TrainingBatch("mixed", [], [])
1365 while test_ds.iteration < args.max_iter:
1366     testing_batches = next(test_ds)
1367
1368     # For architectures that only support one fit call: Sample all batches into one large batch.
1369     if architecture in ("DT", "RF", "ET", "SVM", "kNN", "SVM-Rotor", "[DT,ET,RF,SVM,kNN]"):
1370         for testing_batch in testing_batches:
1371             combined_batch.extend(testing_batch)
1372             if test_ds.iteration < args.max_iter:
1373                 print("Loaded %d ciphertexts." % test_ds.iteration)
1374             continue
1375         test_ds.stop_outstanding_tasks()
1376         print("Loaded %d ciphertexts." % test_ds.iteration)
1377         testing_batches = [combined_batch]
1378
1379     for testing_batch in testing_batches:
1380         statistics, labels = testing_batch.items()
1381
1382         # Decision Tree, Naive Bayes prediction
1383         if architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN"):
1384             prediction = model.predict_proba(statistics)
1385             prediction_metrics[architecture].add_predictions(labels, prediction)
1386         elif architecture == "SVM-Rotor":
1387             prediction = model.predict_proba(statistics)
1388             # add probability 0 to all aca labels that are missing in the prediction
1389             padded_prediction = []
1390             for p in list(prediction):
1391                 padded = [0] * 56 + list(p)
1392                 padded_prediction.append(padded)
1393             prediction_metrics[architecture].add_predictions(labels, padded_prediction)
1394         elif architecture == "[FFNN,NB]":
1395             prediction = model[0].predict(statistics, batch_size=args.batch_size, verbose=1)

```

```

1396         nb_prediction = model[1].predict_proba(statistics)
1397         prediction_metrics["FFNN"].add_predictions(labels, prediction)
1398         prediction_metrics["NB"].add_predictions(labels, nb_prediction)
1399     elif architecture == "[DT,ET,RF,SVM,kNN]":
1400         prediction = model[0].predict_proba(statistics)
1401         prediction_metrics["DT"].add_predictions(labels, prediction)
1402         prediction_metrics["ET"].add_predictions(labels, model[1].predict_proba(statistics))
1403         prediction_metrics["RF"].add_predictions(labels, model[2].predict_proba(statistics))
1404         prediction_metrics["SVM"].add_predictions(labels, model[3].predict_proba(statistics))
1405         prediction_metrics["kNN"].add_predictions(labels, model[4].predict_proba(statistics))
1406
1407     elif architecture == "FFNN" and isinstance(model, FFNN):
1408         preds, labels = predict_torch_ffnn(model, test_ds, args)
1409         # You may want to adapt this to your PredictionPerformanceMetrics usage:
1410         prediction_metrics = {architecture: PredictionPerformanceMetrics(model_name=architecture)}
1411         prediction_metrics[architecture].add_predictions(labels, preds)
1412         for metrics in prediction_metrics.values():
1413             metrics.print_evaluation()
1414         elapsed_prediction_time = datetime.fromtimestamp(time.time()) - datetime.fromtimestamp(start_time)
1415         prediction_stats = 'Prediction time: %d days %d hours %d minutes %d seconds.' % (
1416             elapsed_prediction_time.days, elapsed_prediction_time.seconds // 3600,
1417             (elapsed_prediction_time.seconds // 60) % 60,
1418             elapsed_prediction_time.seconds % 60)
1419         return prediction_stats
1420
1421     elif architecture == "LSTM" and isinstance(model, LSTM):
1422         preds, labels = predict_torch_lstm(model, test_ds, args)
1423         prediction_metrics = {architecture: PredictionPerformanceMetrics(model_name=architecture)}
1424         prediction_metrics[architecture].add_predictions(labels, preds)
1425         for metrics in prediction_metrics.values():
1426             metrics.print_evaluation()
1427         elapsed_prediction_time = datetime.fromtimestamp(time.time()) - datetime.fromtimestamp(start_time)
1428         prediction_stats = 'Prediction time: %d days %d hours %d minutes %d seconds.' % (
1429             elapsed_prediction_time.days, elapsed_prediction_time.seconds // 3600,
1430             (elapsed_prediction_time.seconds // 60) % 60,
1431             elapsed_prediction_time.seconds % 60)
1432         return prediction_stats
1433     else:
1434         prediction = model.predict(statistics, batch_size=args.batch_size, verbose=1)
1435         prediction_metrics[architecture].add_predictions(labels, prediction)
1436
1437         total_len_prediction += len(prediction)
1438         cntr += 1
1439         test_iter = args.train_dataset_size * cntr
1440         test_epoch = test_ds.epoch
1441         if test_epoch > 0:
1442             test_epoch = test_iter // ((test_ds.iteration + test_ds.batch_size * test_ds.dataset_workers) // test_ds.epoch)
1443             print("Prediction Epoch: %d, Iteration: %d / %d" % (test_epoch, test_iter, args.max_iter))
1444             if test_iter >= args.max_iter:
1445                 break
1446         if test_ds.iteration >= args.max_iter:
1447             break
1448
1449     test_ds.stop_outstanding_tasks()
1450     elapsed_prediction_time = datetime.fromtimestamp(time.time()) - datetime.fromtimestamp(start_time)
1451
1452     if total_len_prediction > args.train_dataset_size:
1453         total_len_prediction -= total_len_prediction % args.train_dataset_size
1454         print('\ntest data predicted: %d ciphertexts' % total_len_prediction)
1455
1456     # print prediction metrics
1457     for metrics in prediction_metrics.values():
1458         metrics.print_evaluation()
1459
1460     # print selected feature again
1461     if architecture == "SVM-Rotor":
1462         # print(f"RFE-support: \n{list(model.support_)}\n")
1463         # print(f"Support names: {convert_rfe_support_to_names(model.support_)}")
1464         # print(f"RFE-rank: \n{list(model.ranking_)}\n")
1465         print()
1466
1467     # print("GridSearchCV:")
1468     # print(f"Best score: {model.best_score_}")
1469     # print(f"Best params: {model.best_params_}")
1470
1471     prediction_stats = 'Prediction time: %d days %d hours %d minutes %d seconds with %d iterations and %d epochs.' % (
1472         elapsed_prediction_time.days, elapsed_prediction_time.seconds // 3600,
1473         (elapsed_prediction_time.seconds // 60) % 60,
1474         elapsed_prediction_time.seconds % 60, test_iter, test_epoch)
1475
1476     return prediction_stats
1477
1478 def expand_cipher_groups(cipher_types):
1479     """Turn cipher group identifiers (ACA, MTC3, ROTOR, ALL) into a list of their ciphers."""
1480     expanded = cipher_types
1481     if config.MTC3 in expanded:
1482         del expanded[expanded.index(config.MTC3)]
1483         for i in range(5):
1484             expanded.append(config.CIPHER_TYPES[i])

```

