



(12) BẢN MÔ TẢ SÁNG CHẾ THUỘC BẰNG ĐỘC QUYỀN SÁNG CHẾ

(19) Cộng hòa xã hội chủ nghĩa Việt Nam (VN)
CỤC SỞ HỮU TRÍ TUỆ

(11)



1-0042068

(51)^{2020.01} G06F 40/00

(13) B

(21) 1-2020-01159

(22) 28/02/2020

(45) 25/12/2024 441

(43) 27/09/2021 402

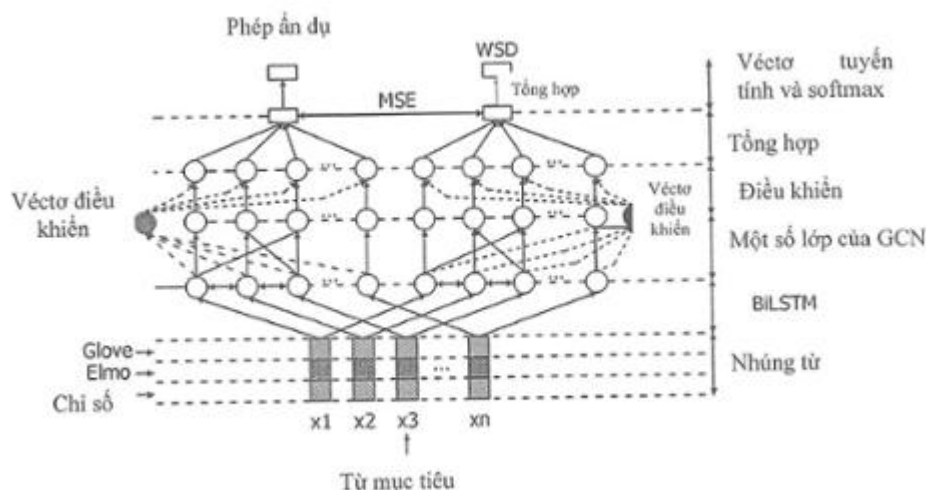
(73) Công ty Cổ phần Nghiên cứu và Ứng dụng Trí tuệ nhân tạo VinAI (VN)
Tòa nhà văn phòng Symphony, đường Chu Huy Mân, KĐT sinh thái Vinhomes
Riverside, phường Phúc Lợi, quận Long Biên, thành phố Hà Nội, Việt Nam

(72) Hung Hai Bui (US); Nguyễn Hữu Thiện (VN); Lê Minh Dương (VN).

(74) Công ty TNHH Tầm nhìn và Liên danh (VISION & ASSOCIATES CO.LTD.)

(54) BỘ MÃ HÓA, HỆ THỐNG VÀ PHƯƠNG PHÁP ĐỂ PHÁT HIỆN PHÉP ẨN DỤ TRONG XỬ LÝ NGÔN NGỮ TỰ NHIÊN

(57) Sáng chế đề cập đến bộ mã hóa, hệ thống và phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên. Hệ thống này bao gồm môđun mã hóa được tạo cấu hình để biến đổi các từ có trong câu thành các vectơ biểu diễn BiLSTM; bộ mã hóa thứ nhất được tạo cấu hình để tạo ra vectơ biểu diễn tổng thể thứ nhất của tác vụ giải quyết làm rõ nghĩa từ (Word Sense Disambiguation, WSD); bộ mã hóa thứ hai được tạo cấu hình để tạo ra vectơ biểu diễn tổng thể thứ hai của tác vụ phát hiện phép ẩn dụ (metaphor detection, MD); và môđun học đa nhiệm được tạo cấu hình để thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất và thứ hai. Trong đó, mỗi trong số các bộ mã hóa thứ nhất và thứ hai này bao gồm môđun mạng nơron tích chập dạng đồ thị (graph convolutional neural network, GCN) được tạo cấu hình để mã hóa liên kết giữa từ mục tiêu và từ cốt lõi để tạo ra các vectơ biểu diễn GCN; mô đun điều khiển được tạo cấu hình để điều chỉnh các vectơ biểu diễn GCN để tạo ra vectơ biểu diễn tổng thể.



Lĩnh vực kỹ thuật được đề cập

Sáng chế đề cập đến bộ mã hóa, hệ thống và phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên.

Tình trạng kỹ thuật của sáng chế

Phép diễn đạt ẩn dụ là phương tiện thuyết phục dùng cho giao tiếp hàng ngày và có thể tạo ra sự sinh động và rõ ràng cho tư duy và trao đổi thông tin của con người.

Ở mức độ nhận thức, phép diễn đạt ẩn dụ có thể hỗ trợ khái niệm hóa các trải nghiệm cụ thể của con người trong thế giới thực và chuyển kiến thức được khái niệm hóa sang các lĩnh vực khác. Cụ thể hơn, phép diễn đạt ẩn dụ có thể được giải thích là hiện tượng trong đó tồn tại liên kết ẩn dụ có hệ thống giữa hai khái niệm và/hoặc lĩnh vực khác nhau được thiết lập trong nhận thức của con người và được biểu diễn bằng ngôn ngữ.

Phát hiện phép diễn đạt ẩn dụ đã trở thành vấn đề quan trọng trong lĩnh vực xử lý ngôn ngữ tự nhiên (Natural Language Processing, NLP) do tính phổ biến của phép ẩn dụ. Nhận diện chính xác phép diễn đạt ẩn dụ có tác động đáng kể đến khả năng hiểu văn bản của các hệ thống NLP dùng cho các mục đích khác nhau (ví dụ, trích xuất thông tin, thảo luận ý kiến và dịch máy). Tuy nhiên, phát hiện phép diễn đạt ẩn dụ là một vấn đề khó trong đó các hệ thống cần hiểu nghĩa đen của phép diễn đạt ẩn dụ ở văn bản cụ thể và phân biệt giữa nghĩa đen với nghĩa bóng trên cơ sở so sánh tương tự.

Các hệ thống học dựa trên quy tắc và các hệ thống học máy đang được sử dụng làm các tác vụ ban đầu nhằm phát hiện văn bản để phát hiện phép diễn đạt ẩn dụ. Hơn nữa, gần đây, các kỹ thuật học sâu và các kỹ thuật diễn đạt từ đại diện được áp dụng để phát hiện phép diễn đạt ẩn dụ.

Bản chất kỹ thuật của sáng chế

Sáng chế nhằm đề xuất hệ thống và phương pháp để nhận dạng các phép diễn đạt và từ ẩn dụ trong văn bản.

Sáng chế còn nhằm đề xuất hệ thống và phương pháp để phát hiện phép diễn đạt ẩn dụ có độ chính xác cao bằng cách sử dụng công nghệ học sâu.

Theo khía cạnh thứ nhất, sáng chế đề xuất bộ mã hóa để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên (natural language processing, NLP) bao gồm: môđun mạng nơron tích chập dạng đồ thị (Graph Convolutional Neural Network, GCN) được tạo cấu hình để mã hóa liên kết giữa từ mục tiêu và từ cốt lõi để tạo ra các vectơ biểu diễn GCN; môđun điều khiển được tạo cấu hình để điều chỉnh các vectơ biểu diễn GCN để tạo ra vectơ biểu diễn tổng thể.

Môđun GCN, để mã hóa liên kết giữa từ mục tiêu và từ cốt lõi, cấp các vectơ biểu diễn BiLSTM vào mạng nơron tích chập dạng đồ thị (GCN) được tạo cấu trúc để tính toán trên cây phụ thuộc của câu đầu vào.

Môđun điều khiển, để điều chỉnh các vectơ biểu diễn GCN, áp dụng hàm kích hoạt (ví dụ, đơn vị tuyến tính được tinh chỉnh (rectified linear unit, ReLU)) cho vectơ biểu diễn BiLSTM của từ mục tiêu để lần lượt tạo ra vectơ điều khiển BiLSTM và vectơ điều khiển GCN; áp dụng vectơ điều khiển BiLSTM và vectơ điều khiển GCN cho các vectơ biểu diễn BiLSTM và các vectơ biểu diễn GCN qua phép nhân từng phần tử để lần lượt tạo ra các vectơ BiLSTM được lọc và các vectơ GCN được lọc; và tổng hợp các vectơ GCN được lọc để tạo ra vectơ biểu diễn tổng thể.

Bộ mã hóa còn bao gồm mạng nơron truyền thẳng sử dụng vectơ biểu diễn tổng thể làm đầu vào và tính toán sự phân bố xác suất của việc liệu từ mục tiêu có là phép diễn đạt ẩn dụ hoặc phép diễn đạt không ẩn dụ hay không bằng cách sử dụng lớp softmax, trong đó hàm mất mát để huấn luyện là hàm log hợp lý âm cho tập dữ liệu học.

Theo khía cạnh thứ hai, sáng chế còn đề xuất hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên bao gồm môđun mã hóa được tạo cấu hình để biến đổi các từ có trong câu thành các vectơ biểu diễn BiLSTM, bộ mã hóa thứ nhất được tạo cấu hình để tạo ra vectơ biểu diễn tổng thể thứ nhất của tác vụ giải quyết

WSD, bộ mã hóa thứ hai được tạo cấu hình để tạo ra vectơ biểu diễn tổng thể thứ hai của tác vụ MD, trong đó các bộ mã hóa thứ nhất và thứ hai này có cùng kiến trúc mạng với bộ mã hóa theo khía cạnh thứ nhất của sáng chế, mô đun học đa nhiệm được tạo cấu hình để thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất và thứ hai.

Mô đun mã hóa, để biến đổi các từ có trong câu thành các vectơ biểu diễn BiLSTM, ghép các vectơ được tính toán bằng cách sử dụng kỹ thuật nhúng từ không có ngữ cảnh đã được huấn luyện trước, nhúng từ có ngữ cảnh và nhúng chỉ số thành vectơ được ghép, và thực thi mạng BiLSTM trên vectơ được ghép nêu trên để tạo ra các vectơ biểu diễn BiLSTM.

Mô đun học đa nhiệm bao gồm bộ phân loại theo tác vụ cụ thể được tạo cấu hình để tính toán sự phân bố xác suất của các nhãn có thể đối với mỗi tác vụ.

Mô đun học đa nhiệm cực tiểu hóa hàm mất mát sau:

$$C(w^t, p^t, y^t) = -\log P^t(y^t | w^t, p^t) + \lambda \|V^{wsd} - V^{md}\|_2^2$$

để cập nhật các tham số cho bộ phân loại theo tác vụ cụ thể và các bộ mã hóa thứ nhất và thứ hai.

Quy trình huấn luyện luân phiên được sử dụng để huấn luyện hệ thống.

