

Mạng tích chập đồ thị cho phân tích tình cảm phản hồi của sinh viên Việt Nam

Nguyễn Thị Phong Dung

Khoa Công nghệ Thông tin – Trường Đại học Nguyễn Tất Thành, Thành phố Hồ Chí Minh, Việt Nam

ntpdung@ntt.edu.vn

Tóm tắt

Phân tích tình cảm đóng vai trò quan trọng trong việc hiểu ý kiến của sinh viên và nâng cao chất lượng dịch vụ giáo dục. Tuy nhiên, phân tích tình cảm tiếng Việt vẫn còn nhiều thách thức do đặc điểm ngôn ngữ và sự hạn chế về dữ liệu được chú thích. Nghiên cứu này đề xuất một phương pháp học sâu dựa trên đồ thị sử dụng mạng tích chập đồ thị (GCN – Graph Convolutional Networks) để phân loại tình cảm phản hồi của sinh viên Việt Nam. Mỗi câu đầu vào được mô hình hóa như một đồ thị đồng xuất hiện từ, trong đó các từ được biểu diễn dưới dạng các nút và các cạnh được xây dựng dựa trên cửa sổ trượt cố định. Các embedding từ được học từ đầu đến cuối trong khuôn khổ GCN, cho phép học biểu diễn ngữ nghĩa hiệu quả. Các thí nghiệm được thực hiện trên corpus phản hồi của sinh viên Việt Nam thuộc Trường Đại học Công nghệ Thông tin. Mô hình đề xuất đạt độ chính xác 88,5 % và F1 là 71,3 %. Kết quả thực nghiệm cho thấy, các mô hình mạng nơ-ron dựa trên đồ thị là một hướng đi đầy triển vọng cho phân tích tình cảm tiếng Việt, đặc biệt là trong việc khai thác phản hồi giáo dục.

© 2026 Journal of Science and Technology - NTTU

Nhận 18/01/2026

Được duyệt 09/02/2026

Công bố 28/07/2026

Từ khóa

Phân tích tình cảm;
mạng nơ-ron tích chập
đồ thị; tiếng Việt;
phản hồi của sinh viên;
học sâu.

1 Giới thiệu

Phân tích tình cảm đã trở thành một nhiệm vụ thiết yếu trong xử lý ngôn ngữ tự nhiên (NLP – Natural Language Processing), nhằm mục đích xác định và phân loại các ý kiến chủ quan được thể hiện trong dữ liệu văn bản. Trong lĩnh vực giáo dục, phân tích tình cảm phản hồi của sinh viên cung cấp những hiểu biết có giá trị để đánh giá chất lượng giảng dạy, hiệu quả khóa học và trải nghiệm học tập tổng thể.

Mặc dù đã có những tiến bộ đáng kể trong phân tích tình cảm đối với các ngôn ngữ có nguồn tài nguyên phong phú như tiếng Anh, phân tích tình cảm tiếng Việt vẫn là một vấn đề đầy thách thức. Những thách thức này xuất phát từ các đặc điểm ngôn ngữ của

tiếng Việt, bao gồm trật tự từ linh hoạt, biểu hiện tình cảm ngầm và sự khan hiếm các tập dữ liệu được chú thích quy mô lớn. Các nghiên cứu gần đây đã chỉ ra rằng các phương pháp học máy truyền thống, dựa nhiều vào các đặc trưng được tạo thủ công, thường gặp khó khăn trong việc nắm bắt các mối quan hệ ngữ nghĩa phức tạp trong văn bản tiếng Việt [1].

Với sự phát triển nhanh chóng của học sâu, các mô hình như mạng nơ-ron hồi quy (RNN – Recurrent Neural Network) và mạng nơ-ron tích chập (CNN – Convolutional Neural Network) đã được áp dụng rộng rãi cho các nhiệm vụ phân tích tình cảm. Mặc dù các mô hình dựa trên chuỗi này đã cho thấy kết quả đầy hứa hẹn, nhưng chúng chủ yếu tập trung vào trật tự từ