```

1485     elif config.ACA in expanded:
1486         del expanded[expanded.index(config.ACA)]
1487         for i in range(56):
1488             expanded.append(config.CIPHER_TYPES[i])
1489     elif config.ROTOR in expanded:
1490         del expanded[expanded.index(config.ROTOR)]
1491         for i in range(56, 61):
1492             expanded.append(config.CIPHER_TYPES[i])
1493     elif config.ALL in expanded:
1494         del expanded[expanded.index(config.ALL)]
1495         for i in range(61):
1496             expanded.append(config.CIPHER_TYPES[i])
1497     return expanded
1498
1499 def main():
1500     # Don't fork processes to keep memory footprint low.
1501     multiprocessing.set_start_method("spawn")
1502
1503     args = parse_arguments()
1504
1505     cpu_count = os.cpu_count()
1506     if cpu_count and cpu_count < args.dataset_workers:
1507         print("WARNING: More dataset_workers set than CPUs available.")
1508
1509     # Print arguments
1510     for arg in vars(args):
1511         print("{:23s}={:s}".format(arg, str(getattr(args, arg))))
1512
1513     args.plaintext_input_directory = os.path.abspath(args.plaintext_input_directory)
1514     args.rotor_input_directory = os.path.abspath(args.rotor_input_directory)
1515     args.ciphers = args.ciphers.lower()
1516     cipher_types = args.ciphers.split(',')
1517     architecture = args.architecture
1518     extend_model = args.extend_model
1519
1520     # Validate inputs
1521     if os.path.splitext(args.model_name)[1] not in ('.h5', '.pth'):
1522         print('ERROR: The model must have extension ".h5" (for Keras) or ".pth" (for PyTorch FFNN).', file=sys.stderr)
1523         sys.exit(1)
1524
1525     if extend_model is not None:
1526         if architecture not in ('FFNN', 'CNN', 'LSTM'):
1527             print('ERROR: Models with the architecture %s can not be extended!' % architecture,
1528                 file=sys.stderr)
1529             sys.exit(1)
1530         if len(os.path.splitext(extend_model)) != 2 or os.path.splitext(extend_model)[1] != '.h5':
1531             print('ERROR: The extended model name must have the ".h5" extension!', file=sys.stderr)
1532             sys.exit(1)
1533
1534     if architecture == "SVM-Rotor" and cipher_types[0] != "rotor":
1535         print(f"When training rotor-only model, the argument 'ciphers' "
1536             f"should equal 'rotor'. Selected ciphers are: '{cipher_types}'")
1537         sys.exit(1)
1538
1539     if args.train_dataset_size * args.dataset_workers > args.max_iter:
1540         print("ERROR: --train_dataset_size * --dataset_workers must not be bigger than --max_iter. "
1541             "In this case it was %d > %d" %
1542             (args.train_dataset_size * args.dataset_workers, args.max_iter),
1543             file=sys.stderr)
1544         sys.exit(1)
1545
1546     # Convert commandline cipher argument (all, aca, mtc3, rotor, etc.) to list of
1547     # all ciphers contained in the provided group. E.g. 'rotor' gets expanded
1548     # into 'enigma', 'm209', etc.
1549     cipher_types = expand_cipher_groups(cipher_types)
1550
1551     # Ensure plaintext dataset is available at 'args.plaintext_input_directory'.
1552     if should_download_plaintext_datasets(args):
1553         download_plaintext_datasets(args)
1554
1555     # Load the datasets for the requested cipher types. If aca and rotor cipher types
1556     # are contained in 'cipher_types', both plaintext and ciphertext datasets are loaded.
1557     train_ds, test_ds = load_datasets_from_disk(args, cipher_types)
1558
1559     # Get the number of cipher classes to predict. Since the label numbers are fixed,
1560     # it must be ensured that the output_layer_size of the neural networks contain
1561     # enough nodes upto the highest wanted class label.
1562     output_layer_size = max([config.CIPHER_TYPES.index(type) for type in cipher_types]) + 1
1563
1564     # Create a model and allow for distributed training on multi-GPU machines
1565     model, strategy = create_model_with_distribution_strategy(
1566         architecture, extend_model, output_layer_size=output_layer_size, max_train_len=args.max_train_len)
1567
1568     early_stopping_callback, train_iter, training_stats = train_model(model, strategy,
1569         args, train_ds)
1570
1571     save_model(model, args)
1572     prediction_stats = predict_test_data(test_ds, model, args, early_stopping_callback, train_iter)
1573

```

```

1574     print(training_stats)
1575     print(prediction_stats)
1576
1577 if __name__ == "__main__":
1578     main()

```

8.2 Full listing of eval.py