Mô tả vắn tắt các hình vẽ

Các đối tượng, dấu hiệu và ưu điểm nêu trên và đối tượng, dấu hiệu và ưu điểm khác của sáng chế sẽ trở nên rõ ràng đối với người có hiểu biết trung bình trong lĩnh vực kỹ thuật bằng cách mô tả các phương án làm ví dụ của sáng chế một cách chi tiết dựa trên các hình vẽ kèm theo, trong đó:

FIG.1 là sơ đồ ý tưởng của hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án.

FIG.2 là sơ đồ khối của cấu hình của hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án.

FIG.3 là sơ đồ khối của cấu hình của bộ mã hóa để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án.

FIG.4 là lưu đồ của phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án.

FIG.5 là lưu đồ của phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án.

FIG.6 là câu ví dụ để giải thích một phương án.

Mô tả chi tiết sáng chế

Sau đây, các phương án làm ví dụ của sáng chế sẽ được mô tả một cách chi tiết. Tuy nhiên, nếu xác định được rằng phần mô tả chi tiết về các chức năng hoặc cấu hình đã biết có liên quan đến sáng chế làm tối nghĩa đối tượng của sáng chế trong phần mô tả về các phương án một cách không cần thiết, thì phần mô tả chi tiết này sẽ được lược bỏ. Ngoài ra, các kích thước của các phần tử trên các hình vẽ có thể được phóng đại để giải thích và không đề cập đến các kích thước thực tế được áp dụng.

Theo các phương án này, khi một câu được cho, thì có thể xác định được liệu mỗi từ có trong câu có phải là từ mục tiêu để phát hiện phép ẩn dụ (metaphor detection, MD) hay không bằng cách sử dụng bài toán phân loại nhị phân.

Theo các phương án này, thuật ngữ " $w=w_1, w_2, \dots, w_n$ " có thể biểu thị câu có chiều dài n , và w_a (ở đây, $a=1, 2, 3, \dots, n$) có thể biểu thị từ mục tiêu để MD.

Hơn nữa, theo các phương án này, từ mục tiêu để MD có thể được gọi là từ mục tiêu hoặc từ được quan tâm.

Ngoài ra, các từ ngữ cảnh quan trọng, mà quan trọng hơn so với các từ ngữ cảnh khác về mặt đóng góp của chúng đối với việc xác định phép ẩn dụ của từ mục tiêu, có thể được gọi là các từ cốt lõi hoặc từ được liên kết.

FIG.1 là sơ đồ ý tưởng của hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án, FIG.2 là sơ đồ khối của cấu hình của hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án và FIG.3 là sơ đồ khối của cấu hình của bộ mã hóa để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án.

Dựa vào FIG.3, theo một phương án, bộ mã hóa 9 để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên (NLP) có thể bao gồm môđun mạng nơron tích chập dạng đồ thị (graph convolutional neural network, GCN) 12 được tạo cấu hình để mã hóa liên kết giữa từ mục tiêu và từ cốt lõi để tạo ra các vectơ biểu diễn GCN; môđun điều khiển 13 được tạo cấu hình để điều chỉnh các vectơ biểu diễn GCN để tạo ra vectơ biểu diễn tổng thể.

Dựa vào các hình vẽ FIG.1 và FIG.2, theo một phương án, hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên (NLP) 10 có thể bao gồm môđun mã hóa 11 được tạo cấu hình để biến đổi các từ có trong câu thành các vectơ biểu diễn BiLSTM, bộ mã hóa thứ nhất 15 được tạo cấu hình để tạo ra vectơ biểu diễn tổng thể thứ nhất của tác vụ giải quyết làm rõ nghĩa từ (Word Sense Disambiguation, WSD), bộ mã hóa thứ hai 16 được tạo cấu hình để tạo ra vectơ biểu diễn tổng thể thứ hai của tác vụ phát hiện phép ẩn dụ (Metaphor Detection, MD), trong đó các bộ mã hóa thứ nhất 15 và thứ hai 16 có thể có cùng kiến trúc mạng với bộ mã hóa 9 trên FIG.3, môđun học đa nhiệm 14 được tạo cấu hình để thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất 15 và thứ hai 16.

Môđun mã hóa 11 có thể biến đổi các từ có trong câu w thành vectơ giá trị thực dùng cho kiến trúc học sâu. Theo một phương án, môđun mã hóa 11 này có thể biến đổi các từ có trong câu thành các vectơ biểu diễn BiLSTM.

Môđun mã hóa 11 có thể biến đổi mỗi trong số các từ có trong câu thành vectơ bằng cách sử dụng kỹ thuật nhúng từ không có ngữ cảnh đã được huấn luyện trước, nhúng từ có ngữ cảnh và nhúng chỉ số.

Ví dụ, môđun mã hóa 11 có thể biến đổi mỗi từ w_i ($w_i \in w$) thành vectơ x_i bằng cách ghép các vectơ được tính toán bằng cách sử dụng ba kỹ thuật nhúng sau.

Vectơ thứ nhất có thể đề cập đến vectơ được tạo ra bằng cách sử dụng kỹ thuật nhúng từ không có ngữ cảnh đã được huấn luyện trước của từ w_i từ mô hình Global Vector (Global Vectors, GloVe) để biểu diễn từ. Môđun mã hóa 11 có thể tìm kiếm vectơ thứ nhất trong bảng nhúng được tạo ra bởi GloVe.

Vectơ thứ hai có thể đề cập đến các vectơ được tạo ra bằng cách sử dụng kỹ thuật nhúng từ có ngữ cảnh của từ w_i từ mô hình kỹ thuật nhúng từ mô hình ngôn

ngữ (Embeddings from Language Model, ELMo). Môđun mã hóa 11 có thể thực thi mô hình ELMo được huấn luyện trước đối với câu đầu vào để tạo ra ba chuỗi vectơ ẩn đối với các từ trong câu. Trong trường hợp này, mỗi trong số các chuỗi này có thể tương ứng với lớp kiến trúc ELMo. Sau đó, môđun mã hóa 11 có thể tính toán tổng có trọng số của các vectơ ẩn ở vị trí i của mỗi chuỗi để tạo ra vectơ biểu diễn tích lũy cho từ w_i . Trọng số của các chuỗi tương ứng và các tham số vô hướng của ELMo có thể được học trong quy trình huấn luyện của toàn bộ mô hình.

Vectơ thứ ba có thể đề cập đến vectơ được tạo ra bởi kỹ thuật nhúng chỉ số. Trong vectơ được tạo ra bởi kỹ thuật nhúng chỉ số này, để chỉ rõ rằng từ w_a là từ quan tâm trong MD, bộ chỉ báo nhị phân b_i có thể được gán. Khi $i=a$, thì bộ chỉ báo nhị phân b_i có thể được gán bằng một, nếu không, bộ chỉ báo nhị phân b_i có thể được gán bằng không. Sau đó, bộ chỉ báo nhị phân b_i có thể được ánh xạ đến vectơ giá trị thực bằng cách sử dụng bảng nhúng chỉ số được khởi tạo một cách ngẫu nhiên và được cố định trong quá trình huấn luyện.

Môđun mã hóa 11 có thể biến đổi câu w thành chuỗi các vectơ giá trị thực $x=x_1, x_2, \dots, x_n$ bằng cách sử dụng ba vectơ được mô tả ở trên.

Theo một phương án, do đặc điểm của mạng bộ nhớ dài - ngắn hạn hai chiều (Bidirectional Long-Short Term Memory, BiLSTM), nên các thành phần của vectơ thứ hai có thể hữu ích để đóng gói thông tin ngữ cảnh của toàn bộ câu vào mỗi vectơ x_i . Tuy nhiên, các thành phần của vectơ thứ nhất độc lập với các thành phần của vectơ thứ ba, và các thành phần của vectơ thứ nhất và các thành phần của vectơ thứ ba không được tích hợp tốt với các thành phần của vectơ thứ hai để tạo ra các vectơ biểu diễn có ngữ cảnh hoàn toàn và giàu thông tin đối với các từ trong câu đầu vào.

Do đó, để kết hợp ba vectơ thành phần của vectơ x_i một cách hiệu quả, thì môđun mã hóa 11 có thể thực thi mạng BiLSTM trên chuỗi vectơ $x_1, x_2, \dots, x_n$ và tạo ra chuỗi vectơ ẩn $h=h_1, h_2, \dots, h_n$. Chuỗi vectơ ẩn này có thể được sử dụng làm vectơ biểu diễn.

Do các đặc điểm lặp lại và hai chiều của mạng BiLSTM, nên thông tin ngữ cảnh trong toàn bộ câu x ở mỗi vectơ h_i có thể được mã hóa và các thành phần còn

lại của véc tơ x_i có thể được đóng gói một cách thuận lợi thành một không gian diễn đạt được ngữ cảnh hóa duy nhất.

Môđun GCN 12 có thể mã hóa liên kết giữa từ mục tiêu và từ cốt lõi để tạo ra các véc tơ biểu diễn GCN.

Môđun GCN 12 có thể sử dụng GCN với cây phân tích cú pháp phụ thuộc của câu để liên kết trực tiếp từ mục tiêu với từ cốt lõi về mặt ngữ cảnh để MD. Ở đây, GCN có thể đề xuất cơ chế điều khiển mới để lọc véc tơ biểu diễn được huấn luyện để duy trì thông tin quan trọng nhất để MD.

Theo một phương án, cây phân tích cú pháp phụ thuộc có thể đề cập đến dữ liệu thu được bằng cách lập đồ thị mối quan hệ liên kết giữa các từ có trong câu.

Môđun GCN 12 có thể thực hiện mô hình hóa đối với các từ cốt lõi trong văn bản. Khi từ quan tâm, là mục tiêu để phát hiện phép diễn đạt ẩn dụ, được cho trong văn bản, thì một số từ có thể là các từ cốt lõi để MD của từ quan tâm, trong khi các từ còn lại có thể là các từ không quan trọng để MD.

Ví dụ, trong câu ví dụ trên FIG.6, từ “nhà” có thể là từ cốt lõi để phát hiện phép ẩn dụ của từ “nằm” là từ quan tâm. Điều này là vì từ “nằm” thường được sử dụng cho động vật sống thay vì đối tượng tĩnh chẳng hạn như nhà.

Mặt khác, các từ còn lại (ví dụ, “đường như,” “nổi,” “sương mù,” và các từ tương tự) có thể là các từ không quan trọng để MD và có thể được nhận dạng là nhiễu trong phép diễn đạt được học bởi mô hình học sâu. Theo sáng chế này, từ cốt lõi có thể được phát hiện bằng cách thực hiện tính toán GCN để MD trên cơ sở cây phân tích cú pháp phụ thuộc của câu.

