

XÂY DỰNG HỆ THỐNG PHÁT HIỆN TIN GIẢ MẠO DẠNG VĂN BẢN BẰNG PHƯƠNG PHÁP HỌC SÂU BUILDING THE SYSTEM FOR DETECTING FAKE INFORMATION IN THE TEXT FORM BY DEEP LEARNING METHOD

NGUYỄN TRƯỜNG VƯƠNG¹, CÙ VĨNH LỘC²

¹ Khoa Công nghệ Thông tin, Trường Đại học SPKT Vĩnh Long;

² Khoa Công nghệ Thông tin và Truyền Thông, Trường Đại học Cần Thơ;

Email: tr.vuong.vl@gmail.com; cvloc@cit.ctu.edu.vn

Nhận bài (Received): 07/5/2022; Phản biện (Reviewed): 17/5/2022; Chấp nhận (Accepted): 01/7/2022

TÓM TẮT:

Trong bài báo này, chúng tôi trình bày cách tiếp cận đối với việc phát hiện tin tức giả mạo dạng văn bản bằng cách sử dụng phương pháp học sâu kết hợp với mô hình huấn luyện trước dành riêng cho tiếng Việt (PhoBert) để nhận dạng và phát hiện tin giả mạo. Các mô hình đề xuất của chúng tôi được huấn luyện và đánh giá trên bộ dữ liệu tiếng Việt [1]. Từ dữ liệu đầu vào qua các bước xử lý dữ liệu đầu vào, chúng tôi cần nhúng từ (word embedding) để chuyển đổi dữ liệu văn bản sang mô hình vector. Sau đó vector này sẽ đi qua mô hình biến đổi PhoBert để có được các vector đặc trưng làm đầu ra. Vector đầu ra này là đầu vào của một mạng truyền thẳng MLP (Multi-layer Perceptron) để sử dụng cho quá trình phân lớp. Chúng tôi sử dụng hàm softmax() cho quá trình phân lớp. Hàm này trả ra xác suất trên hai tập nhãn fake hoặc real. Trong bài báo này chúng tôi sử dụng ba mô hình khác nhau như CNN, LSTM và PhoBert để thực hiện huấn luyện mô hình trên cùng một tập dữ liệu, sau đó so sánh kết quả đạt được và độ chính xác của từng mô hình để từ đó lựa chọn mô hình tốt nhất, có độ chính xác nhất để xây dựng hệ thống phát hiện tin giả mạo. Kết quả thử nghiệm cho thấy mô hình PhoBert đã đạt được kết quả vượt trội với 89% độ chính xác phát hiện tin giả mạo.

Từ khóa: Phát hiện tin giả mạo, CNN, LSTM, PhoBert, MLP.

ABSTRACT:

In this paper, we present an approach to textual fake information detection using deep learning method combined with a Vietnamese-specific pre-training model (PhoBert) to recognize and detect fake information. Our proposed models are tested and evaluated on Vietnamese dataset [1]. From input data through input data processing steps, we need to apply word embedding to convert text data into vector model. Then this vector will go through the PhoBert transform model to get the feature vectors as output. This output vector is the input of an MLP (Multi-layer Perceptron) feedforward network to be used for the classification process. The softmax() function for the classification process is used. This function returns the probability on two sets of fake or real labels. In this paper, three different models such as CNN, LSTM and PhoBert are used to perform model training on the same data set, then compare the obtained results and the accuracy of each model to

then select the best and most accurate model to build a fake information detection system. The test results show that the PhoBert model has achieved outstanding results with 89% accuracy in detecting fake information.

Keywords: Fake information detection, CNN, LSTM, PhoBert, MLP.