```

1 import multiprocessing
2 from pathlib import Path
3 import argparse
4 import random
5 import sys
6 import os
7 import pickle
8 import functools
9 import numpy as np
10 from datetime import datetime
11
12 import torch
13 import torch.nn.functional as F
14 import torch.optim as optim
15
16 # This environ variable must be set before all tensorflow imports!
17 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '3'
18 import tensorflow as tf
19 import tensorflow_datasets as tfds
20 from tensorflow.keras.optimizers import Adam
21 from tensorflow.keras.metrics import SparseTopK_categorical_accuracy
22 sys.path.append("../")
23 from util.utils import map_text_into_numberspace
24 from util.utils import print_progress
25 import cipherTypeDetection.config as config
26 from cipherTypeDetection.cipherStatisticsDataset import CipherStatisticsDataset, PlaintextPathsDatasetParameters,
27     RotorCipherTextsDatasetParameters, calculate_statistics, pad_sequences
28 from cipherTypeDetection.predictionPerformanceMetrics import PredictionPerformanceMetrics
29 from cipherTypeDetection.rotorDifferentiationEnsemble import RotorDifferentiationEnsemble
30 from cipherTypeDetection.ensembleModel import EnsembleModel
31 from cipherTypeDetection.transformer import MultiHeadSelfAttention, TransformerBlock, TokenAndPositionEmbedding
32 from util.utils import get_model_input_length
33 from cipherImplementations.cipher import OUTPUT_ALPHABET, UNKNOWN_SYMBOL_NUMBER
34 tf.debugging.set_log_device_placement(enabled=False)
35 # always flush after print as some architectures like RF need very long time before printing anything.
36 print = functools.partial(print, flush=True)
37
38 for device in tf.config.list_physical_devices('GPU'):
39     tf.config.experimental.set_memory_growth(device, True)
40
41 def str2bool(v):
42     return v.lower() in ("yes", "true", "t", "1")
43
44 def benchmark(args, model, architecture):
45     cipher_types = args.ciphers
46     args.plaintext_folder = os.path.abspath(args.plaintext_folder)
47     if args.dataset_size * args.dataset_workers > args.max_iter:
48         print("ERROR: --dataset_size * --dataset_workers must not be bigger than --max_iter. In this case it was %d > %d" % (
49             args.dataset_size * args.dataset_workers, args.max_iter), file=sys.stderr)
50     sys.exit(1)
51     if args.download_dataset and not os.path.exists(args.plaintext_folder) and args.plaintext_folder == os.path.abspath(
52         '../data/gutenberg_en'):
53         print("Downloading Datasets...")
54         tfds.download.add_checksums_dir('../data/checksums/')
55         download_manager = tfds.download.DownloadManager(download_dir='../data/', extract_dir=args.plaintext_folder)
56         download_manager.download_and_extract(
57             'https://drive.google.com/uc?id=1bF5sSVjxTx3DB-P5wxn87nxWnRhK_V&export=download')
58         path = os.path.join(args.plaintext_folder, 'ZIP_ucid_1bF5sSVjxTx-P5wxn87nxWn_V_export_downloadR9Cwhunev5CvJ-ic__'
59             'HawxhTtGOISdcCrr04fxfEI8A', os.path.basename(args.plaintext_folder))
60         dir_name = os.listdir(path)
61         for name in dir_name:
62             p = Path(os.path.join(path, name))
63             parent_dir = p.parents[2]
64             p.rename(parent_dir / p.name)
65         os.rmdir(path)
66         os.rmdir(os.path.dirname(path))
67         print("Datasets Downloaded.")
68
69     print("Loading Datasets...")
70     def validate_ciphertext_path(ciphertext_path, cipher_types):
71         file_name = Path(ciphertext_path).stem.lower()
72         if not file_name in cipher_types:
73             raise Exception(f"Filename must equal one of the expected cipher types. Expected cipher types are: {cipher_types}. Current
74                 filename is '{file_name}'.")
75
76     plaintext_files = []
77     dir_name = os.listdir(args.plaintext_folder)

```