Theo một phương án, liên kết giữa từ đầu và từ bổ nghĩa có thể liên kết từ cốt lõi và từ quan tâm, là từ mục tiêu để MD. Ví dụ, trong câu ví dụ trên FIG.6, từ cốt lõi “nhà” có thể được đặt cạnh từ “nằm” là từ quan tâm về mặt cú pháp bằng mối quan hệ “ n_{subj} ”. Trong GCN được thực hiện đối với cấu trúc phụ thuộc, véc tơ biểu diễn cho từ ở một giai đoạn có thể được tính toán dựa vào véc tơ biểu diễn của từ liền kề về mặt cú pháp ở giai đoạn trước đó. Kết quả là, véc tơ biểu diễn cho từ của GCN có thể bao gồm thông tin của từ cốt lõi để MD trong câu. Theo đó, các từ không được liên kết, ngoại trừ từ cốt lõi và từ quan tâm, có thể được lọc và hiệu

suất của mô hình có thể được cải thiện. Hơn nữa, bằng cách tạo ra hàm tùy chỉnh cho vectơ biểu diễn GCN để phân loại một từ hoặc một phép diễn đạt duy nhất trong câu tại một thời điểm để phát hiện phép ẩn dụ, nên chỉ thông tin được liên kết với từ quan tâm có thể được duy trì và được truyền qua bước tính toán tiếp theo. Cụ thể là, vectơ điều khiển được tạo ra từ vectơ biểu diễn của từ quan tâm có thể được sử dụng làm bộ lọc vectơ GCN để thực hiện MD.

Khi môđun GCN 12 nhận các vectơ ẩn từ môđun mã hóa 11, thì môđun GCN 12 có thể sử dụng các vectơ ẩn này để tính toán các vectơ biểu diễn GCN để MD. Trong trường hợp này, môđun GCN 12 có thể nhận dạng từ cốt lõi được liên kết về mặt ngữ cảnh với từ mục tiêu.

Như được mô tả trên, quy trình nhận dạng từ cốt lõi hiện tại không được thực hiện đối với mạng BiLSTM. Do đó, môđun GCN 12 có thể cấp các vectơ ẩn (nghĩa là, các vectơ biểu diễn BiLSTM) h_i vào GCN được tạo cấu trúc để tính toán dựa vào cây phân tích cú pháp phụ thuộc của câu đầu vào w .

Phép tính của GCN yêu cầu ma trận liên kết A để mã hóa liên kết của các từ trong cây phân tích cú pháp phụ thuộc cho vectơ x . Môđun GCN 12 có thể thêm cạnh nghịch đảo và vòng tự lặp vào cây ngoài liên kết trực tiếp vốn có của cây phân tích cú pháp phụ thuộc để tạo ra ma trận liên kết. Cạnh nghịch đảo đã thêm có thể cho phép từ cốt lõi của từ w_i và chính từ w_i ở cây phân tích cú pháp phụ thuộc góp phần vào việc tính toán vectơ biểu diễn GCN cho từ w_i qua phép toán tích chập. Cạnh nghịch đảo đã thêm giúp làm giàu vectơ biểu diễn GCN bằng ngữ cảnh quan trọng về mặt cú pháp để cải thiện hiệu suất của MD.

Môđun GCN 12 có thể bao gồm tích chập đa lớp. Mỗi lớp có thể sử dụng ma trận H_i ($i \geq 0$) của lớp trước đó i làm đầu vào để tính toán ma trận H_{i+1} ở lớp hiện tại dựa vào phương trình 1 bên dưới.

[Phương trình 1]

$$H_{i+1} = g(AH_i W_i^g)$$

Trong phương trình 1, $H_0 = [h_1, h_2, \dots, h_n]$ có thể biểu thị ma trận có các hàng là các vectơ ẩn từ BiLSTM được tạo ra bởi môđun mã hóa 11. Hơn nữa, W_i^g có thể biểu thị ma trận trọng số của lớp thứ i và g có thể biểu thị hàm không tuyến tính.

Môđun GCN 12 có thể tối ưu số lượng lớp của môđun GCN 12 trên cơ sở tập dữ liệu hợp lệ bằng cách sử dụng phương trình 1.

Sau đây, theo một phương án, các vectơ hàng của ma trận ở lớp tích chập cuối cùng của môđun GCN 12 sẽ được gọi là $h_g = h_1^g, h_2^g, \dots, h_n^g$.

Để thực hiện tác vụ MD và tác vụ giải quyết làm rõ nghĩa từ (Word-Sense Disambiguation, WSD) đối với từ mục tiêu, môđun điều khiển 13 có thể tính toán mỗi vectơ biểu diễn tổng thể và môđun học đa nhiệm 14 có thể cấp vectơ biểu diễn tổng thể cho bộ phân loại theo tác vụ cụ thể để tính toán sự phân bố xác suất của các nhãn có thể đối với mỗi tác vụ và thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất và thứ hai sử dụng tính tương tự giữa tác vụ MD và tác vụ giải quyết WSD.

Hơn nữa, môđun điều khiển 13 có thể điều chỉnh các vectơ biểu diễn GCN để tạo ra vectơ biểu diễn tổng thể.

Hơn nữa, môđun điều khiển 13 có thể tính toán các vectơ điều khiển bằng cách áp dụng hàm kích hoạt (ví dụ, đơn vị tuyến tính được tinh chỉnh (ReLU)) cho vectơ biểu diễn BiLSTM của từ mục tiêu để lần lượt tạo ra vectơ điều khiển BiLSTM và vectơ điều khiển GCN, áp dụng vectơ điều khiển BiLSTM và vectơ điều khiển GCN cho các vectơ biểu diễn BiLSTM và các vectơ biểu diễn GCN qua phép nhân từng phần tử để lần lượt tạo ra các vectơ BiLSTM được lọc và các vectơ GCN được lọc, và tổng hợp các vectơ GCN được lọc để tạo ra vectơ biểu diễn tổng thể.

Để môđun điều khiển 13 thực hiện MD đối với từ mục tiêu w_a , thì mô hình học sâu có thể tính toán vectơ biểu diễn tổng thể V .

Theo một phương án, môđun điều khiển 13 có thể tổng hợp các vectơ biểu diễn GCN $h_g = h_1^g, h_2^g, \dots, h_n^g$ bằng cách tổng hợp (ví dụ, lớn nhất hoặc trung bình) và các cơ chế chú ý để tạo ra vectơ biểu diễn tổng thể.

Trong trường hợp này, có nhược điểm là tất cả các chiều của vectơ biểu diễn được giả sử có cùng mức độ quan trọng. Vì các chiều của vectơ biểu diễn hoàn toàn chưa bị ràng buộc, nên một số chiều có ảnh hưởng nhiều hơn các chiều còn lại để MD. Ngoài ra, vì các vectơ biểu diễn BiLSTM và GCN được trích xuất từ vectơ gốc x_i thông qua BiLSTM và GCN, là các lớp ẩn, nên thông tin (cụ thể là, thành

phần nhúng chỉ số của x_i) về vị trí của từ mục tiêu có thể được thể hiện không rõ ràng, dẫn đến sự nhầm lẫn của vectơ biểu diễn đối với từ mục tiêu để MD.

Theo một phương án, môđun điều khiển 13 có thể giải quyết vấn đề được mô tả ở trên bằng cách điều chỉnh vectơ biểu diễn GCN cụ thể hơn và nhận biết được từ mục tiêu w_a .

Vì sự điều chỉnh được thực hiện bởi môđun điều khiển 13 ở cấp độ chiều (nghĩa là, theo từng dấu hiệu), nên các chiều này có thể được định lượng một cách phù hợp tùy thuộc vào mức độ quan trọng của chúng đối với bài toán MD.

Trước tiên, môđun điều khiển 13 có thể tính toán vectơ điều khiển trên cơ sở vectơ biểu diễn h_a của từ mục tiêu w_a và sau đó áp dụng vectơ điều khiển, làm bộ lọc theo từng dấu hiệu đối với các vectơ biểu diễn BiLSTM và GCN khác.

Môđun điều khiển 13 có thể tính toán vectơ điều khiển BiLSTM c_h đối với vectơ biểu diễn $h=h_1, h_2, \dots, h_n$ bằng cách sử dụng phương trình 2 bên dưới.

[Phương trình 2]

$$c_h = Relu(W_h h_a)$$

Môđun điều khiển 13 có thể lọc từ ít liên quan đến từ mục tiêu w_a trong vectơ biểu diễn BiLSTM h_i bằng phép nhân từng phần tử bằng cách sử dụng phương trình 3 bên dưới.

[Phương trình 3]

$$\hat{h}_i = c_h \odot h_i \quad \forall 1 \leq i \leq n$$

Trong phương trình 3, $\hat{h}_i$ có thể biểu thị vectơ BiLSTM được lọc, và phép nhân từng phần tử có thể đề cập đến cơ chế để đạt được việc điều chỉnh từng dấu hiệu và/hoặc mức chiều của vectơ biểu diễn.

Vectơ điều khiển GCN c_g có thể được điều chỉnh bởi tổng có trọng số m của vectơ biểu diễn h chẳng hạn.

Vì thông tin đã được biểu diễn trong vectơ biểu diễn BiLSTM h có thể được chuyển thành vectơ điều khiển GCN c_g , nên thông tin quan trọng của vectơ biểu

diễn BiLSTM có thể giữ nguyên khi véctor điều khiển GCN c_g được áp dụng cho véctor biểu diễn GCN h_g .

Môđun điều khiển 13 có thể sử dụng các véctor được lọc $\hat{h}_i$ để thu được các trọng số α cho các véctor biểu diễn h sao cho các trọng số này được tùy chỉnh của từ mục tiêu w_a bằng cách sử dụng phương trình 4 bên dưới.

[Phương trình 4]

$$\alpha_i = \frac{\exp(W_\alpha \hat{h}_i)}{\sum_{j=1}^n \exp(W_\alpha \hat{h}_j)}$$

$$m = \sum_{i=1}^n \alpha_i h_i, c_g = \text{Relu}(W_g[h_a, m])$$

Tiếp theo, môđun điều khiển 13 có thể áp dụng véctor điều khiển GCN c_g cho véctor biểu diễn GCN h_i^g bằng phép nhân từng phần tử sử dụng phương trình 5 bên dưới để tạo ra véctor biểu diễn được lọc GCN $\hat{h}_i^g$.