1. Giới thiệu

Phát hiện tin giả mạo có vai trò quan trọng trong việc kiểm tra tính đúng đắn của thông tin cũng như ngăn chặn những hậu quả xấu mà nó gây ra. Trên thế giới cũng có nhiều nghiên cứu về phát hiện tin giả của nhiều nhóm tác giả như: Ahmed và cộng sự [2] đã sử dụng TF-IDF (Tần số tài liệu nghịch đảo thuật ngữ) như một phương pháp ngoại trừ tính năng với các mô hình học máy khác nhau. Các thí nghiệm mở rộng đã thực hiện với LR (Mô hình hồi quy tuyến tính) và thu được độ chính xác là 89,00%. Ghanem và cộng sự [3] đã áp dụng phương pháp nhúng từ khác nhau, bao gồm các tính năng n-gram để phát hiện lập trường của các bài báo giả mạo, họ đã thu được độ chính xác là 48,80%. Ruchansky và cộng sự [4] đã sử dụng mô hình sâu để phân loại tin giả. Họ đã sử dụng mối quan hệ người dùng tin tức như một yếu tố cần thiết và đạt được độ chính xác là 89,20%. Kaliyar [5] và cộng sự đã trình bày một hệ thống sử dụng sự kết hợp của ba khối song song của Mạng thần kinh kết hợp một lớp (CNN) cùng với BERT để đạt được độ chính xác là 98,90% trên dữ liệu thử nghiệm. Hiện nay những công trình nghiên cứu phát hiện tin giả trên tiếng Việt còn rất ít. Xuất phát từ nhu cầu cuộc sống hằng ngày về tìm hiểu tin tức, trau dồi kiến thức và kế thừa những thành công của những công trình nghiên cứu trong và ngoài nước, chúng tôi đề xuất mô hình "xây dựng hệ thống phát hiện tin giả mạo dạng văn bản bằng phương pháp học sâu" nhằm giúp cho độc giả phát hiện và ngăn chặn sự phát tán của tin giả mạo. Trong bài báo này, chúng tôi sử dụng mô hình PhoBert làm bộ mã hóa câu, ưu điểm

là có thể lấy chính xác trình bày ngữ cảnh của một câu. Công việc này trái ngược với các công trình nghiên cứu trước đây là chỉ xem xét một chuỗi văn bản theo một hướng theo cách (từ trái sang phải hoặc từ phải sang trái để đào tạo trước). Chúng tôi sử dụng bộ dữ liệu của VFND [1] để huấn luyện và phát hiện tin giả mạo. VFND là bộ dữ liệu chứa 223 bản dữ liệu các bài báo và tên miền của các trang đã đăng các bài báo đó và được gán nhãn 0 là tin thật, 1 là tin giả về các tin tức thật và tin tức giả bằng ngôn ngữ tiếng Việt được tập hợp trong khoảng thời gian từ 2017 đến nay. Các tin tức được đưa vào đây được phân loại tin thật hoặc tin giả dựa trên một số nguồn tin, tham chiếu chéo đến các nguồn tin được dẫn hoặc được phân loại bởi cộng đồng. Các tin tức sử dụng trong bộ dữ liệu đều là tin tức tường thuật về một sự kiện. Các chủ đề mà bộ dữ liệu tập trung là: Thể thao, Văn hóa, Xã hội, Kinh tế, Pháp luật. Các tin tức sẽ được kiểm tra chéo về nguồn gốc, nội dung, sự kiện để xác định thật và giả.

2. Cơ sở lý thuyết

2.1. Các đặc trưng của tin giả [6]

Tin giả thường có tiêu đề hấp dẫn, thu hút người đọc, viết in hoa kèm dấu ký tự mang tính chất khẳng định, thường có vẻ khó xảy ra trong thực tế nhưng được viết dưới dạng khẳng định để thu hút sự chú ý của người đọc. Tin giả thường không được chú trọng về cấu trúc ngữ pháp, thể thức văn bản, dễ có lỗi chính tả và ngữ pháp, không thống nhất. Hình ảnh sử dụng trong bài viết đa phần là ảnh trên mạng hoặc được chỉnh sửa cho phù hợp với nội dung nguồn tin. Về mốc thời gian sử dụng trong bài viết, tin giả được biên soạn và định dạng mốc thời

gian không trùng với thực tế. Về các luận cứ, luận chứng trong bài viết, thông thường các tin giả được tạo ra được dựa trên một câu chuyện, tình tiết có thực nhưng được làm giả ở những nội dung quan trọng nhất.