```

77     for name in dir_name:
78         path = os.path.join(args.plaintext_folder, name)
79         if os.path.isfile(path):
80             plaintext_files.append(path)
81
82     def find_ciphertext_paths_in_dir(folder_path):
83         """Loads all .txt files in the given folder and checks that their names match the
84         known cipher types."""
85         file_names = os.listdir(folder_path)
86         result = []
87         for name in file_names:
88             path = os.path.join(folder_path, name)
89             file_name, file_type = os.path.splitext(name)
90             if os.path.isfile(path) and file_name.lower() in cipher_types and file_type == ".txt":
91                 validate_ciphertext_path(path, config.ROTOR_CIPHER_TYPES)
92
93             result.append(path)
94         return result
95
96     eval_rotor_ciphertext_paths = find_ciphertext_paths_in_dir(args.rotor_ciphertext_folder)
97
98     # Calculate batch size for rotor ciphers. The amount of samples per rotor cipher should be
99     # equal to the amount of samples per aca cipher.
100     number_of_rotor_ciphers = len(config.ROTOR_CIPHER_TYPES)
101     number_of_aca_ciphers = len(config.CIPHER_TYPES) - number_of_rotor_ciphers
102     amount_of_samples_per_cipher = args.dataset_size // number_of_aca_ciphers
103     rotor_train_dataset_size = amount_of_samples_per_cipher * number_of_rotor_ciphers
104
105     plaintext_dataset_params = PlaintextPathsDatasetParameters(cipher_types[:56],
106                                                                args.dataset_size,
107                                                                args.min_text_len,
108                                                                args.max_text_len,
109                                                                args.keep_unknown_symbols,
110                                                                args.dataset_workers,
111                                                                generate_evaluation_data=True)
112     rotor_dataset_params = RotorCiphertextsDatasetParameters(config.ROTOR_CIPHER_TYPES,
113                                                             rotor_train_dataset_size,
114                                                             args.dataset_workers,
115                                                             args.min_text_len,
116                                                             args.max_text_len,
117                                                             generate_evaluation_data=True)
118     dataset = CipherStatisticsDataset(plaintext_files, plaintext_dataset_params, eval_rotor_ciphertext_paths,
119                                     rotor_dataset_params, generate_evaluation_data=True)
120
121     if args.dataset_size % dataset.key_lengths_count != 0:
122         print("WARNING: the --dataset_size parameter must be dividable by the amount of --ciphers and the length configured KEY_LENGTHS in
123               ↪ config.py. The current key_lengths_count is %d" % dataset.key_lengths_count, file=sys.stderr)
124     print("Datasets loaded.\n")
125
126     print('Evaluating model...')
127
128     import time
129     start_time = time.time()
130     iteration = 0
131     epoch = 0
132     results = []
133     prediction_metrics = PredictionPerformanceMetrics(model_name=architecture)
134
135     while dataset.iteration < args.max_iter:
136         batches = next(dataset)
137
138         for index, batch in enumerate(batches):
139             statistics, labels, ciphertexts = batch.items()
140
141             if architecture == "FFNN":
142                 if hasattr(model, "evaluate"): # Keras model
143                     results.append(model.evaluate(statistics, labels, batch_size=args.batch_size, verbose=1))
144                 else: # PyTorch model
145                     stats_np = statistics.numpy()
146
147                     x = torch.tensor(stats_np, dtype=torch.float32)
148                     y = torch.tensor(labels.numpy(), dtype=torch.long)
149
150                     with torch.no_grad():
151                         outputs = model(x)
152                         loss = F.cross_entropy(outputs, y)
153                         top1 = torch.argmax(outputs, dim=1)
154                         acc = (top1 == y).float().mean()
155
156                     # Calc top-3
157                     top3 = torch.topk(outputs, k=3, dim=1).indices
158                     y_expanded = y.unsqueeze(1).expand_as(top3)
159                     k3_acc = (top3 == y_expanded).any(dim=1).float().mean()
160
161                     results.append((loss.item(), acc.item(), k3_acc.item()))
162
163             elif architecture in ("CNN", "LSTM", "Transformer"):
164                 results.append(model.evaluate(ciphertexts, labels, batch_size=args.batch_size, verbose=1))
165

```



```

166         elif architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN"):
167             results.append(model.score(statistics, labels))
168             print("accuracy: %f" % (results[-1]))
169         elif architecture == "Ensemble":
170             results.append(model.evaluate(statistics, ciphertexts, labels, args.batch_size, prediction_metrics, verbose=1))
171
172         iteration = dataset.iteration - len(batch) * (len(batches) - index - 1)
173         epoch = dataset.epoch
174         if epoch > 0:
175             epoch = iteration // (dataset.iteration // dataset.epoch)
176             print("Epoch: %d, Iteration: %d" % (epoch, iteration))
177             if iteration >= args.max_iter:
178                 break
179
180         if dataset.iteration >= args.max_iter:
181             break
182
183         elapsed_evaluation_time = datetime.fromtimestamp(time.time()) - datetime.fromtimestamp(start_time)
184         print('Finished evaluation in %d days %d hours %d minutes %d seconds with %d iterations and %d epochs.\n' % (
185             elapsed_evaluation_time.days, elapsed_evaluation_time.seconds // 3600, (elapsed_evaluation_time.seconds // 60) % 60,
186             elapsed_evaluation_time.seconds % 60, iteration, epoch))
187
188     if architecture in ("FFNN", "CNN", "LSTM", "Transformer"):
189         avg_loss = 0
190         avg_acc = 0
191         avg_k3_acc = 0
192         for loss, acc_pred, k3_acc in results:
193             avg_loss += loss
194             avg_acc += acc_pred
195             avg_k3_acc += k3_acc
196         avg_loss = avg_loss / len(results)
197         avg_acc = avg_acc / len(results)
198         avg_k3_acc = avg_k3_acc / len(results)
199         print("Average evaluation results: loss: %f, accuracy: %f, k3_accuracy: %f\n" % (avg_loss, avg_acc, avg_k3_acc))
200     elif architecture in ("DT", "NB", "RF", "ET", "Ensemble", "SVM", "kNN"):
201         avg_test_acc = 0
202         avg_k3_acc = 0
203         for acc, k3_acc in results:
204             avg_test_acc += acc
205             avg_k3_acc += k3_acc
206         avg_test_acc = avg_test_acc / len(results)
207         avg_k3_acc = avg_k3_acc / len(results)
208         print("Average evaluation results from %d iterations: avg_test_acc=%f, k3_accuracy: %f\n" % (iteration, avg_test_acc, avg_k3_acc))
209
210     print("Detailed results:")
211     prediction_metrics.print_evaluation()
212
213
214 def evaluate(args, model, architecture):
215     results_list = []
216     dir_name = os.listdir(args.data_folder)
217     dir_name.sort()
218     cntnr = 0
219     iterations = 0
220     for name in dir_name:
221         if iterations > args.max_iter:
222             break
223         path = os.path.join(args.data_folder, name)
224         if os.path.isfile(path):
225             if iterations > args.max_iter:
226                 break
227             batch = []
228             batch_ciphertexts = []
229             labels = []
230             results = []
231             dataset_cnt = 0
232             input_length = get_model_input_length(model, args.architecture)
233             with open(path, "rb") as fd:
234                 lines = fd.readlines()
235                 for line in lines:
236                     # remove newline
237                     line = line.strip(b'\n').decode()
238                     if line == '':
239                         continue
240                     split_line = line.split(' ')
241                     labels.append(int(split_line[0]))
242                     statistics = [float(f) for f in split_line[1].split(',')]
243                     batch.append(statistics)
244                     ciphertext = [int(j) for j in split_line[2].split(',')]
245                     if input_length is not None:
246                         if len(ciphertext) < input_length:
247                             ciphertext = pad_sequences([ciphertext], maxlen=input_length)[0]
248                             # if the length is too high, we need to strip it..
249                         elif len(ciphertext) > input_length:
250                             ciphertext = ciphertext[:input_length]
251                     batch_ciphertexts.append(ciphertext)
252                     iterations += 1
253                     if iterations == args.max_iter:
254                         break

```