[Phương trình 5]

$$\hat{h}_i^g = c_g \odot h_i^g$$

Tiếp theo, môđun điều khiển 13 có thể tổng hợp véctor biểu diễn được lọc GCN $\hat{h}_i^g$ sử dụng véctor ghép bằng phương trình 6 bên dưới để tạo ra véctor biểu diễn tổng thể V để MD.

[Phương trình 6]

$$V = [\hat{h}_a^g, \max(\hat{h}_1^g, \hat{h}_2^g, \dots, \hat{h}_n^g)]$$

Trong phương trình 6, $\hat{h}_a^g$ có thể được sử dụng để thu thập thông tin ngữ cảnh về từ mục tiêu w_a , trong khi $\max(\hat{h}_1^g, \hat{h}_2^g, \dots, \hat{h}_n^g)$ có thể được sử dụng để nâng cao véctor biểu diễn tổng thể V bằng cách sử dụng thông tin ngữ cảnh quan trọng nhất từ các từ còn lại.

Bộ mã hóa 9 có thể sử dụng véctor biểu diễn tổng thể V làm đầu vào của mạng nơron truyền thẳng (không được thể hiện trên hình vẽ) và, cuối cùng, tính toán sự phân bố xác suất của việc liệu từ mục tiêu có là phép diễn đạt ẩn dụ hoặc phép diễn đạt không ẩn dụ hay không bằng cách sử dụng lớp softmax.

Theo một phương án, bộ mã hóa 9 có thể sử dụng hàm log-hợp lý âm cho tập dữ liệu học làm hàm mất mát để huấn luyện mô hình.

Môđun học đa nhiệm 14 có thể bao gồm bộ phân loại theo tác vụ cụ thể được tạo cấu hình để tính toán sự phân bố xác suất của các nhãn có thể đối với mỗi tác vụ.

Môđun học đa nhiệm 14 có thể sử dụng chương trình khung học đa nhiệm trong đó kiến thức được chuyển giữa hai tác vụ bằng cách sử dụng tính tương tự giữa giải quyết WSD và phát hiện phép diễn đạt ẩn dụ.

Nhận xét đầu tiên về kỹ thuật học đa nhiệm sử dụng giải quyết WSD là giải quyết WSD có liên quan đến tác vụ MD và hiệu suất phát hiện phép diễn đạt ẩn dụ có thể được cải thiện bằng cách chuyển kiến thức trong quy trình suy luận của giải quyết WSD. Theo một phương án, môđun học đa nhiệm 14 có thể nhận dạng nghĩa chính xác của từ và/hoặc phép diễn đạt về mặt ngữ cảnh trong số các nghĩa khả thi mà từ này thường có thể có thông qua giải quyết WSD.

Ở mức độ mô hình hóa, giải quyết WSD và MD có thể thực hiện các bài toán phân loại đối với các từ và/hoặc phép diễn đạt về mặt ngữ cảnh của câu. Ở mức ngữ nghĩa, các nghĩa ẩn dụ khác nhau của từ được ghi trong kho của WordNet. Ví dụ, nghĩa của từ “chìm” ở cụm từ “chìm trong công việc” là phép ẩn dụ của từ “chìm” ở cụm từ “bị chìm trong nước.”

Kết quả là, khi hệ thống học sâu có thể học sự biểu diễn hiệu quả để phân biệt giữa các nghĩa của các từ về mặt ngữ cảnh, thì các biểu diễn được suy ra từ hệ thống có thể còn hỗ trợ để phát hiện biểu diễn ẩn dụ do sự liên quan chặt chẽ trong mô hình hóa ngữ nghĩa về mặt ngữ cảnh.

Theo phương án này, bằng cách đề xuất chương trình khung học đa nhiệm mới khớp với các biểu diễn được suy ra để giải quyết WSD và MD để hỗ trợ chuyển kiến thức giữa hai tác vụ dựa vào tính tương tự giữa giải quyết WSD và MD.

Chương trình khung theo một phương án huấn luyện hai mạng đối với hai tác vụ sao cho các biểu diễn của hai mạng được thực hiện là tương tự khi cùng câu và/hoặc ngữ cảnh được biểu diễn, và do đó vấn đề trong đó tập dữ liệu này chỉ được gán cho một tác vụ duy nhất có thể được giải quyết. Nghĩa là, kiến thức ở tập dữ

liệu đối với giải quyết WSD có thể được chuyển bởi môđun học đa nhiệm 14 và do đó hiệu suất của tác vụ MD có thể được cải thiện.

Môđun học đa nhiệm 14 có thể đồng thời giải quyết một số tác vụ khác nhau nhưng liên quan đến nhau bằng cách sử dụng tập dữ liệu học sẵn có trong các thiết lập học đa nhiệm của xử lý ngôn ngữ tự nhiên (NLP).

Môđun học đa nhiệm 14 có thể thực thi thủ tục huấn luyện luân phiên để thay thế quy trình huấn luyện được sử dụng cho tác vụ được liên kết khi tập dữ liệu cho tác vụ được liên kết bao gồm các văn bản đầu vào khác nhau (ví dụ, khi câu của tập dữ liệu cho tác vụ chỉ có nhãn cho tác vụ tương ứng).

Một tác vụ có thể được chọn với xác suất nhất định bởi một thủ tục tính toán và môđun học đa nhiệm 14 có thể lấy mẫu nhóm nhỏ của tập dữ liệu đối với tác vụ tương ứng để tính toán mô hình mất mát và cập nhật.

Các tham số của các bộ mã hóa thứ nhất 15 và thứ hai 16 được chia sẻ bởi tất cả các tác vụ và cập nhật mỗi khi các tham số này được tính toán theo cách lặp lại, trong khi chỉ các tham số cho bộ phân loại 17 của tác vụ được chọn hiện tại bị tác động ở một lần học.

Chương trình khung học đa nhiệm để giải quyết WSD và MD theo một phương án có thể tương ứng với kịch bản sau vì các tập dữ liệu sẵn có cho hai tác vụ này không chia sẻ câu đầu vào.

Do đó, môđun học đa nhiệm 14 có thể sử dụng thủ tục huấn luyện luân phiên làm ngưỡng cơ sở của chương trình khung học đa nhiệm trong tác vụ tương ứng.

Vấn đề trong giải pháp ngưỡng cơ sở là một mô hình học sâu duy nhất có thể được sử dụng làm bộ mã hóa cho một số tác vụ xem xét được liên kết.

Mặc dù có mức độ tương tự nhất định giữa giải quyết WSD và MD về mặt phân loại ngữ nghĩa của các từ và/hoặc phép diễn đạt về mặt ngữ cảnh, nhưng vectơ biểu diễn của bộ mã hóa đối với giải quyết WSD có thể yêu cầu nhiều thông tin kỹ lưỡng hoặc tinh vi hơn MD.

Điều này là do nhãn của giải quyết WSD thường cụ thể và tỉ mỉ hơn so với nhãn của MD. Cụ thể, một từ của giải quyết WSD có thể có mười hoặc nhiều hơn mười nghĩa khác nhau, trong khi từ của MD có thể chỉ được gán cho hai nhãn

(nghĩa là, ẩn dụ hoặc không ẩn dụ). Mặt khác, một số nghĩa ẩn dụ có thể không thể hiện trong WordNet đối với giải quyết WSD, do đó vectơ biểu diễn của MD đòi hỏi phải nắm bắt được thông tin ngữ nghĩa mà có khả năng khác với vectơ biểu diễn của giải quyết WSD. Kết quả là, khi một bộ mã hóa duy nhất được sử dụng để suy ra các vectơ biểu diễn của giải quyết WSD và MD, thì bộ mã hóa này rất khó xác định thông tin ngữ nghĩa nào cần được tập trung vào, và do đó chất lượng biểu diễn có thể bị hạ thấp.

Mô đun học đa nhiệm 14 theo một phương án có thể giải quyết các vấn đề được mô tả ở trên bằng cách sử dụng hai bộ mã hóa 15 và 16 riêng biệt để tính toán các vectơ biểu diễn của giải quyết WSD và MD, thay vì sử dụng một bộ mã hóa học sâu duy nhất.

Hai bộ mã hóa 15 và 16 sử dụng kiến trúc mạng được mô tả ở trên, và khi cùng câu đầu vào được hiển thị, thì hai bộ mã hóa 15 và 16 này có thể tạo ra các vectơ biểu diễn tương tự để kiến thức có thể được chuyển giữa các bộ mã hóa 15 và 16.

Hai mạng mã hóa độc lập có thể tạo ra độ linh hoạt để học các đặc điểm cụ thể đối với các tác vụ cụ thể. Hơn nữa, cơ chế chuyển kiến thức giữa các bộ mã hóa tạo điều kiện cho việc tích hợp thông tin vào giải quyết WSD để MD.

Theo một phương án, E^{wsd} có thể biểu thị bộ mã hóa thứ nhất 15 được áp dụng để giải quyết WSD và E^{md} có thể biểu thị bộ mã hóa thứ hai 16 được áp dụng để MD.

Hơn nữa, (w^t, p^t, y^t) có thể biểu thị một ví dụ về bộ dữ liệu của giải quyết WSD hoặc MD trong đó t ($t \in \{wsd, md\}$) có thể biểu thị bộ chỉ báo tác vụ. Trong bộ dữ liệu này, w^t có thể biểu thị câu đầu vào, p^t có thể biểu thị vị trí của từ mục tiêu và y^t có thể biểu thị nhãn của từ w^t đối với tác vụ t .

Mô đun điều khiển 13 có thể cấp văn bản đầu vào (w^t, p^t) đến cả bộ mã hóa E^{wsd} và E^{md} bằng cách sử dụng phương trình 7 bên dưới để thực hiện chuyển kiến thức và do đó có thể tạo ra vectơ biểu diễn tổng thể V^{wsd} của giải quyết WSD và vectơ biểu diễn tổng thể V^{md} của MD.

[Phương trình 7]

$$V^{wsd} = E^{wsd}(w^t, p^t), V^{md} = E^{md}(w^t, p^t)$$

Đối với tác vụ t , vectơ biểu diễn V^t có thể được truyền đến bộ phân loại theo tác vụ cụ thể 17 Ft (ví dụ, mạng nơron truyền thẳng được theo sau bởi lớp softmax tiếp theo) để tính toán phân bố xác suất $P^t(\cdot | w^t, p^t)$ đối với các nhãn có thể cho t .

Môđun học đa nhiệm 14 có thể cập nhật các tham số của bộ phân loại theo tác vụ cụ thể 17 Ft và các bộ mã hóa thứ nhất E^{wsd} và thứ hai E^{md} bằng cách cực tiểu hóa hàm mất mát tiếp theo sử dụng phương trình 8 bên dưới.