tuyến tính và có thể không nắm bắt được các phụ thuộc ngữ nghĩa tầm xa hoặc phi cục bộ giữa các từ [2]. Các mô hình transformer (như PhoBERT) dù mạnh mẽ nhưng đòi hỏi tài nguyên tính toán lớn và thường coi văn bản là chuỗi phẳng. Việc sử dụng GCN cho phép mô hình hóa cấu trúc đồ thị không gian, giúp bắt được các phụ thuộc phi tuyến tính giữa các từ mà không cần kiến trúc quá phức tạp. Để khắc phục hạn chế này, mạng nơ-ron dựa trên đồ thị, đặc biệt là GCN, đã nổi lên như một phương pháp hiệu quả để mô hình hóa các mối quan hệ cấu trúc trong văn bản. Bằng cách biểu diễn câu dưới dạng đồ thị trong đó các từ được coi là các nút và các cạnh mã hóa các mối quan hệ cú pháp hoặc ngữ nghĩa, GCN có thể tổng hợp thông tin ngữ cảnh từ các từ lân cận và học được các biểu diễn ngữ nghĩa phong phú hơn [3, 4].

Các công trình gần đây đã chứng minh hiệu quả của các mô hình dựa trên GCN trong cả thiết lập phân tích tình cảm đơn ngữ và đa ngữ [5]. Nghiên cứu này, đề xuất một mô hình phân tích tình cảm dựa trên GCN cho phản hồi của sinh viên Việt Nam. Mỗi câu được biểu diễn dưới dạng đồ thị đồng xuất hiện từ được xây dựng bằng cơ chế cửa sổ trượt. Mô hình đề xuất được đánh giá trên corpus phản hồi của sinh viên Việt Nam thuộc Trường Đại học Công nghệ Thông tin (UIT Vietnamese Student Feedback Corpus – UIT-VSFC). Mục tiêu chính của nghiên cứu này gồm:

- Đề xuất một khung phân tích tình cảm dựa trên đồ thị cho tiếng Việt bằng cách sử dụng GCN.
- Thực nghiệm mở rộng trên một tập dữ liệu phản hồi của sinh viên Việt Nam công khai.
- Phân tích chuyên sâu về hiệu suất mô hình và các mẫu lỗi, làm nổi bật những thách thức của việc phân loại tình cảm trung tính.

Khác với các mô hình transformer như PhoBERT vốn xử lý văn bản theo dạng chuỗi phẳng, GCN cho phép nắm bắt các quan hệ ngữ nghĩa không gian giữa các từ thông qua cấu trúc đồ thị, giúp giải quyết tốt hơn bài toán phân tích tình cảm trong các câu phản hồi ngắn và có cấu trúc phi chính thức của sinh viên.

2 Tổng quan nghiên cứu

2.1 Phân tích tình cảm tiếng Việt

Các nghiên cứu ban đầu về phân tích tình cảm tiếng Việt chủ yếu dựa vào các phương pháp học máy

truyền thống như máy hỗ trợ vector (SVM – Support Vector Machine) và bộ phân loại Naïve Bayes kết hợp với các đặc trưng từ vựng và cú pháp được tạo thủ công. Mặc dù các phương pháp này đạt được hiệu suất hợp lý, nhưng chúng đòi hỏi kỹ thuật đặc trưng rộng rãi và rất nhạy cảm với nhiễu và ngôn ngữ không chính thức thường thấy trong phản hồi của sinh viên và văn bản trên mạng xã hội [6]. Với sự phát triển của học sâu, các mô hình dựa trên mạng nơ-ron như mạng bộ nhớ ngắn-dài hạn (LSTM – Long Short-Term Memory), LSTM hai chiều và kiến trúc dựa trên CNN ngày càng được áp dụng cho các nhiệm vụ phân loại tình cảm tiếng Việt [7, 8]. Các mô hình này cải thiện đáng kể hiệu suất bằng cách tự động học các biểu diễn từ phân tán và các đặc trưng ngữ cảnh. Tuy nhiên, vì chúng xử lý văn bản dưới dạng chuỗi tuyến tính, chúng thường gặp khó khăn trong việc nắm bắt các phụ thuộc tầm xa và các mối quan hệ ngữ nghĩa phức tạp, đặc biệt là trong các câu có biểu hiện tình cảm ngầm [9].