```

255         if len(labels) == args.dataset_size:
256             if architecture == "FFNN":
257                 results.append(model.evaluate(tf.convert_to_tensor(batch), tf.convert_to_tensor(labels), args.batch_size, verbose
258                                     ↪=0))
259             elif architecture in ("CNN", "LSTM", "Transformer"):
260                 results.append(model.evaluate(tf.convert_to_tensor(batch_ciphertexts), tf.convert_to_tensor(labels),
261                                     ↪args.batch_size, verbose=0))
262             elif architecture == "Ensemble":
263                 results.append(model.evaluate(tf.convert_to_tensor(batch), tf.convert_to_tensor(batch_ciphertexts), tf.
264                                     ↪convert_to_tensor(labels),
265                                     ↪args.batch_size, verbose=0))
266             elif architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN"):
267                 results.append(model.score(batch, tf.convert_to_tensor(labels)))
268                 batch = []
269                 batch_ciphertexts = []
270                 labels = []
271                 dataset_cnt += 1
272         if len(labels) > 0:
273             if architecture == "FFNN":
274                 results.append(model.evaluate(tf.convert_to_tensor(batch), tf.convert_to_tensor(labels), args.batch_size, verbose=0))
275             elif architecture in ("CNN", "LSTM", "Transformer"):
276                 results.append(
277                     model.evaluate(tf.convert_to_tensor(batch_ciphertexts), tf.convert_to_tensor(labels), args.batch_size, verbose=0))
278             elif architecture == "Ensemble":
279                 results.append(
280                     model.evaluate(tf.convert_to_tensor(batch), tf.convert_to_tensor(batch_ciphertexts), tf.convert_to_tensor(labels),
281                                     ↪args.batch_size, verbose=0))
282             elif architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN"):
283                 results.append(model.score(batch, tf.convert_to_tensor(labels)))
284         if architecture in ("FFNN", "CNN", "LSTM", "Transformer"):
285             avg_loss = 0
286             avg_acc = 0
287             avg_k3_acc = 0
288             for loss, acc_pred, k3_acc in results:
289                 avg_loss += loss
290                 avg_acc += acc_pred
291                 avg_k3_acc += k3_acc
292             result = [avg_loss / len(results), avg_acc / len(results), avg_k3_acc / len(results)]
293         elif architecture in ("DT", "NB", "RF", "ET", "Ensemble", "SVM", "kNN"):
294             avg_test_acc = 0
295             for acc in results:
296                 avg_test_acc += acc
297             result = avg_test_acc / len(results)
298         results_list.append(result)
299         cnt += 1
300         if args.evaluation_mode == 'per_file':
301             if architecture in ("FFNN", "CNN", "LSTM", "Transformer"):
302                 print("%s (%d lines) test_loss: %f, test_accuracy: %f, test_k3_accuracy: %f (progress: %d%%)" % (
303                     os.path.basename(path), len(batch) + dataset_cnt * args.dataset_size, result[0], result[1], result[2], max(
304                         int(cnt / len(dir_name) * 100), int(iterations / args.max_iter * 100)))
305             elif architecture in ("DT", "NB", "RF", "ET", "Ensemble", "SVM", "kNN"):
306                 print("%s (%d lines) test_accuracy: %f (progress: %d%%)" % (
307                     os.path.basename(path), len(batch) + dataset_cnt * args.dataset_size, result,
308                     max(int(cnt / len(dir_name) * 100), int(iterations / args.max_iter * 100)))
309             else:
310                 print_progress("Evaluating files: ", cnt, len(dir_name), factor=5)
311         if iterations == args.max_iter:
312             break
313         if architecture in ("FFNN", "CNN", "LSTM", "Transformer"):
314             avg_test_loss = 0
315             avg_test_acc = 0
316             avg_test_acc_k3 = 0
317             for loss, acc, acc_k3 in results_list:
318                 avg_test_loss += loss
319                 avg_test_acc += acc
320                 avg_test_acc_k3 += acc_k3
321             avg_test_loss = avg_test_loss / len(results_list)
322             avg_test_acc = avg_test_acc / len(results_list)
323             avg_test_acc_k3 = avg_test_acc_k3 / len(results_list)
324             print("\n\nAverage evaluation results from %d iterations: avg_test_loss=%f, avg_test_acc=%f, avg_test_acc_k3=%f" % (
325                 iterations, avg_test_loss, avg_test_acc, avg_test_acc_k3))
326         elif architecture in ("DT", "NB", "RF", "ET", "Ensemble", "SVM", "kNN"):
327             avg_test_acc = 0
328             for acc in results_list:
329                 avg_test_acc += acc
330             avg_test_acc = avg_test_acc / len(results_list)
331             print("\n\nAverage evaluation results from %d iterations: avg_test_acc=%f" % (iterations, avg_test_acc))
332
333 def predict_single_line(args, model, architecture):
334     cipher_id_result = ''
335     ciphertexts = []
336     result = []
337     if args.ciphertext is not None:
338         ciphertexts.append(args.ciphertext.encode())
339     else:
340         ciphertexts = open(args.file, 'rb')
341
342     print()

```