2.2. Mô hình biến đổi (Transformers)

Cấu trúc biến đổi (transformer) [7] dựa trên “sự tự chú ý” (self attention) để tính toán sự biểu diễn của dữ liệu đầu vào mà không sử dụng cấu trúc mạng neural hồi quy (recurrent neural networks - RNN). Cấu trúc biến đổi dựa trên bộ mã hoá và giải mã (encoder - decoder) kết hợp với cơ chế “tự chú ý (self attention)”, và các lớp mạng neural truyền thẳng. Cơ chế “chú ý (attention)” cho phép cấu trúc biến đổi tính toán sự ánh xạ của một truy vấn (query) và một tập các từ khoá - giá trị (key - value) cho đầu ra. Sau đó, đầu ra được tính toán bằng cộng có trọng số của các giá trị.

2.3. Mô hình PhoBert

PhoBert [8] là mô hình huấn luyện trước, đặc biệt chỉ huấn luyện dành riêng cho tiếng Việt và dựa trên sự biến đổi (transformer), cho phép biểu diễn ngữ cảnh của một từ bằng cách dựa trên mối quan hệ của từ đó với các từ xung quanh [19]. Tương tự như Bert [9], PhoBert khác biệt với các mô hình một chiều (unidirectional) khi chỉ học các biểu diễn từ trái qua phải hoặc từ phải qua trái. PhoBert được sử dụng để huấn luyện mô hình ngôn ngữ mặt nạ (masked language model) bằng cách học hai bài toán cùng một lúc là bài toán dự đoán từ và bài toán dự đoán câu. Với bài toán dự đoán từ, các từ trong một câu sẽ được che giấu (masked). Quá trình huấn luyện sẽ dự đoán từ bị che dấu bằng cách dựa vào các từ xung quanh. PhoBert cũng có hai phiên bản là PhoBert base với 12 transformers block và PhoBert large với 24 transformers block. Trong nghiên cứu này, bài báo thử nghiệm với mô hình

PhoBert base. Bài báo sử dụng bpe của mô hình để encode một tin tức thành một danh sách các 22 subword. Mô hình có dict chứa từ điển sẵn có của PhoBert. Bài báo sẽ sử dụng từ điển này để giúp ánh xạ ngược từ subword về id của nó trong bộ từ vựng được cung cấp sẵn. Xây dựng mô hình huấn luyện PhoBert có hai lựa chọn là Fairseq và Transformer. Ở đây bài báo lựa chọn thử nghiệm với Transformer [9] và sử dụng BertForSequenceClassification để đào tạo mô hình.

2.4. Hàm softmax()

Hàm softmax dùng để tính toán xác suất xảy ra của một sự kiện. Nói một cách khái quát, hàm softmax sẽ tính khả năng xuất hiện của một class trong tổng số tất cả các class có thể xuất hiện. Sau đó, xác suất này sẽ được sử dụng để xác định class mục tiêu cho các input. Cụ thể, hàm softmax biến vector k chiều có các giá trị thực bất kỳ thành vector k chiều có giá trị thực có tổng bằng 1. Giá trị nhập có thể dương, âm, bằng 0 hoặc lớn hơn 1, nhưng hàm softmax sẽ luôn biến chúng thành một giá trị nằm trong khoảng $(0;1]$. Như vậy, chúng có thể được gọi là “xác suất”. Nếu một trong các giá trị nhập rất nhỏ hoặc âm, hàm softmax biến chúng thành 1 xác suất nhỏ. Còn nếu một giá trị nhập lớn thì nó sẽ được chuyển thành một xác suất lớn. Nhưng xác suất luôn lớn hơn 0 và nhỏ hơn 1, hoặc bằng 1. Hàm Softmax được sử dụng trong mô hình hồi quy logistic phân loại. Trong xây dựng neural network (mạng thần kinh nhân tạo), hàm softmax được sử dụng trong nhiều lớp.