2.2 Mạng nơ-ron dựa trên đồ thị cho xử lý ngôn ngữ tự nhiên

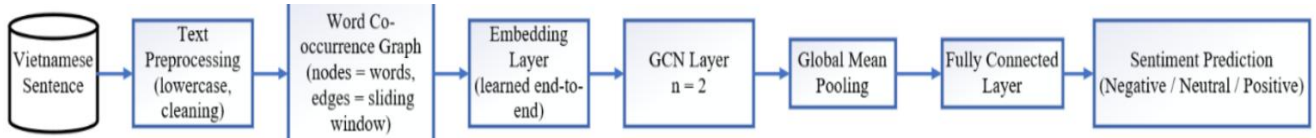
Mạng nơ-ron đồ thị (GNN – Graph Neural Network) gần đây đã nổi lên như một khung mạnh mẽ để mô hình hóa dữ liệu có cấu trúc và đã được áp dụng thành công cho nhiều nhiệm vụ xử lý ngôn ngữ tự nhiên khác nhau, bao gồm phân loại văn bản, trích xuất mối quan hệ và phân tích tình cảm [10]. Trong số đó, mạng nơ-ron tích chập đồ thị (GCN) mở rộng các phép toán tích chập cho đầu vào có cấu trúc đồ thị, cho phép các nút tổng hợp thông tin từ các vùng lân cận cục bộ của chúng. Trong phân tích tình cảm, các phương pháp dựa trên đồ thị thường biểu diễn văn bản dưới dạng đồ thị phụ thuộc hoặc đồ thị đồng xuất hiện từ, cho phép các mô hình nắm bắt các mối quan hệ cú pháp và ngữ nghĩa vượt ra ngoài thứ tự từ tuần tự [11, 12]. Một số nghiên cứu đã chứng minh rằng, các mô hình dựa trên GCN vượt trội hơn các kiến trúc dựa trên chuỗi trong việc nắm bắt các phụ thuộc ngữ cảnh không cục bộ. Tuy nhiên, nghiên cứu về việc áp dụng mạng nơ-ron dựa trên đồ thị vào phân tích tình cảm tiếng Việt, đặc biệt là trong lĩnh vực phản hồi giáo dục, vẫn còn tương đối hạn chế [1], điều này thúc đẩy phương pháp được đề xuất trong công trình này.

3 Phương pháp đề xuất

3.1 Tổng quan kiến trúc

Quy trình phân tích tình cảm dựa trên GCN được đề xuất bao gồm năm giai đoạn chính: tiền xử lý, xây dựng đồ thị, học biểu diễn, gộp và phân loại. Văn bản phản hồi của sinh viên Việt Nam được tiền xử lý và mã hóa thành từ. Mỗi câu sau đó được biểu diễn dưới dạng đồ thị đồng xuất hiện từ, trong đó các nút tương ứng với các từ và các cạnh được hình thành dựa trên cửa sổ

trượt có kích thước cố định. Các từ nhúng được học từ đầu đến cuối và được sử dụng làm đặc trưng nút cho nhiều lớp GCN, cho phép mô hình nắm bắt các mối quan hệ ngữ cảnh và cấu trúc giữa các từ. Một phép toán gộp trung bình toàn cục tổng hợp các đặc trưng cấp nút thành một biểu diễn câu có độ dài cố định, cuối cùng được chuyển đến bộ phân loại softmax để dự đoán nhãn tình cảm (tiêu cực, trung tính hoặc tích cực). Quy trình đề xuất được minh họa trong Hình 1.



Hình 1 Quy trình tổng thể của mô hình phân tích tình cảm dựa trên GCN được đề xuất

3.2 Đề xuất bài toán

Xét bài toán phân tích tình cảm văn bản tiếng Việt trên tập dữ liệu phản hồi sinh viên UIT-VSFC. Cho tập dữ liệu huấn luyện:

$$D = \{(S_i, y_i)\}_{i=1}^N$$

trong đó mỗi mẫu dữ liệu bao gồm: $S_i = \{w_1, w_2, \dots, w_{n_i}\}$ là một câu văn bản tiếng Việt, được biểu diễn dưới dạng chuỗi các từ; $y_i \in \{0, 1, 2\}$ là nhãn tình cảm tương ứng với ba lớp: tiêu cực (negative), trung tính (neutral) và tích cực (positive).