```

343 for line in ciphertxts:
344     # remove newline
345     line = line.strip(b'\n')
346     if line == b'':
347         continue
348     # evaluate aca features file
349     # label = line.split(b' ')[0]
350     # statistics = ast.literal_eval(line.split(b' ')[1].decode())
351     # ciphertext = ast.literal_eval(line.split(b' ')[2].decode())
352     # print(config.CIPHER_TYPES[int(label.decode())], "length: %d" % len(ciphertext))
353
354     # Append ciphertext to itself. This improves the reliability of the results.
355     while len(line) < 1000:
356         line = line + line
357     # Limit line to at most 1000 characters to limit the execution time
358     line = line[:1000]
359
360     print(line)
361     ciphertext = map_text_into_numberspace(line, OUTPUT_ALPHABET, UNKNOWN_SYMBOL_NUMBER)
362     try:
363         statistics = calculate_statistics(ciphertext)
364     except ZeroDivisionError:
365         print("\n")
366         continue
367     results = None
368     if architecture == "FFNN":
369         result = model.predict(tf.convert_to_tensor([statistics]), args.batch_size, verbose=0)
370     elif architecture in ("CNN", "LSTM", "Transformer"):
371         input_length = get_model_input_length(model, architecture)
372         if len(ciphertext) < input_length:
373             ciphertext = pad_sequences([list(ciphertext)], maxlen=input_length)[0]
374         split_ciphertext = [ciphertext[input_length*j:input_length*(j+1)] for j in range(len(ciphertext) // input_length)]
375         results = []
376         if architecture in ("LSTM", "Transformer"):
377             for ct in split_ciphertext:
378                 results.append(model.predict(tf.convert_to_tensor([ct]), args.batch_size, verbose=0))
379         elif architecture == "CNN":
380             for ct in split_ciphertext:
381                 results.append(
382                     model.predict(tf.reshape(tf.convert_to_tensor([ct]), (1, input_length, 1)), args.batch_size, verbose=0))
383         result = results[0]
384         for res in results[1:]:
385             result = np.add(result, res)
386         result = np.divide(result, len(results))
387     elif architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN"):
388         result = model.predict_proba(tf.convert_to_tensor([statistics]))
389     elif architecture == "Ensemble":
390         result = model.predict(tf.convert_to_tensor([statistics]), [ciphertext], args.batch_size, verbose=0)
391
392     if isinstance(result, list):
393         result_list = list(result[0])
394     else:
395         result_list = result[0].tolist()
396     if results is not None and architecture not in ('Ensemble', 'LSTM', 'Transformer', 'CNN'):
397         for j in range(len(result_list)):
398             result_list[j] /= len(results)
399     if args.verbose:
400         for cipher in args.ciphers:
401             print("[:23s] {:.f}%".format(cipher, result_list[config.CIPHER_TYPES.index(cipher)]*100))
402         max_val = max(result_list)
403         cipher = config.CIPHER_TYPES[result_list.index(max_val)]
404     else:
405         max_val = max(result_list)
406         cipher = config.CIPHER_TYPES[result_list.index(max_val)]
407         print("[:s] {:.f}%".format(cipher, max_val * 100))
408     print()
409     cipher_id_result += cipher[0].upper()
410
411 if args.file is not None:
412     ciphertxts.close()
413
414 # return a list of probabilities (does only return the last one in case a file is used)
415 res_dict = {}
416 if len(result) != 0:
417     for j, val in enumerate(result[0]):
418         res_dict[args.ciphers[j]] = val * 100
419 return res_dict
420
421
422 def load_model(architecture, args, model_path, cipher_types):
423     strategy = args.strategy
424     model_list = args.models
425     architecture_list = args.architectures
426
427     model = None
428
429     if architecture == "FFNN" and model_path.endswith(".pth"):
430         from cipherTypeDetection.train import FFNN

```

```

432     checkpoint = torch.load(model_path, map_location=torch.device("cpu"))
433
434     model = FFNN(
435         input_size=checkpoint['input_size'],
436         hidden_size=checkpoint['hidden_size'],
437         output_size=checkpoint['output_size'],
438         num_hidden_layers=checkpoint['num_hidden_layers']
439     )
440     model.load_state_dict(checkpoint['model_state_dict'])
441     model.eval()
442
443     config.FEATURE_ENGINEERING = True
444     config.PAD_INPUT = False
445
446     return model
447
448     elif architecture in ("FFNN", "CNN", "LSTM", "Transformer"):
449         if architecture == "Transformer":
450             if not hasattr(config, "maxlen"):
451                 raise ValueError("maxlen must be defined in the config when loading a Transformer model!")
452             model = tf.keras.models.load_model(args.model, custom_objects={
453                 'TokenAndPositionEmbedding': TokenAndPositionEmbedding,
454                 'TransformerBlock': TransformerBlock})
455             else:
456                 model = tf.keras.models.load_model(args.model)
457             if architecture in ("CNN", "LSTM", "Transformer"):
458                 config.FEATURE_ENGINEERING = False
459                 config.PAD_INPUT = True
460             else:
461                 config.FEATURE_ENGINEERING = True
462                 config.PAD_INPUT = False
463                 optimizer = Adam(learning_rate=config.learning_rate, beta_1=config.beta_1, beta_2=config.beta_2, epsilon=config.epsilon,
464                                 amsgrad=config.amsgrad)
465                 model.compile(optimizer=optimizer, loss="sparse_categorical_crossentropy",
466                             metrics=["accuracy", SparseTopKCategoricalAccuracy(k=3, name="k3_accuracy")])
467         elif architecture in ("DT", "NB", "RF", "ET", "SVM", "kNN"):
468             config.FEATURE_ENGINEERING = True
469             config.PAD_INPUT = False
470             with open(model_path, "rb") as f:
471                 model = pickle.load(f)
472         elif architecture == "Ensemble":
473             cipher_indices = []
474             for cipher_type in cipher_types:
475                 cipher_indices.append(config.CIPHER_TYPES.index(cipher_type))
476             model = EnsembleModel(model_list, architecture_list, strategy, cipher_indices)
477         else:
478             raise ValueError("Unknown architecture: %s" % architecture)
479
480     # Controlla se ci sono cifrari rotor tra quelli richiesti
481     has_rotor_ciphers = any(c in config.ROTOR_CIPHER_TYPES for c in cipher_types)
482
483     # Se ci sono cifrari rotor, carica anche il modello rotor_only
484     if has_rotor_ciphers:
485         rotor_only_model_path = args.rotor_only_model
486         if not os.path.exists(rotor_only_model_path):
487             raise FileNotFoundError(f"Rotor-only model is required but not found at {rotor_only_model_path}")
488         with open(rotor_only_model_path, "rb") as f:
489             rotor_only_model = pickle.load(f)
490         return RotorDifferentiationEnsemble(architecture, model, rotor_only_model)
491
492     # Se non ci sono cifrari rotor:
493     # - se è un ensemble, restituisci direttamente l'ensemble
494     # - altrimenti restituisci il modello normale
495     return model
496
497
498 def expand_cipher_groups(cipher_types):
499     """Turn cipher group identifiers (ACA, MTC3) into a list of their ciphers"""
500     expanded = cipher_types
501     if config.MTC3 in expanded:
502         del expanded[expanded.index(config.MTC3)]
503         for i in range(5):
504             expanded.append(config.CIPHER_TYPES[i])
505     elif config.ACA in expanded:
506         del expanded[expanded.index(config.ACA)]
507         for i in range(56):
508             expanded.append(config.CIPHER_TYPES[i])
509     elif "aca+rotor" in expanded:
510         del expanded[expanded.index("aca+rotor")]
511         for i in range(61):
512             expanded.append(config.CIPHER_TYPES[i])
513     return expanded
514
515 def parse_arguments():
516     parser = argparse.ArgumentParser(
517         description='CANN Ciphertext Detection Neuronal Network Evaluation Script', formatter_class=argparse.RawTextHelpFormatter)
518     sp = parser.add_subparsers()
519     bench_parser = sp.add_parser('benchmark',
520                                  help='Use this argument to create ciphertexts on the fly, \nlike in training mode, and evaluate them with

```