[Phương trình 8]

$$C(w^t, p^t, y^t) = -\log P^t(y^t | w^t, p^t) + \lambda \|V^{wsd} - V^{md}\|_2^2$$

Trong phương trình 8, λ có thể biểu thị tham số đánh đổi. Nguyên lý của số hạng thứ hai là V^{wsd} và V^{md} cần tương tự vì V^{wsd} và V^{md} biểu thị các vectơ biểu diễn đối với cùng câu đầu vào w^t của các bộ mã hóa thứ nhất E^{wsd} và thứ hai E^{md} .

Hai bộ mã hóa có thể truyền thông với nhau để kiến thức từ một tác vụ (ví dụ, giải quyết WSD) có thể được chuyển đến tác vụ khác (ví dụ, MD) để cải thiện chất lượng của vectơ biểu diễn.

Hơn nữa, quy trình huấn luyện luân phiên có thể còn được sử dụng để huấn luyện hệ thống 10.

FIG.4 là lưu đồ của phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án. Môđun mạng nơron tích chập dạng đồ thị (graph convolutional neural network, GCN) 12 có thể mã hóa liên kết giữa từ mục tiêu và từ cốt lõi để tạo ra các vectơ biểu diễn GCN bằng cách cấp các vectơ biểu diễn BiLSTM vào mạng nơron tích chập dạng đồ thị (GCN) được tạo cấu trúc để tính toán trên cây phụ thuộc của câu đầu vào (S401). Môđun điều khiển 13 có thể điều chỉnh các vectơ biểu diễn GCN để tạo ra vectơ biểu diễn tổng thể bằng cách áp dụng hàm kích hoạt (ví dụ, đơn vị tuyến tính được tinh chỉnh (ReLU)) cho vectơ biểu diễn BiLSTM của từ mục tiêu để lần lượt tạo ra vectơ điều khiển BiLSTM và vectơ điều khiển GCN; áp dụng vectơ điều khiển BiLSTM và vectơ điều khiển GCN cho các vectơ biểu diễn BiLSTM và các vectơ biểu diễn GCN qua phép nhân

từng phần tử để lần lượt tạo ra các vectơ BiLSTM được lọc và các vectơ GCN được lọc; và tổng hợp các vectơ GCN được lọc để tạo ra vectơ biểu diễn tổng thể (S402).

Phương pháp này có thể còn bao gồm các bước nhập vectơ biểu diễn tổng thể vào mạng nơron truyền thẳng và tính toán sự phân bố xác suất của việc liệu từ mục tiêu có là phép diễn đạt ẩn dụ hoặc phép diễn đạt không ẩn dụ hay không bằng cách sử dụng lớp softmax, trong đó hàm mất mát để huấn luyện mô hình là hàm log hợp lý âm cho tập dữ liệu học.

FIG.5 là lưu đồ của phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án. Môđun mã hóa 11 có thể biến đổi các từ có trong câu thành các vectơ biểu diễn BiLSTM bằng cách ghép các vectơ được tính toán bằng cách sử dụng kỹ thuật nhúng từ không có ngữ cảnh đã được huấn luyện trước, nhúng từ có ngữ cảnh và nhúng chỉ số thành vectơ được ghép, và thực thi mạng BiLSTM trên vectơ được ghép nêu trên để tạo ra các vectơ biểu diễn BiLSTM (S501). Bộ mã hóa thứ nhất 15 có thể tạo ra vectơ biểu diễn tổng thể thứ nhất của tác vụ giải quyết WSD (S502). Bộ mã hóa thứ hai 16 có thể tạo ra vectơ biểu diễn tổng thể thứ hai của tác vụ MD (S503). Các bộ mã hóa thứ nhất 15 và thứ hai 16 này có thể thực hiện phương pháp như được mô tả trên FIG.4. Môđun học đa nhiệm 14 có thể thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất 15 và thứ hai 16. Bộ phân loại theo tác vụ cụ thể 17 có thể tính toán sự phân bố xác suất của các nhãn có thể đối với mỗi tác vụ. Môđun học đa nhiệm 14 có thể cực tiểu hóa hàm mất mát sau:

$$C(w^t, p^t, y^t) = -\log P^t(y^t | w^t, p^t) + \lambda \|V^{wsd} - V^{md}\|_2^2$$

để cập nhật các tham số cho bộ phân loại theo tác vụ cụ thể 17 và các bộ mã hóa thứ nhất 15 và thứ hai 16 (S504).

Trong thử nghiệm theo bảng 1 bên dưới, sử dụng bộ dữ liệu được mô tả ở trên, hiệu suất của các mô hình MD khác nhau so sánh với hiệu suất của hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án.

Mô hình	VUA All POS				VUA VERB				MOH-X				TroFi			
	P	R	F1	Acc	P	R	F1	Acc	P	R	F1	Acc	P	R	F1	Acc
Ngưỡng cơ sở	-	-	-	-	67,9	40,7	50,9	76,4	39,1	26,7	31,3	43,6	72,4	55,7	62,9	71,4

Lexical																
SimNet	-	-	-	-	-	-	-	-	73,6	76,1	74,2	74,8	-	-	-	-
CNN+BiLSTM+	60,8	70,0	65,1	-	60,0	76,3	67,2	-	-	-	-	-	-	-	-	-
RNN_CLS	-	-	-	-	53,4	65,6	58,9	69,1	75,3	84,3	79,1	78,5	68,7	74,6	72,0	73,7
RNN_SEQ_ELMo+	71,6	73,6	72,6	93,1	68,2	71,3	69,7	81,4	79,1	73,5	75,6	77,2	70,7	71,6	71,1	74,6
RNN_SEQ_BERT+	71,5	71,9	71,7	92,9	66,7	71,5	69,0	80,7	75,1	81,8	78,2	78,1	70,3	67,1	68,7	73,4
RNN_HG+	71,8	76,3	74,0	93,6	69,3	72,3	70,8	82,1	79,7	79,8	79,8	79,7	67,4	77,8	72,2	74,9
RNN_MHCA+	73,0	75,7	74,3	93,8	66,3	75,2	70,5	81,8	77,5	83,1	80,0	79,8	68,6	76,8	72,4	75,2
MUL_GCN	74,8	75,5	75,1	93,8	72,5	70,9	71,7	83,2	79,7	80,5	79,6	79,9	73,1	73,6	73,2	76,4

<Bảng 1>

Dựa vào bảng 1, MUL_GCN đề cập đến hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên theo một phương án, và Lexical Baseline, SimNet, CNN+BiLSTM, RNN CLS, RNN SEQ ELMo, RNN SEQ BERT, RNN HG và RNN MHCA đề cập đến hệ thống để phát hiện phép ẩn dụ thông thường trong xử lý ngôn ngữ tự nhiên theo các ví dụ so sánh.

Lexical Baseline là hệ thống đơn giản dựa vào tần số ẩn dụ của các từ (Gao và cộng sự 2018). SimNet là các mạng nơron tương tự sử dụng kỹ thuật nhúng từ skip-gram chẳng hạn (Rei và cộng sự 2017). CNN+BiLSTM là mô hình tập hợp được với các mạng nơron tích chập (Convolutional Neural Network, CNN) và BiLSTM chẳng hạn (Wu và cộng sự 2018). RNN CLS là mô hình BiLSTM với sự chú ý chẳng hạn (Gao và cộng sự 2018) đối với cách thiết lập phân loại. RNN SEQ ELMo là mô hình BiLSTM chẳng hạn (Gao và cộng sự 2018) đối với cách thiết lập dự đoán liên tiếp. Và RNN SEQ BERT (được báo cáo ở (Mao, Lin và Guerin 2019)) tương tự với RNN SEQ ELMo ngoại trừ các kỹ thuật nhúng ELMo được thay thế bằng các kỹ thuật nhúng BERT (Devlin và cộng sự 2019). RNN HG là mô hình BiLSTM dựa trên nguyên tắc MIP (Mao, Lin, và Guerin 2019). RNN MHCA là mô hình BiLSTM với sự chú ý về mặt ngữ cảnh và nguyên tắc SPV (Mao, Lin và Guerin 2019).

Hơn nữa, dựa vào bảng 1, trong thử nghiệm hiện tại, để tương thích với công nghệ trước đó, hoạt động phát hiện phép diễn đạt ẩn dụ đã được thực hiện sử dụng ba kiểu bộ dữ liệu, văn bản ra ẩn dụ VU Amsterdam (VU Amsterdam, VUA),

MOH-X và bộ tìm kiếm chuyên nghĩa (Trope Finder, TroFi), mà được sử dụng để MD.

VUA thể hiện bộ dữ liệu đánh giá công khai lớn nhất để phát hiện phép ẩn dụ được sử dụng bởi tác vụ chia sẻ phép ẩn dụ NAACL-2018. Chú thích của bộ dữ liệu này dựa vào MIP mà mỗi từ trong các câu được gắn nhãn để nhận dạng phép ẩn dụ. Tiếp theo tình trạng kỹ thuật trước đó, hai phiên bản của bộ dữ liệu này cũng được xem xét, nghĩa là, VUA ALL POS trong đó các từ của tất cả các kiểu (ví dụ, danh từ, động từ, tính từ) được gắn nhãn, và VUA VERB chỉ tập trung vào các động từ để phát hiện phép ẩn dụ. Đối với MOH-X, các câu ngắn hơn và đơn giản hơn so với các câu trong các bộ dữ liệu khác do chúng được lấy mẫu từ WordNet. Chỉ một động từ duy nhất được gắn nhãn ở mỗi câu trong MOH-X. Cuối cùng, TroFi bao hàm các câu từ bản phát hành tạp trí phổ Wall 1987-89 1. Tương tự với MOHX, TroFi cũng chỉ được chú thích cho một động từ mục tiêu duy nhất. Theo các cách thiết lập ở tình trạng kỹ thuật trước đó, sự xác thực chéo 10 lần đối với MOH-X và TroFi được thực hiện và các bộ dữ liệu VUA được chia thành các bộ huấn luyện, thiết lập xác thực và thử nghiệm. Cùng sự phân chia dữ liệu được sử dụng cho tất cả ba bộ dữ liệu như tình trạng kỹ thuật trước đó để so sánh công bằng. Bộ dữ liệu Semcor được sử dụng cho bộ dữ liệu WSD theo sáng chế này. Bộ dữ liệu này bao gồm các câu có các từ đã được chú thích bằng tay có các id cảm giác WordNet. Do số lượng từ trong Semcor lớn hơn nhiều so với số lượng từ trong các bộ dữ liệu này để phát hiện phép ẩn dụ, chỉ một phần của Semcor được lấy mẫu để huấn luyện các mô hình ở sáng chế này.