Khác với các phương pháp truyền thống chỉ xem câu văn bản như một chuỗi tuyến tính, trong nghiên cứu này, mỗi câu S_i được biểu diễn dưới dạng một đồ thị có cấu trúc:

$$G_i = (V_i, E_i)$$

trong đó: V_i là tập các nút đại diện cho các từ trong câu; E_i là tập các cạnh mô tả mối quan hệ ngữ nghĩa và/hoặc cú pháp giữa các từ, được xây dựng dựa trên thông tin ngữ cảnh.

Trên mỗi đồ thị G_i , mô hình học sâu dựa trên đồ thị được sử dụng để học biểu diễn đặc trưng toàn cục của câu: $h_i = \Phi(G_i)$

trong đó $\Phi(\cdot)$ là hàm ánh xạ từ đồ thị sang không gian biểu diễn tiềm ẩn.

Mục tiêu của bài toán là xác định một hàm dự đoán:

$$f: S_i \rightarrow y_i$$

sao cho hàm mất mát được tối thiểu hóa trên tập huấn luyện, đồng thời tối ưu hóa độ chính xác (Accuracy) và chỉ số F1-macro trên tập dữ liệu kiểm thử.

Bài toán này có thể được xem như một bài toán phân loại tình cảm đa lớp (*multi-class sentiment classification*) với đầu vào là đồ thị văn bản có cấu trúc, nhằm khai thác tốt hơn các mối quan hệ ngữ nghĩa trong tiếng Việt.

3.3 Tiền xử lý và nhúng từ

Cho một câu văn bản đầu vào $S = \{w_1, w_2, \dots, w_n\}$, quá trình tiền xử lý bao gồm chuẩn hóa văn bản, loại bỏ nhiễu và ánh xạ các từ về một từ điển cố định V với kích thước $|V| = 50.000$.

Mỗi từ w_i được ánh xạ sang một vector liên tục thông qua ma trận nhúng:

$$E \in \mathbb{R}^{|V| \times d}$$

Trong đó: $d = 300$ là chiều của không gian nhúng.

Biểu diễn vector của từ w_i được xác định là:

$$x_i = E(w_i) \in \mathbb{R}^d$$

Kết quả, câu văn bản S được biểu diễn dưới dạng ma trận đặc trưng ban đầu:

$$X = [x_1, x_2, \dots, x_n]^T \in \mathbb{R}^{n \times d}$$

3.4 Xây dựng đồ thị ngữ cảnh

Để nắm bắt các mối quan hệ cục bộ giữa các từ trong câu, văn bản được chuyển đổi thành một đồ thị vô hướng:

$$G = (V, E)$$

Trong đó: tập nút V bao gồm n nút, mỗi nút tương ứng với một từ trong câu. Tập cạnh E được xây dựng dựa trên chiến lược cửa sổ trượt (sliding window) với kích thước $k = 2$ với đặc thù phản hồi của sinh viên thường ngắn và trực diện, giúp tập trung vào các quan hệ từ

vững cục bộ mạnh nhất, tránh tạo ra các cạnh nhiễu giữa các từ quá xa nhau.

Cụ thể, một cạnh $(i, j) \in E$ tồn tại nếu:

$$|idx(w_i) - idx(w_j)| \leq k$$

trong đó $idx(w)$ biểu thị vị trí của từ w trong câu.

Chiến lược này cho phép mô hình khai thác các mối quan hệ ngữ cảnh cục bộ giữa các từ lân cận trong văn bản tiếng Việt.

3.5 Các tầng tích chập đồ thị

Mô hình sử dụng hai tầng GCN để lan truyền và cập nhật đặc trưng giữa các nút trong đồ thị, với số lớp GCN = 2 cho phép thông tin lan truyền đến các “nút lân cận của lân cận”, đủ để bao quát ngữ cảnh của một câu phản hồi trung bình mà không gây ra hiện tượng “quá mịn” (over-smoothing). Quy trình lan truyền đặc trưng tại tầng thứ l được xác định bởi công thức sau:

$$H^{(l+1)} = \sigma\left(\tilde{D}^{-\frac{1}{2}} \tilde{A} \tilde{D}^{-\frac{1}{2}} H^{(l)} W^{(l)}\right)$$

trong đó: A là ma trận kề của đồ thị G ; $\tilde{A} = A + I_n$ là ma trận kề mở rộng với các vòng lặp tự thân; $\tilde{D}$ là ma trận bậc của $\tilde{A}$; $W^{(l)}$ là ma trận trọng số có thể học được tại tầng l ; $\sigma(\cdot)$ là hàm kích hoạt phi tuyến ReLU. $H^{(l+1)}$ là ma trận biểu diễn đặc trưng của các nút tại tầng (lớp) thứ l trong mạng GCN, trong đó mỗi hàng tương ứng với vector đặc trưng của một nút sau khi được cập nhật ở tầng này.