```

521         'the model. \nThis option is optimized for large throughput to test the model.')
522 eval_parser = sp.add_parser('evaluate', help='Use this argument to evaluate cipher types for single files or directories.')
523 single_line_parser = sp.add_parser('single_line', help='Use this argument to predict a single line of ciphertext.')
524
525 parser.add_argument('--batch_size', default=128, type=int,
526                     help='Batch size for training.')
527 parser.add_argument('--max_iter', default=1000000000, type=int,
528                     help='the maximal number of iterations before stopping evaluation.')
529 parser.add_argument('--model', default='../data/models/ml.h5', type=str,
530                     help='Path to the model file. The file must have the .h5 extension.')
531 parser.add_argument('--rotor_only_model', default='../data/models/svm_rotor_only.h5', type=str,
532                     help='Path to rotor only model. This model is used in conjunction '
533                         'with the normal model in an ensemble to improve the recognition '
534                         'of rotor ciphers. The file must have the .h5 extension.')
535 parser.add_argument('--architecture', default='FFNN', type=str, choices=[
536     'FFNN', 'CNN', 'LSTM', 'DT', 'NB', 'RF', 'ET', 'Transformer', 'SVM', 'kNN', 'Ensemble'],
537                     help='The architecture to be used for training. \n'
538                         'Possible values are:\n'
539                         '- FFNN\n'
540                         '- CNN\n'
541                         '- LSTM\n'
542                         '- DT\n'
543                         '- NB\n'
544                         '- RF\n'
545                         '- ET\n'
546                         '- Transformer\n'
547                         '- SVM\n'
548                         '- kNN\n'
549                         '- Ensemble')
550 parser.add_argument('--ciphers', '--ciphers', default='aca', type=str,
551                     help='A comma separated list of the ciphers to be created.\n'
552                         'Be careful to not use spaces or use \' to define the string.\n'
553                         'Possible values are:\n'
554                         '- mtc3 (contains the ciphers Monoalphabetic Substitution, Vigenere,\n'
555                         '  Columnar Transposition, Plaisfair and Hill)\n'
556                         '- aca (contains all currently implemented ciphers from \n'
557                         '  https://www.cryptogram.org/resource-area/cipher-types/)\n'
558                         '- aca+rotor\n'
559                         '- simple_substitution\n'
560                         '- vigenere\n'
561                         '- columnar_transposition\n'
562                         '- playfair\n'
563                         '- hill\n')
564
565 parser.add_argument('--models', action='append', default=None,
566                     help='A list of models to be used in the ensemble model. The length of the list must be the same like the one in '
567                         'the --architectures argument.')
568 parser.add_argument('--architectures', action='append', default=None,
569                     help='A list of the architectures to be used in the ensemble model. The length of the list must be the same like '
570                         'the one in the --models argument.')
571 parser.add_argument('--strategy', default='weighted', type=str, choices=['mean', 'weighted'],
572                     help='The algorithm used for decisions.\n- Mean voting adds the probabilities from every class and returns the mean '
573                         'value of it. The highest value wins.\n- Weighted voting uses pre-calculated statistics, like for example '
574                         'precision, to weight the output of a specific model for a specific class.')
575 parser.add_argument('--dataset_size', default=16000, type=int,
576                     help='Dataset size per evaluation. This argument should be dividable \nby the amount of --ciphers.')
577
578 bench_parser.add_argument('--download_dataset', default=True, type=bool)
579 bench_parser.add_argument('--dataset_workers', default=1, type=int)
580 bench_parser.add_argument('--plaintext_folder', default='../data/gutenberg_en', type=str)
581 bench_parser.add_argument('--rotor_ciphertext_folder', default='../data/rotor_ciphertexts', type=str)
582 bench_parser.add_argument('--keep_unknown_symbols', default=False, type=bool)
583 bench_parser.add_argument('--min_text_len', default=50, type=int)
584 bench_parser.add_argument('--max_text_len', default=-1, type=int)
585
586 bench_group = parser.add_argument_group('benchmark')
587 bench_group.add_argument('--download_dataset', help='Download the dataset automatically.')
588 bench_group.add_argument('--dataset_workers', help='The number of parallel workers for reading the input files.')
589 bench_group.add_argument('--plaintext_folder', help='Input folder of the plaintexts.')
590 bench_group.add_argument('--keep_unknown_symbols', help='Keep unknown symbols in the plaintexts. Known \n'
591                     'symbols are defined in the alphabet of the cipher.')
592 bench_group.add_argument('--min_text_len', help='The minimum length of a plaintext to be encrypted in the evaluation process.\n'
593                     'If this argument is set to -1 no lower limit is used.')
594 bench_group.add_argument('--max_text_len', help='The maximum length of a plaintext to be encrypted in the evaluation process.\n'
595                     'If this argument is set to -1 no upper limit is used.')
596
597 eval_parser.add_argument('--evaluation_mode', nargs='?', choices=('summarized', 'per_file'), default='summarized', type=str)
598 eval_parser.add_argument('--data_folder', default='../data/gutenberg_en', type=str)
599
600 eval_group = parser.add_argument_group('evaluate')
601 eval_group.add_argument('--evaluation_mode',
602                         help='To create an single evaluation result over all iterated data files use the \'summarized\' option.\n'
603                             '\n This option is to be preferred over the benchmark option, if the tests should be reproducable.\n'
604                             '- To create an evaluation for every file use \'per_file\' option. This mode allows the \n'
605                             '  calculation of the \n - average value of the prediction \n'
606                             '  - lower quartile - value at the position of 25 percent of the sorted predictions\n'
607                             '  - median - value at the position of 50 percent of the sorted predictions\n')

```