Cụ thể, việc lấy mẫu được thực hiện sao cho số lượng ví dụ trong WSD và các bộ dữ liệu phát hiện phép ẩn dụ sẽ tương tự nhau. Đối với các bộ dữ liệu phát hiện phép ẩn dụ chỉ bao hàm các động từ làm các mục tiêu (nghĩa là, VUA VERB, MOH-X, TroFi), chỉ các động từ trong Semcor cũng được lấy mẫu đối với WSD phù hợp. Liên quan đến các kỹ thuật nhúng từ được huấn luyện trước, các vectơ Glove 300d và vectơ ELMo 1024d còn được sử dụng chẳng hạn (Gao và cộng sự 2018; Mao, Lin và Guerin 2019). Kích thước của các kỹ thuật nhúng chỉ số được thiết lập thành 50 trong cách thiết lập phân loại để so sánh công bằng. Các siêu tham số của mô hình được đề xuất đối với mỗi bộ dữ liệu được tinh chỉnh dẫn đến

các giá trị tham số như sau. Số lượng đơn vị ẩn đối với cả hai mạng BiLSTM và mạng GCN là 200 trong khi số lượng lớp GCN được thiết lập thành 2. Các mô hình được huấn luyện với nhóm nhỏ được bố trí lại có kích thước là 32, sử dụng trình tối ưu hóa Adam để cập nhật các tham số này. Tham số đánh đổi λ đối với việc học đa nhiệm ở phương trình 4 của tất cả VUA ALL POS, VUA VERB, MOH-X và TroFi được thiết lập thành 1.

Bảng 1 thể hiện hiệu suất trong đó F1 là phép đo quan trọng nhất đối với tác vụ này.

Có hai cách thiết lập/tiếp cận khác nhau để phát hiện phép ẩn dụ trong tài liệu, nghĩa là, cách thiết lập gắn nhãn liên tiếp và cách thiết lập phân loại. Trong cách thiết lập gắn nhãn liên tiếp, các mô hình được huấn luyện để dự báo trình tự các nhãn nhị phân để biểu thị tính ẩn dụ của các từ trong các câu (nghĩa là, các mô hình trong bảng 1) trong khi cách thiết lập phân loại xác định tính ẩn dụ của các từ trong các câu một cách độc lập là vấn đề phân loại từ (nghĩa là, cách phát hiện phép ẩn dụ được xây dựng mô hình theo sáng chế này). Một mặt, bảng 1 thể hiện rằng trong số mô hình có cách thiết lập phân loại, mô hình được đề xuất vượt trội hơn đáng kể so với mô hình mới nhất trước đó (nghĩa là, RNN CLS). Khoảng cách hiệu suất có nghĩa và có giá trị so với VUA VERB và TroFi. Mặt khác, so sánh các mô hình gắn nhãn liên tiếp trước đó với mô hình được đề xuất MUL GCN, có thể thấy rằng MUL GCN còn có điểm số F1 tốt hơn đáng kể so với các mô hình trước đó (ví dụ, hệ thống mới nhất hiện tại RNN MHCA) đối với ba trên bốn bộ dữ liệu được xem xét ($p < 0,01$). Loại trừ duy nhất là đối với bộ dữ liệu MOHX trong đó MUL GCN đạt được hiệu suất so sánh được với RNN MHCA. Các bằng chứng này rõ ràng giúp chứng minh các ưu điểm của mô hình được đề xuất so với mô hình trong tình trạng kỹ thuật trước đó. Một điểm thú vị là dự đoán các nhãn ẩn dụ của các từ ngữ cảnh (ở cách thiết lập gắn nhãn liên tiếp) được đề xuất bởi tình trạng kỹ thuật trước đó là cách tốt hơn để phát hiện phép ẩn dụ so với thiết lập phân loại. Tuy nhiên, theo sáng chế này, điều ngược lại là cách thiết lập phân loại có thể vẫn tạo ra các mô hình phát hiện phép ẩn dụ với hiệu suất mới nhất được thể hiện. Thành tựu này được chỉ định cho đề xuất học đa nhiệm, cơ chế điều khiển và các GCN giúp tăng hiệu suất của mô hình theo sáng chế này một cách đáng kể.

Các thành phần chính của mô hình mạng nơron theo sáng chế này bao gồm mô hình BiLSTM trong môđun mã hóa, môđun GCN và môđun điều khiển. Phần này đánh giá hiệu quả của các thành phần này khi chúng được loại bỏ khỏi toàn bộ mô hình. Đối với môđun GCN, mô hình khi môđun GCN được thay thế bằng lớp tự chú ý đa đầu phổ biến từ bộ biến đổi được đánh giá thêm để chứng minh sự cần thiết của các GCN theo sáng chế này. Tương tự với tình trạng kỹ thuật trước đó, tập dữ liệu VERB VUA được sử dụng cho các nghiên cứu lược bỏ này. Bảng 2 thể hiện hiệu suất của các mô hình trong các bộ dữ liệu thử nghiệm của VUA VERB.

Mô hình	P	R	F1	Acc
MUL_GCL	72,5	70,9	71,7	83,2
Lớp BiLSTM - MUL_GCL	65,9	69,3	67,6	80,0
Môđun điều khiển - MUL_GCL	69,0	67,6	68,3	81,2
Môđun GCN - MUL_GCL	74,6	60,9	67,0	82,0
Thay thế GCN bằng tự chú ý	72,6	67,4	69,9	82,6

<Bảng 2>

Như được thể hiện từ bảng 2, mỗi thành phần (nghĩa là, BiLSTM, các môđun điều khiển và GCN) là quan trọng đối với mô hình được đề xuất MUL GCN do loại trừ thành phần bất kỳ trong số chúng sẽ làm giảm hiệu suất đáng kể. Việc thay thế GCN bằng sự tự chú ý còn làm xấu hơn mô hình đáng kể mà giúp chứng minh thêm về lợi ích của các GCN để chọn ngữ cảnh thích hợp đối với việc học biểu diễn trong khi phát hiện phép ẩn dụ.

Môđun quan trọng khác theo sáng chế này là chương trình khung học đa nhiệm giữa WSD và học ẩn dụ. Để chứng minh tính hiệu quả của môđun này để phát hiện phép ẩn dụ, phần này đánh giá kỹ thuật cơ bản sau để huấn luyện các mô hình: (1) Một mạng duy nhất: chỉ một mạng duy nhất để phát hiện phép ẩn dụ được huấn luyện (nghĩa là, hoàn thành việc bỏ qua mạng đối với WSD), (2) Huấn luyện trước: một mạng duy nhất được huấn luyện đối với tập dữ liệu WSD trước tiên và sau đó huấn luyện lại đối với tập dữ liệu ẩn dụ sau đó, và (3) Thay thế: một mô hình duy nhất được huấn luyện đối với cả phát hiện phép ẩn dụ và WSD, theo thủ tục

huấn luyện luân phiên. Bảng 3 thể hiện hiệu suất của các phương pháp trên bộ kiểm tra VUA VERB.

Phương pháp	P	R	F1	Acc
Mul_GCN (được đề xuất)	72,5	70,9	71,7	83,2
Một mạng duy nhất	69,7	68,1	68,9	81,5
Huấn luyện trước	70,8	68,5	69,6	82,1
Thay thế	72,9	65,4	69,0	82,3

<Bảng 3>

Như được thể hiện từ bảng 3, chương trình khung học đa nhiệm có thể cải thiện đáng kể một phương pháp mạng duy nhất với điểm số F1 tốt hơn đáng kể. Điều này minh lợi ích của WSD để phát hiện phép ẩn dụ. Hiển nhiên là phương pháp Mul GCN được đề xuất vượt trội đáng kể so với các ngưỡng cơ sở học đa nhiệm nhờ biên lớn trên điểm số F1 (nghĩa là, cải thiện lên đến 2,7% so với F1 tuyệt đối), nhờ đó chứng thực các ưu điểm của cơ chế học đa nhiệm theo sáng chế này để phát hiện phép ẩn dụ.

Để đạt được tính tương tự của hai vectơ V^{wsd} và V^{md} trong mô đun đa nhiệm (nghĩa là, phương trình 8), mô hình được đề xuất sử dụng sai số bình phương trung bình (mean squared error, MSE) là phép đo về sự không đồng dạng cần được cực tiểu hóa qua hàm mất mát tổng thể như sau:

$$M = \|V^{wsd} - V^{md}\|_2^2$$

Trong thực tế, có một số phép đo không đồng dạng thay thế M có thể được thêm vào hàm mất mát đối với mục đích này.

Trong phần này, các phép đo không đồng dạng M sau để chuyển kiến thức vào mô đun học đa nhiệm được khảo sát thêm để hiểu rõ hơn về hiệu quả của các lựa chọn phép đo này đối với mô hình theo sáng chế này dưới dạng:

· Phân kỳ Kullback-Leibler (KL):

$$M = KL(S^{wsd}, S^{md}) = -\sum_i S_i^{wsd} \log \frac{S_i^{wsd}}{S_i^{md}}$$

(trong đó $S^{wsd} = \text{softmax}(V^{wsd})$ và $S^{md} = \text{softmax}(V^{md})$)

·Cosin (Cosin):

$$M = 1 - \cos(V^{wsd}, V^{md})$$

·Mất mát biên (Biên):

$$M = 1 - s_{wsd} + s_{md}$$

(trong đó $S^{wsd} = \text{sigmoid}(\text{FF}(V^{wsd}))$ và $S^{md} = \text{sigmoid}(\text{FF}(V^{md}))$ với FF là hàm chuyển tiếp –cấp liệu để biến đổi các vectơ V^{wsd} và V^{md} thành vô hướng)

Bảng 4 báo cáo hiệu suất của phép đo không đồng dạng này trên bộ kiểm tra VUA VERB khi chúng được sử dụng trong mô hình được đề xuất (nghĩa là, thay thế phép đo MSE). Điều này rõ ràng từ bảng mà MSE tốt hơn đáng kể so với phép đo không đồng dạng khác đối với mô hình, khẳng định đối với sự lựa chọn MSE của họ theo sáng chế này.

Phép đo	P	R	F1	Acc
MSE (được đề xuất)	72,5	70,9	71,7	83,2
KL	73,8	66,0	69,7	82,8
Cosin	71,5	65,5	68,4	81,8
Biên	72,6	64,4	68,3	82,0

<Bảng 4>

Theo một phương án, có hiệu quả kỹ thuật mà vectơ ẩn có thể được tạo ra bằng cách sử dụng kỹ thuật nhúng từ và BiLSTM và vectơ ẩn có thể được nhập vào GCN được tạo cấu trúc để tính toán trên cây phân tích cú pháp phụ thuộc để tạo ra vectơ biểu diễn GCN.