Đặc trưng đầu vào của tầng GCN đầu tiên được xác định là:

$$H^{(0)} = X$$

Sau hai tầng GCN, mỗi nút được biểu diễn bằng một vector ẩn có chiều $h = 128$.

3.6 Gộp đặc trưng toàn cục

Để thu được biểu diễn cố định cho toàn bộ câu văn bản, mô hình áp dụng Gộp trung bình toàn cục (Global Mean Pooling) trên tập các nút để giữ lại thông tin trung bình của tất cả các từ trong câu, phù hợp cho bài toán phân loại tình cảm tổng quát:

$$h_G = \frac{1}{n} \sum_{i=1}^n h_i$$

trong đó $h_i \in \mathbb{R}^{128}$ là vector đặc trưng của nút thứ i sau các tầng GCN. Vector h_G đóng vai trò là biểu diễn ngữ nghĩa toàn cục của câu văn bản.

3.7 Phân loại và tối ưu hóa

Biểu diễn câu h_G được đưa qua một tầng tuyến tính để dự đoán phân phối xác suất của các lớp tình cảm:

$$\hat{y} = \text{Softmax}(W_c h_G + b_c)$$

Trong đó, $W_c \in \mathbb{R}^{3 \times 128}$ và b_c là các tham số của bộ phân loại; $\hat{y}$ là vector xác suất dự đoán cho ba lớp tình cảm.

Hàm mất mát được sử dụng là Cross-Entropy Loss, được định nghĩa như sau:

$$\mathcal{L} = - \sum_{c=1}^3 y_c \log(\hat{y}_c)$$

Trong đó: y_c là giá trị nhãn thực tế của lớp thứ c ở dạng one-hot (tức là chỉ có một phần tử bằng 1 và các phần tử còn lại bằng 0); $\hat{y}_c$ là xác suất dự đoán của mô hình cho lớp thứ c , được tính từ hàm Softmax; Tổng được tính trên 3 lớp cảm xúc: tích cực, tiêu cực và trung tính.

Thuật toán 1: xây dựng đồ thị ngữ cảnh từ văn bản

Thuật toán này chuyển đổi một câu văn bản thô thành cấu trúc đồ thị phù hợp cho GCN.

Algorithm 1: Text-to-Graph Conversion

Input:

- Câu văn bản thô S
- Từ điển $Vocab$
- Kích thước cửa sổ trượt k

Output:

- Đồ thị $G = (X, E, y)$

1. CleanText:

Chuyển câu S sang chữ thường; loại bỏ URL, ký tự đặc biệt và nhiễu.

2. Tokenization:

Tách câu thành danh sách từ: $\{w_1, w_2, \dots, w_n\}$

3. Indexing:

Chuyển mỗi từ w_i thành chỉ số id_i dựa trên từ điển:

- Nếu $w_i \in Vocab$, gán $id_i = Vocab(w_i)$
- Ngược lại, gán $id_i = ID_{UNK}$

Thu được danh sách chỉ số: $X = \{id_1, id_2, \dots, id_n\}$

4. Initialize Edge List:

5. Graph Construction:

$E \leftarrow \emptyset$

6. for $i = 0$ to $n - 1$ do

7. for $j = i - k$ to $i + k$ do

8. if $(0 \leq j < n)$ and $(i \neq j)$ then



9. $E \leftarrow E \cup \{(i, j)\}$

10. end if

11. end for

12. end for

13. Handle Isolated Nodes:

Nếu $E = \emptyset$, thêm vòng lặp tự thân:

$E \leftarrow \{(0,0)\}$

14. Return:

$G = (X, E, y)$

Thuật toán 2: huấn luyện Mô hình GCN

Thuật toán này mô tả quá trình lan truyền xuôi và tối ưu hóa tham số của mạng GCN.