2.5. Mô hình phân lớp

Trong phạm vi bài báo, chúng tôi chỉ đề cập đến mạng perceptron đa lớp (MLP). Về cơ bản, MLP có tối thiểu 3 lớp: một lớp ngõ vào, một lớp ngõ ra và một hoặc nhiều lớp ẩn. Trong đó, lớp ngõ vào tiếp nhận

dữ liệu vào mạng neuron, lớp ẩn (hidden layer) dùng để tính toán các thông số dựa trên dữ liệu từ lớp ngõ vào và chuyển tiếp kết quả đến lớp tiếp theo và lớp ngõ ra thực hiện đánh giá và đưa ra kết quả dựa trên các thuật toán thích hợp. Các mô hình mạng MLP sử dụng ngõ ra lớp trước làm ngõ vào lớp sau, không có sự phản hồi từ lớp ngõ ra về lại lớp ngõ vào được gọi là mạng *feed-forward* và sự liên kết các neuron giữa các lớp được gọi là *fullyconnected* hay mạng kết nối hoàn toàn. Mạng fully connected có ưu điểm là đơn giản, dễ triển khai tuy

nhien nhược điểm là số lượng thông số và số lượng phép tính trong mạng rất lớn. Chúng tôi sử dụng hàm *softmax()* cho quá trình phân lớp. Hàm này trả ra xác suất trên hai tập nhãn (fake hoặc real).

3. Thực nghiệm và kết quả đạt được

3.1. Phương pháp sử dụng mạng nơ-ron

Với phương pháp sử dụng mạng nơ-ron như: CNN, LSTM, RNN,... thì quá trình phân loại dữ liệu dạng văn bản cũng gồm hai giai đoạn:

- Giai đoạn huấn luyện:



Hình 3.1: Mô hình giai đoạn huấn luyện sử dụng mạng nơ-ron

- Giai đoạn phân lớp:

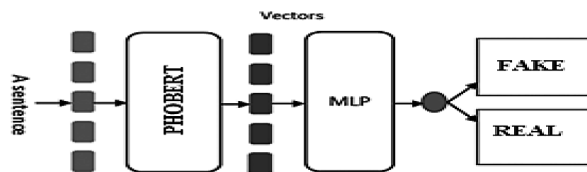


Hình 3.2: Mô hình giai đoạn phân lớp sử dụng mạng nơ-ron

3.2. Mô hình

Chúng tôi giới thiệu mô hình phát hiện tin giả mạo dựa trên PhoBERT ở Hình 3.3. Sau khi thực hiện các bước xử lý dữ liệu đầu vào, chúng tôi cần nhúng từ (word embedding) để chuyển đổi dữ liệu văn bản sang mô hình vector. Các

vector này sẽ được biến đổi bằng PhoBERT để có được một vector đầu ra. Vector này sẽ là đầu vào của một mạng truyền thẳng (feed-forward network). Mạng này sẽ cho ra một vector cuối cùng cho phân lớp. Kết quả ở bộ phân lớp cho biết tin tức đó là fake hoặc real



Hình 3.3. Mô hình phân loại tin tức sử dụng PhoBERT

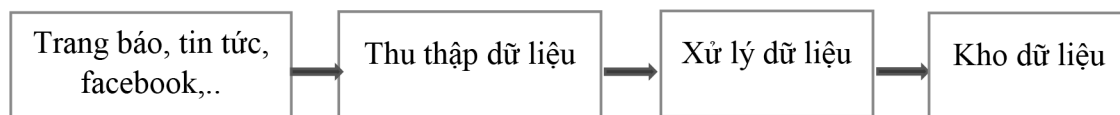
3.3. Huấn luyện

Chúng tôi sử dụng mô hình PhoBERT, CNN, LSTM cho quá trình huấn luyện. Sau khi huấn luyện sẽ so sánh độ chính xác của ba mô hình để từ đó đưa ra nhận xét mô hình tốt nhất. Quá trình huấn luyện trong 50 lần lặp.