Hơn nữa, có hiệu quả kỹ thuật trong đó từ mục tiêu của MD trong văn bản có thể được nhận diện cụ thể hơn bằng cách điều chỉnh vectơ biểu diễn GCN ở mức độ đặc trưng và/hoặc chiều.

Hơn nữa, có hiệu quả kỹ thuật mà, để cải thiện tính tương tự giữa các vectơ biểu diễn của mạng kép ở cùng câu, hai bộ mã hóa của MD và giải quyết WSD có thể được sử dụng để thực hiện giải quyết WSD và MD một cách độc lập.

Theo sáng chế này, ở hệ thống và phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên, các phép diễn đạt ẩn dụ và từ có thể được xác định trong văn bản.

Hơn nữa, học sâu có thể được sử dụng để phát hiện phép diễn đạt ẩn dụ với độ chính xác cao.

Hơn nữa, cơ chế mới có thể được tạo ra để cải thiện hiệu suất của mô hình học sâu để phát hiện phép diễn đạt ẩn dụ.

Trong khi sáng chế đã được mô tả dựa vào các phương án làm ví dụ của sáng chế, sáng chế này sẽ được hiểu bởi người có hiểu biết trung bình trong lĩnh vực kỹ thuật mà các thay đổi và sửa đổi khác nhau có thể được thực hiện mà không lệch khỏi mục đích và phạm vi của sáng chế như được định nghĩa bởi các điểm yêu cầu bảo hộ kèm theo.

YÊU CẦU BẢO HỘ

1. Bộ mã hóa để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên (natural language processing, NLP) bao gồm:

môđun mạng nơron tích chập dạng đồ thị (graph convolutional neural network, GCN) được tạo cấu hình để mã hóa liên kết giữa từ mục tiêu và từ cốt lõi để tạo ra các vectơ biểu diễn GCN;

môđun điều khiển được tạo cấu hình để điều chỉnh các vectơ biểu diễn GCN để tạo ra vectơ biểu diễn tổng thể;

trong đó môđun GCN, để mã hóa liên kết giữa từ mục tiêu và từ cốt lõi, cấp các vectơ biểu diễn BiLSTM vào mạng nơron tích chập dạng đồ thị (GCN) được tạo cấu trúc để tính toán trên cây phụ thuộc của câu đầu vào;

trong đó môđun điều khiển, để điều chỉnh các vectơ biểu diễn GCN,

áp dụng hàm kích hoạt (ví dụ, đơn vị tuyến tính được tinh chỉnh (rectified linear unit, ReLU)) cho vectơ biểu diễn BiLSTM của từ mục tiêu để lần lượt tạo ra vectơ điều khiển BiLSTM và vectơ điều khiển GCN;

áp dụng vectơ điều khiển BiLSTM và vectơ điều khiển GCN cho các vectơ biểu diễn BiLSTM và các vectơ biểu diễn GCN qua phép nhân từng phần tử để lần lượt tạo ra các vectơ BiLSTM được lọc và các vectơ GCN được lọc; và

tổng hợp các vectơ GCN được lọc để tạo ra vectơ biểu diễn tổng thể.

2. Bộ mã hóa theo điểm 1, còn bao gồm:

mạng nơron truyền thẳng sử dụng vectơ biểu diễn tổng thể làm đầu vào và tính toán sự phân bố xác suất của việc liệu từ mục tiêu có là phép diễn đạt ẩn dụ hoặc phép diễn đạt không ẩn dụ hay không bằng cách sử dụng lớp softmax;

trong đó hàm mất mát để huấn luyện là hàm log hợp lý âm cho tập dữ liệu học.

3. Hệ thống để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên bao gồm:

môđun mã hóa được tạo cấu hình để biến đổi các từ có trong câu thành các véctơ biểu diễn BiLSTM;

bộ mã hóa thứ nhất được tạo cấu hình để tạo ra véctơ biểu diễn tổng thể thứ nhất của tác vụ giải quyết làm rõ nghĩa từ (Word Sense Disambiguation, WSD);

bộ mã hóa thứ hai được tạo cấu hình để tạo ra véctơ biểu diễn tổng thể thứ hai của tác vụ phát hiện phép ẩn dụ (metaphor detection, MD), trong đó các bộ mã hóa thứ nhất và thứ hai có cùng kiến trúc mạng làm bộ mã hóa theo điểm 1; và

môđun học đa nhiệm được tạo cấu hình để thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất và thứ hai.

4. Hệ thống theo điểm 3, trong đó môđun mã hóa, để biến đổi các từ có trong câu thành các véctơ biểu diễn BiLSTM,

ghép các véctơ được tính toán bằng cách sử dụng kỹ thuật nhúng từ không có ngữ cảnh đã được huấn luyện trước, nhúng từ có ngữ cảnh và nhúng chỉ số thành véctơ được ghép; và

thực thi mạng BiLSTM trên véctơ được ghép nêu trên để tạo ra các véctơ biểu diễn BiLSTM.

5. Hệ thống theo điểm bất kỳ trong số các điểm 3 hoặc 4 trong đó môđun học đa nhiệm bao gồm bộ phân loại theo tác vụ cụ thể được tạo cấu hình để tính toán sự phân bố xác suất của các nhãn có thể đối với mỗi tác vụ.

6. Hệ thống theo điểm 5, trong đó môđun học đa nhiệm cực tiểu hóa hàm mất mát sau:

$$C(w^t, p^t, y^t) = -\log P^t(y^t | w^t, p^t) + \lambda \|V^{wsd} - V^{md}\|_2^2$$

để cập nhật các tham số cho bộ phân loại theo tác vụ cụ thể và các bộ mã hóa thứ nhất và thứ hai.

7. Hệ thống theo điểm bất kỳ trong số các điểm từ 3 đến 6, trong đó quy trình huấn luyện luân phiên được sử dụng để huấn luyện hệ thống.

8. Phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên, phương pháp bao gồm các bước:

mã hóa, bởi môđun mạng nơron tích chập dạng đồ thị (GCN), liên kết giữa từ mục tiêu và từ cốt lõi để tạo ra các vectơ biểu diễn GCN;

điều chỉnh, bởi môđun điều khiển, các vectơ biểu diễn GCN để tạo ra vectơ biểu diễn tổng thể;

trong đó bước mã hóa liên kết giữa từ mục tiêu và từ cốt lõi bao gồm bước cấp các vectơ biểu diễn BiLSTM vào mạng nơron tích chập dạng đồ thị (GCN) được tạo cấu trúc để tính toán trên cây phụ thuộc của câu đầu vào;

trong đó bước điều chỉnh các vectơ biểu diễn GCN bao gồm các bước:

áp dụng hàm kích hoạt (ví dụ, đơn vị tuyến tính được tinh chỉnh (ReLU)) cho vectơ biểu diễn BiLSTM của từ mục tiêu để lần lượt tạo ra vectơ điều khiển BiLSTM và vectơ điều khiển GCN;

áp dụng vectơ điều khiển BiLSTM và vectơ điều khiển GCN cho các vectơ biểu diễn BiLSTM và vectơ biểu diễn GCN qua phép nhân từng phần tử để lần lượt tạo ra các vectơ BiLSTM được lọc và vectơ GCN được lọc; và

tổng hợp các vectơ GCN được lọc để tạo ra vectơ biểu diễn tổng thể.

9. Phương pháp theo điểm 8, phương pháp này còn bao gồm các bước:

nhập vectơ biểu diễn tổng thể vào mạng nơron truyền thẳng và tính toán sự phân bố xác suất của việc liệu từ mục tiêu có là phép diễn đạt ẩn dụ hoặc phép diễn đạt không ẩn dụ hay không bằng cách sử dụng lớp softmax;

trong đó hàm mất mát để huấn luyện là hàm log hợp lý âm cho tập dữ liệu học.

10. Phương pháp để phát hiện phép ẩn dụ trong xử lý ngôn ngữ tự nhiên bao gồm các bước:

biến đổi, bởi môđun mã hóa, các từ có trong câu thành các vectơ biểu diễn BiLSTM;

tạo ra, bởi bộ mã hóa thứ nhất, vectơ biểu diễn tổng thể thứ nhất của tác vụ giải quyết WSD;

tạo ra, bởi bộ mã hóa thứ hai, vectơ biểu diễn tổng thể thứ hai của tác vụ MD, trong đó các bộ mã hóa thứ nhất và thứ hai được tạo cấu hình để thực hiện phương pháp theo điểm 8; và

thực hiện, bởi môđun học đa nhiệm, chuyển kiến thức giữa các bộ mã hóa thứ nhất và thứ hai.

11. Phương pháp theo điểm 10, trong đó bước biến đổi các từ có trong câu thành các vectơ biểu diễn BiLSTM bao gồm các bước:

ghép các vectơ được tính toán bằng cách sử dụng kỹ thuật nhúng từ không có ngữ cảnh đã được huấn luyện trước, nhúng từ có ngữ cảnh và nhúng chỉ số thành vectơ được ghép; và

thực thi mạng BiLSTM trên vectơ được ghép nêu trên để tạo ra các vectơ biểu diễn BiLSTM.

12. Phương pháp theo điểm bất kỳ trong số các điểm từ 10 đến 11, trong đó bước thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất và thứ hai bao gồm bước tính toán, bởi bộ phân loại theo tác vụ cụ thể, sự phân bố xác suất của các nhãn có thể đối với mỗi tác vụ.

13. Phương pháp theo điểm 12, trong đó bước thực hiện chuyển kiến thức giữa các bộ mã hóa thứ nhất và thứ hai bao gồm bước cực tiểu hóa, bởi môđun học đa nhiệm, hàm mất mát sau:

$$C(w^t, p^t, y^t) = -\log P^t(y^t | w^t, p^t) + \lambda \|V^{wsd} - V^{md}\|_2^2$$

để cập nhật các tham số cho bộ phân loại theo tác vụ cụ thể và các bộ mã hóa thứ nhất và thứ hai.

14. Phương pháp theo điểm bất kỳ trong số các điểm từ 10 đến 13, trong đó phương pháp này sử dụng quy trình huấn luyện luân phiên để huấn luyện.