```

608         ' - upper quartile - value at the position of 75 percent of the sorted predictions\n'
609         ' With these statistics an expert can classify a ciphertext document to a specific cipher.')
```

```

610     eval_group.add_argument('--data_folder', help='Input folder of the data files with labels and calculated features.')
```

```

611
612     single_line_parser.add_argument('--verbose', default=True, type=bool)
613     data = single_line_parser.add_mutually_exclusive_group(required=True)
614     data.add_argument('--ciphertext', default=None, type=str)
615     data.add_argument('--file', default=None, type=str)
616
617     single_line_group = parser.add_argument_group('single_line')
618     single_line_group.add_argument('--ciphertext', help='A single line of ciphertext to be predicted by the model.')
619     single_line_group.add_argument('--file', help='A file with lines of ciphertext to be predicted line by line by the model.')
620     single_line_group.add_argument('--verbose', help='If true all predicted ciphers are printed. \n'
621                                     'If false only the most accurate prediction is printed.')
```

```

622
623     return parser.parse_args()
624
625 def main():
626     multiprocessing.set_start_method("spawn")
627
628     args = parse_arguments()
629
630     for arg in vars(args):
631         print("{:23s} = {}".format(arg, str(getattr(args, arg))))
632     m = os.path.splitext(args.model)
633     if os.path.splitext(args.model)[1] not in ('.h5', '.pth'):
634         print('ERROR: The model must have extension ".h5" (for Keras) or ".pth" (for PyTorch FFNN).', file=sys.stderr)
635         sys.exit(1)
636
637     architecture = args.architecture
638     model_path = args.model
639     args.ciphers = args.ciphers.lower()
640     cipher_types = args.ciphers.split(',')
641     args.ciphers = expand_cipher_groups(cipher_types)
642     if architecture == 'Ensemble':
643         if not hasattr(args, 'models') or not hasattr(args, 'architectures'):
644             raise ValueError("Please use the 'ensemble' subroutine if specifying the ensemble architecture.")
645         if len(args.models) != len(args.architectures):
646             raise ValueError("The length of --models must be the same like the length of --architectures.")
647         models = []
648         for i in range(len(args.models)):
649             model = args.models[i]
650             arch = args.architectures[i]
651             if not os.path.exists(os.path.abspath(model)):
652                 raise ValueError("Model in %s does not exist." % os.path.abspath(model))
653             if arch not in ('FFNN', 'CNN', 'LSTM', 'DT', 'NB', 'RF', 'ET', 'Transformer', 'SVM', 'kNN'):
654                 raise ValueError("Unallowed architecture %s" % arch)
655             if arch in ('FFNN', 'CNN', 'LSTM', 'Transformer') and not os.path.abspath(model).endswith('.h5'):
656                 raise ValueError("Model names of the types %s must have the .h5 extension." % ['FFNN', 'CNN', 'LSTM', 'Transformer'])
657     elif args.models is not None or args.architectures is not None:
658         raise ValueError("It is only allowed to use the --models and --architectures with the Ensemble architecture.")
659
660     print("Loading Model...")
661     # There are some problems regarding the loading of models on multiple GPU's.
662     # gpu_count = len(tf.config.list_physical_devices('GPU'))
663     # if gpu_count > 1:
664     #     strat = tf.distribute.MirroredStrategy()
665     #     with strat.scope():
666     #         model = load_model()
667     # else:
668     #     model = load_model()
669     model = load_model(architecture, args, model_path, cipher_types)
670     print("Model Loaded.")
671
672     # Model is now always an ensemble
673     #architecture = "Ensemble"
674
675     # the program was started as in benchmark mode.
676     if args.download_dataset is not None:
677         benchmark(args, model, architecture)
678     # the program was started in single_line mode.
679     elif args.ciphertext is not None or args.file is not None:
680         predict_single_line(args, model, architecture)
681     # the program was started in prediction mode.
682     else:
683         evaluate(args, model, architecture)
684
685 if __name__ == "__main__":
686     main()

```