3.4. Dữ liệu và cấu hình thực nghiệm

3.4.1. Dữ liệu

- Xây dựng dữ liệu: Việc thực hiện xây dựng kho dữ liệu bài báo đã thực hiện theo từng giai đoạn trong mô hình dưới đây:



Hình 3.4: Mô hình xây dựng kho dữ liệu

- Tiền xử lý dữ liệu đầu vào: Là bước rất quan trọng và cần làm. Việc xử lý dữ liệu đầu vào là quá trình chuẩn hóa dữ liệu và loại bỏ các thành phần không có ý nghĩa cho việc phân loại tin tức dạng văn bản. Tiền xử lý dữ liệu tiếng Việt cho bài toán phân loại tin tức dạng văn bản gồm các bước sau: Xóa HTML code (nếu có), chuẩn hóa bảng mã Unicode (đưa về Unicode tổ hợp dựng sẵn), thực hiện tách từ tiếng Việt (ở đây tác giả sử dụng thư viện tách từ pyvi), chuyển đổi chữ hoa thành chữ thường, xóa các ký tự đặc biệt: “.”, “:”, “&”,... Một điểm đặc biệt trong văn bản tiếng Việt đó là một từ có thể được kết hợp bởi nhiều tiếng khác nhau, ví dụ như: sử_dụng, bắt_đầu, ... khác với tiếng Anh và một số ngôn ngữ khác, các từ được phân cách nhau bằng khoảng trắng: use some examples, i love you... Vì vậy chúng tôi cần tách từ để có thể đảm bảo ý nghĩa của từ được toàn vẹn.

3.4.2 Cấu hình thực nghiệm

Chúng tôi sử dụng phương pháp k -fold cross-validation, tức là bộ dữ liệu sẽ được chia thành k phần bằng nhau, sau đó mô hình sẽ được lần lượt huấn luyện trên $k-1$ phần và kiểm tra trên phần còn lại. Quá trình lặp lại k lần. Kết quả cuối cùng của mô hình là trung bình cộng điểm ROUGE trên toàn bộ các phần.

Hệ thống được cài đặt bằng ngôn ngữ Python và chạy trên cùng một môi trường Google Colab miễn phí, cấu hình RAM 12.5 GB và dùng GPU. Thư viện hỗ trợ huấn luyện mô hình mạng là tensorflow-addons-0.16.1

3.4.3. Độ đo chính xác

- Để đánh giá kết quả của việc xác định thực thể và thuộc tính ta đánh giá thông qua độ chính xác (precision), độ bao phủ (recall) và F1. Để so sánh giá trị giữa các mô hình, bài báo dùng các chỉ số đánh giá như sau:

True Positive (TP): dự đoán các mẫu tin tức giả thực sự được chú thích là tin tức giả;

True Negative (TN): dự đoán các mẫu tin tức thực sự được chú thích là tin tức thực sự;

False Negative (FN): dự đoán các mẫu tin tức thực sự được chú thích là tin giả

False Positive (FP): dự đoán các mẫu tin tức giả thực sự được chú thích là tin tức thật