FIG.1

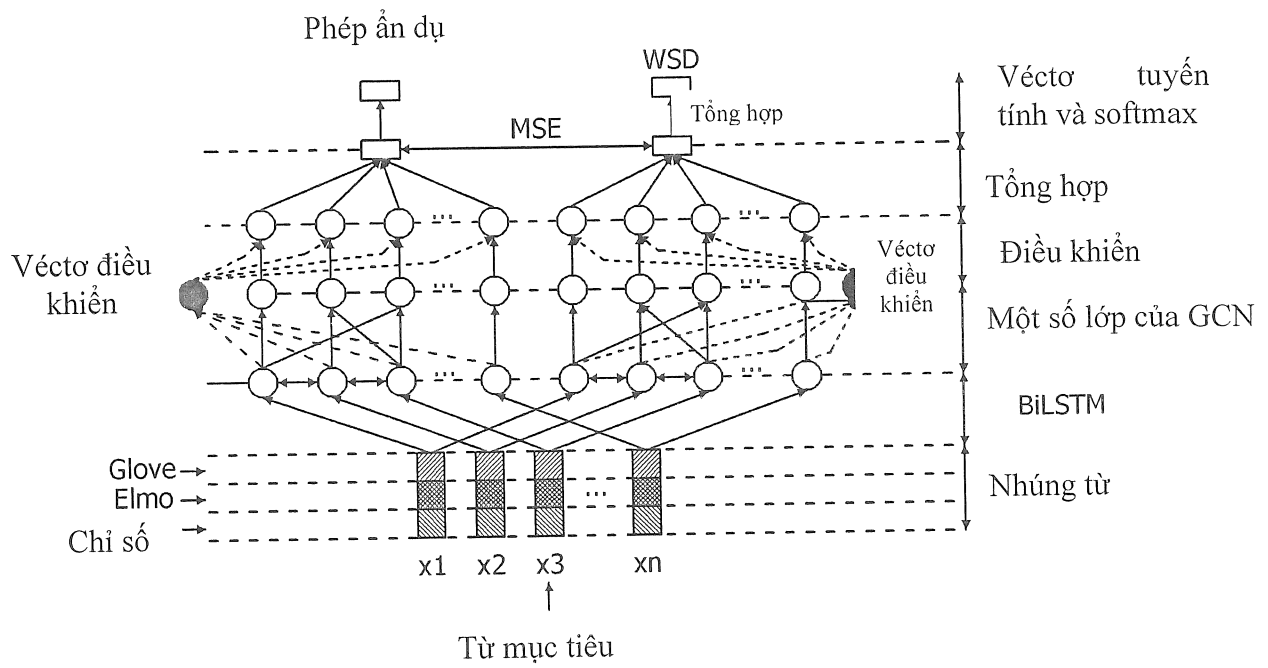


FIG.2

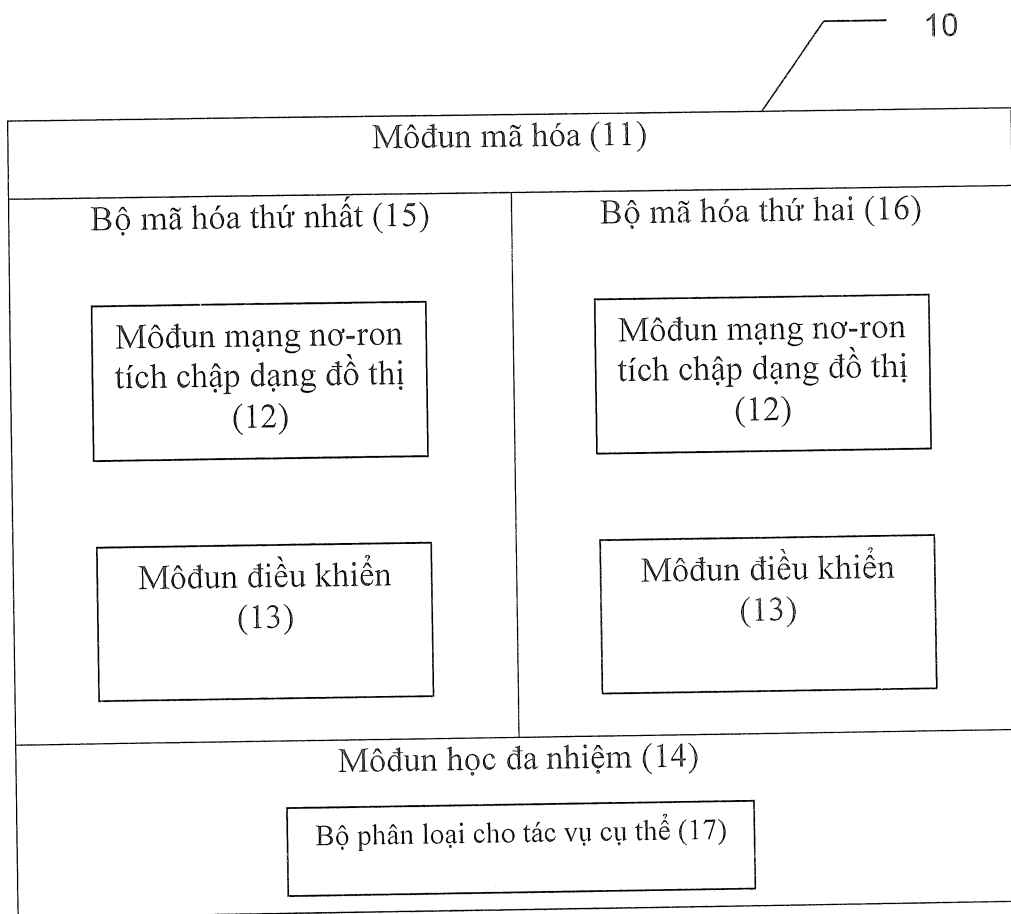


FIG.3

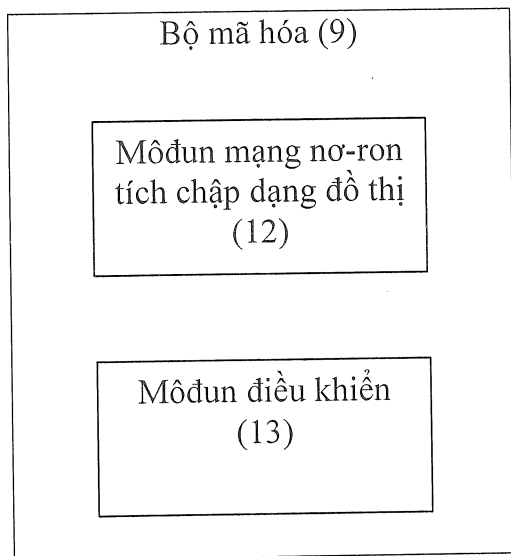


FIG.4

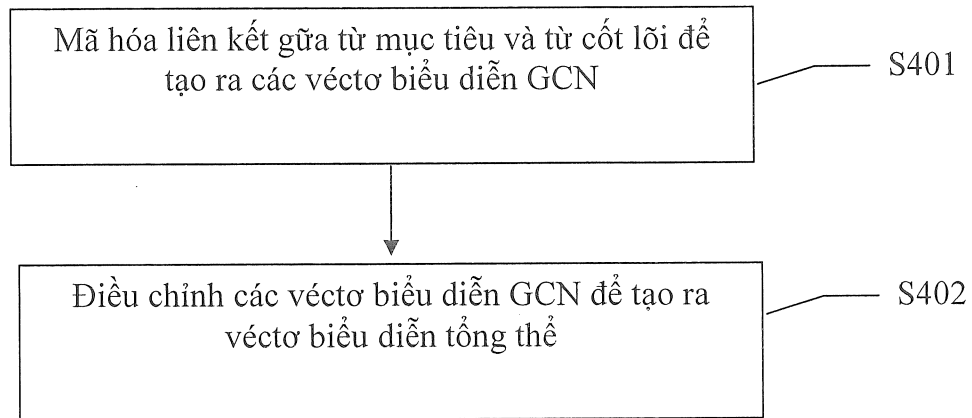


FIG.5

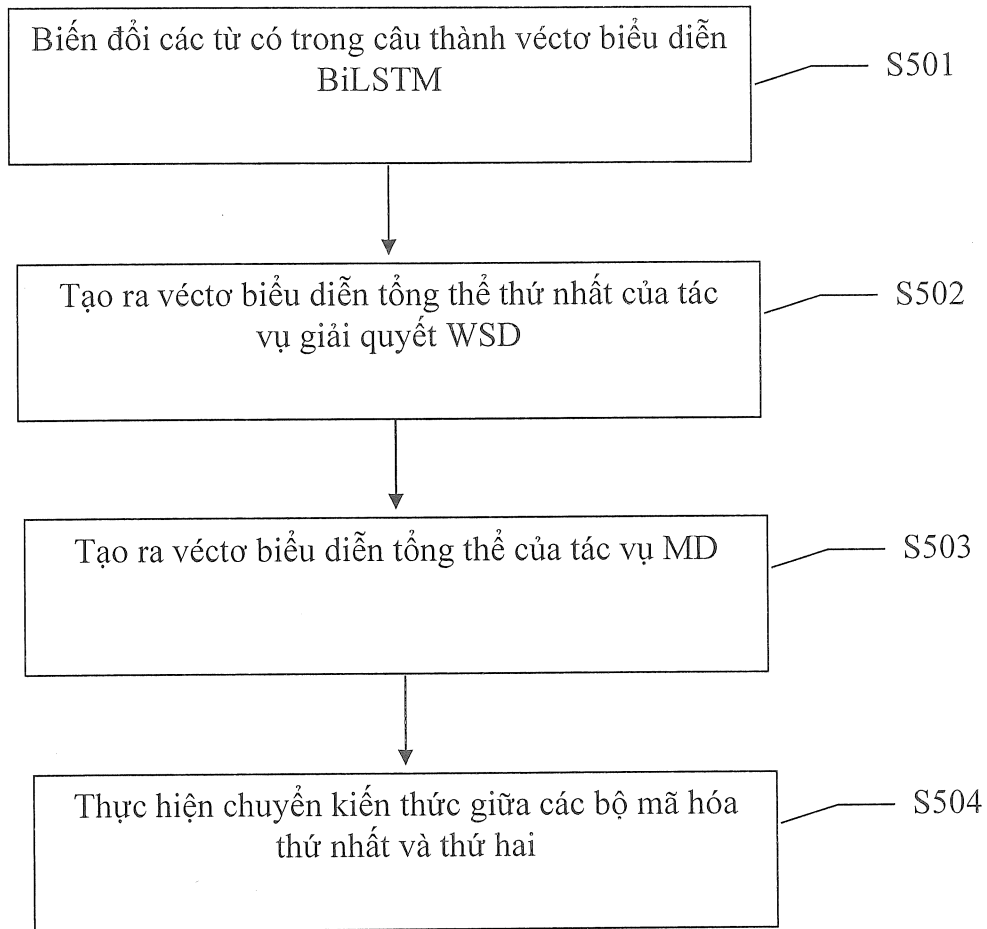


FIG.6