Precision

$$Precision = \frac{|TP|}{|TP| + |FP|}$$

Accuracy

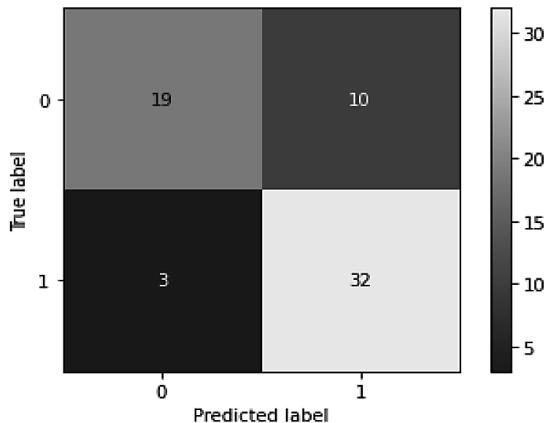
$$Accuracy = \frac{|TP| + |TN|}{|TP| + |TN| + |FP| + |FN|}$$

3.4.4. Kết quả thực nghiệm

Chúng tôi so sánh kết quả mô hình đề xuất với các mô hình khác trên cùng một bộ dữ liệu. Các mô hình gồm:

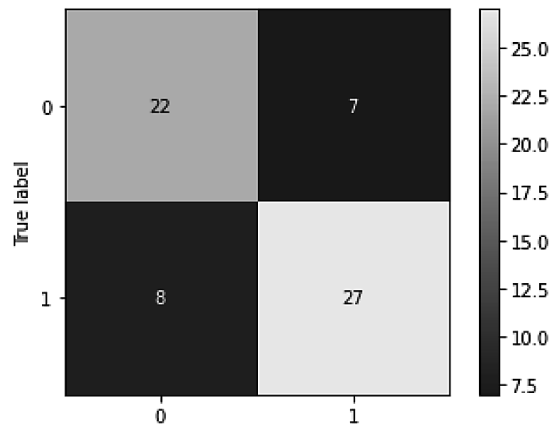
- Mô hình CNN(convolutional neural networks): Các văn bản đầu vào sẽ được tách từ thông qua thư viện pyvi và được làm sạch, tách bỏ stopword, xóa bỏ các ký tự đặc biệt, dấu câu. Tiếp theo chúng tôi sẽ chuyển văn bản sang không gian vector. Mỗi từ sẽ đưa về dạng vector chỉ bao gồm một số 1 và các thành phần còn lại bằng 0. Khi sử dụng mô hình CNN trong xử lý ngôn

ngữ tự nhiên (NLP), chúng tôi sử dụng bộ lọc có giá trị chiều rộng bằng với chiều rộng ma trận đầu vào, sau khi có đầu vào của mô hình là một ma trận được chuyển đổi từ số lượng đặc trưng sau khi được trích chọn. Trải qua các phép tích chập, đầu ra của mô hình là kết quả dự đoán tin thật hay tin giả (real hoặc fake).



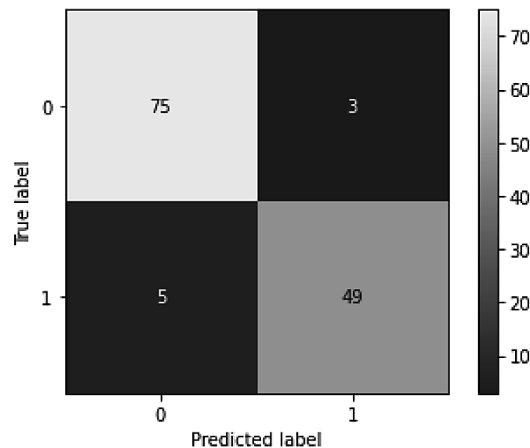
Hình 3.5: Độ chính xác phân lớp của mô hình CNN qua các lần lặp (Epochs)

- Mô hình LSTM (long-term dependencies): Cũng tương tự như mô hình CNN. Sau khi xử lý, chuẩn hóa dữ liệu đầu vào, chúng tôi tiến hành Word Embedding để đưa các từ trong câu về dạng vector để mô hình toán có thể hiểu được và có khả năng miêu tả được mối liên hệ, sự tương đồng về mặt ngữ nghĩa. Cụ thể là từ dạng text, các từ sẽ được chuyển về dạng vector đặc trưng để đưa vào mô hình LSTM. Khi một tin tức được đưa vào, trước tiên nó sẽ được nhúng theo số chỉ số tương ứng của nó trong từ điển. Sau đó, dựa trên từ điển và kết quả Word2vector thu được embedding toàn bộ câu dưới dạng ma trận. Trong mô hình này max_seq_len là độ dài của câu, tất cả các câu trong tập huấn luyện đều được cắt hoặc nối để có cùng độ dài max_seq_len.



Hình 3.6: Độ chính xác phân lớp của mô hình LSTM qua các lần lặp (Epochs)

- Mô hình PhoBert: Sau khi thực hiện các bước xử lý dữ liệu đầu vào, chúng tôi cần nhúng từ (word embedding) để chuyển đổi dữ liệu văn bản sang mô hình vector. Sau đó vector này sẽ đi qua mô hình biến đổi PhoBert để có được các vector đặc trưng làm đầu ra. Vector đầu ra này là đầu vào của một mạng truyền thẳng MLP (Multi-layer Perceptron) để sử dụng cho quá trình phân lớp. Chúng tôi sử dụng hàm softmax() cho quá trình phân lớp. Hàm này trả ra xác suất trên hai tập nhãn fake hoặc real. Cụ thể sẽ bao gồm các bước như sau: Xử lý dữ liệu đầu vào và nhúng từ (word embedding) để chuyển đổi dữ liệu văn bản sang mô hình vector. Phân đoạn văn bản trước khi đưa vào mô hình PhoBert (do PhoBert yêu cầu). Mã hóa dữ liệu bằng bộ mã hóa của PhoBert. Chú ý rằng khi mã hóa ta sẽ thêm hai token đặc biệt là [CLS] và [SEP] vào đầu và cuối câu. Đưa bản tin đã được mã hóa vào mô hình kèm theo chú ý mask. Lấy dữ liệu đầu ra và lấy vector đầu ra đầu tiên (chính là ở vị trí token đặc biệt [CLS]) làm đặc trưng cho câu để đào tạo hoặc để nhận dạng. Đưa vector đặc trưng thu được từ bước trên vào mô hình phân lớp MLP và sử dụng hàm softmax() để trả ra kết quả dự đoán cho quá trình phân lớp là fake hoặc real.



Hình 3.7: Độ chính xác phân lớp của mô hình PhoBERT kết hợp mạng MLP qua các lần lặp (Epochs)

Bảng 1: So sánh kết quả phân loại nhị phân (binary) của ba mô hình

(Chữ đậm thể hiện mô hình tốt nhất, chữ nghiêng thể hiện mô hình tốt thứ hai.)

Mô hình	PRECISION(%)		RECALL(%)		F1(%)	
	Tin thật	Tin giả	Tin thật	Tin giả	Tin thật	Tin giả
CNN	<i>0.76</i>	<i>0.82</i>	<i>0.91</i>	<i>0.66</i>	<i>0.83</i>	<i>0.75</i>
LSTM	0.79	0.73	0.77	0.76	0.78	0.75
PhoBERT + MLP	0.94	0.84	0.86	0.93	0.90	0.89

Từ bảng kết quả nhận thấy với độ đo F1 thì mô hình PhoBERT cho kết quả tốt nhất 90% với tin thật và 89% với tin giả. Kết quả từ Bảng 1 cho thấy mô hình của chúng tôi cho kết quả khả quan hơn hai mô hình còn lại. Trên mô hình PhoBERT đạt độ chính xác 94% với tin thật và 84% với tin giả. Độ chính xác của tin thật cao hơn 18% khi so sánh với mô hình CNN và cao hơn 15% đối với mô hình LSTM, tương tự độ chính xác của tin giả cao hơn 2% khi so sánh với mô hình CNN và 9% khi so sánh với mô hình LSTM. Điều này là do mô hình tận dụng sức mạnh của PhoBERT, được huấn luyện trên bộ dữ liệu rất lớn. Quá trình huấn luyện cho phép biểu diễn yếu tố ngữ cảnh của từ trong câu. Điều này cho phép mô hình chúng tôi có thể dự đoán đúng những câu quan trọng.

4. Kết luận

Trong bài báo này, để có cái nhìn tổng quan và đánh giá chương trình. Chúng tôi đã sử dụng ba mô hình là CNN, LSTM và PhoBERT để huấn luyện mô hình trên cùng một tập dữ liệu tiếng Việt. Sử dụng mô hình PhoBERT để tận dụng khía cạnh ngữ cảnh của các từ do PhoBERT được huấn luyện trên một tập dữ liệu lớn dành riêng cho tiếng Việt. Bằng cách sử dụng PhoBERT, mô hình có thể biểu diễn tốt mối quan hệ của các từ trong một câu, từ đó nâng cao chất lượng quá trình học các mẫu tiềm ẩn trong dữ liệu. Qua đó, giúp cải tiến quá trình phân lớp. Kết quả thực nghiệm sử dụng ba mô hình trên cùng một bộ dữ liệu tiếng Việt cho thấy mô hình PhoBERT cho kết quả khả quan so với hai mô hình còn lại.

TÀI LIỆU THAM KHẢO

- [1] Thanh, Ho Quang, Ho Quang Thanh and ninh-pm-se, “thanhhocse96/vfnd-vietnamese-fake-news-bộ dữ liệu: Tập hợp các bài báo tiếng Việt và các bài post Facebook phân loại 2 nhãn Thật & Giả (228 bài)”. Zenodo, 27-Feb-2019, 2019.
- [2] Ahmed H, Traore I, Saad S, Ahmed H, Traore I, Saad S (2017) Detection of online fake news using N-gram analysis and machine learning techniques. In: International conference on intelligent, secure, and dependable systems in distributed and cloud environments. Springer, Cham, pp 127, 2017.
- [3] Ghanem B, Rosso P, Rangel F, Ghanem B, Rosso P, Rangel F (2018) Stance detection in fake news a combined feature representation. In: Proceedings of the first workshop on fact extraction and VERification (FEVER), pp 66–71, 2018.
- [4] Ruchansky N, Seo S, Liu Y, Ruchansky N, Seo S, Liu Y (2017) Csi: A hybrid deep model for fake news detection. In: Proceedings of the 2017 ACM on conference on information and knowledge management. ACM, pp 797–806, 2017.
- [5] Kaliyar RK, Goswami A, Narang P, Sinha S, Kaliyar RK, Goswami A, Narang P, Sinha S (2020) FNDNetA deep convolutional neural network for fake news detection. Cognitive Systems Research 61:32–44, 2020.
- [6] QĐND, QĐND, “qdnd.vn,” 22 11 2021. [Online]. Available: <https://www.qdnd.vn/phong-chong-dien-bien-hoa-binh/tin-gia-hiem-hoa-that-678175..>
- [7] M.-W. C. K. L. a. K. T. B. P.-t. o. Jacob Devlin, Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova, Bert: Pre-training of, 2019.
- [8] Nguyen, Dat Quoc and Anh Gia-Tuan Nguyen. “PhoBERT: Pre-trained, Nguyen, Dat Quoc and Anh Gia-Tuan Nguyen. “PhoBERT: Pre-trained.
- [9] Jwa H, Oh D, Park K, Kang JM, Lim H, Jwa H, Oh D, Park K, Kang JM, Lim H (2019) exBAKE: Automatic fake news detection model based on bidirectional encoder representations from transformers (BERT). Appl Sci 9(19):4062, 2019.
- [10] Karimi H, Roy P, Saba-Sadiya S, Tang J (2018) Multi-source multi-class fake news detection. In: Proceedings of the 27th international conference on computational linguistics, pp 1546–1557, 2018.