Algorithm 2: GCN Training Process

Input:

- Tập huấn luyện D_{train}
- Tốc độ học α
- Số epoch tối đa E_{max}

Output:

- Mô hình đã tối ưu θ

1. Initialization:

Khởi tạo:

- Ma trận nhúng E
- Hai tầng tích chập đồ thị $GCNConv_1, GCNConv_2$
- Bộ tối ưu Adam với learning rate α

2. Training Loop:

3. for epoch = 1 to E_{max} do

4. for each mini-batch $B \subset D_{train}$ do

5. Forward Pass:

- Nhúng từ: $h^{(0)} = E(X)$
- GCN tầng 1: $h^{(1)} = \text{ReLU}(GCNConv_1(h^{(0)}, E))$
- GCN tầng 2: $h^{(2)} = \text{ReLU}(GCNConv_2(h^{(1)}, E))$
- Gộp đặc trưng toàn cục: $h_G = \text{mean}(h^{(2)})$
- Dự đoán nhãn: $\hat{y} = \text{Linear}(h_G)$

6. Loss Computation: $\mathcal{L} = \text{CrossEntropy}(\hat{y}, y)$

7. Backpropagation & Update:

- Tính gradient của $\mathcal{L}$
- Cập nhật tham số θ bằng Adam optimizer

8. End Training:

9. end for

10. end for

11. Return: θ

4 Thực nghiệm và đánh giá

4.1 Bộ dữ liệu

Trong nghiên cứu này, thực hiện các thực nghiệm trên bộ dữ liệu UIT-VSFC [13]. Đây là bộ dữ liệu chuẩn cho bài toán phân tích tình cảm tiếng Việt, bao gồm các câu phản hồi của sinh viên được gán nhãn thủ công. Bộ dữ liệu bao gồm 16.175 mẫu phản hồi của sinh viên Việt Nam. Mỗi mẫu được phân loại vào một trong ba nhóm tình cảm: tiêu cực, trung lập và tích cực. Bộ dữ liệu được chia chính thức thành 11.426 mẫu huấn luyện, 1583 mẫu xác thực và 3166 mẫu kiểm tra, đảm bảo quá trình huấn luyện mô hình đáng tin cậy và đánh giá công bằng. Các văn bản phản hồi được viết bằng tiếng Việt tự nhiên, không trang trọng, chứa các biểu hiện chuyên ngành và ý kiến chủ quan, điều này làm cho việc phân loại tình cảm trở nên thách thức và thực tế hơn.

4.2 Độ đo đánh giá

Để đánh giá hiệu năng của mô hình GCN đề xuất, sử dụng hai chỉ số đo lường phổ biến trong các bài toán phân loại văn bản:

- Accuracy (độ chính xác tổng thể): đo lường tỷ lệ các dự đoán đúng trên tổng số mẫu thử nghiệm.
- Macro F1-score: chỉ số này được ưu tiên sử dụng để đánh giá khách quan khả năng phân loại của mô hình trên cả 3 lớp, đặc biệt khi có sự mất cân bằng về số lượng mẫu giữa các nhãn tình cảm.

Bảng 1 Tham số thực nghiệm

Thành phần	Tham số	Giá trị
Embedding (nhúng từ)	Dimension (chiều)	300
GCN Layers	Hidden Dimension (chiều ẩn)	128
Optimization (tối ưu)	Algorithm / LR	Adam/ 10^{-3}
Training (huấn luyện)	Epochs / Batch Size	30 / 32

4.3 Kết quả

Trong phần này, nghiên cứu so sánh mô hình GCN đề xuất với các kết quả thực nghiệm được công bố bởi nhóm tác giả của bộ dữ liệu UIT-VSFC. Mặc dù nghiên cứu gốc tập trung vào cả phân tích khía cạnh, nghiên cứu trích xuất kết quả phân phân tích tình cảm để làm đối chứng [13, 14].

Bảng 2 So sánh mô hình GCN

Phương pháp	Đặc trưng	Accuracy (%)	Macro F1-score (%)
Maximum Entropy	n-grams	87,00	63,30
SVM	n-grams + TF-IDF	88,10	68,20
Bi-GRU	Word Embedding	90,40	70,60
Bi-LSTM	Word Embedding	88,10	69,50
CNN	Word Embedding	87,20	68,40
PhoBERT (base)	Pretrained Transformer (PhoBERT)	90,40	72,10
GCN (Đề xuất)	Graph + Embedding	88,50	71,34

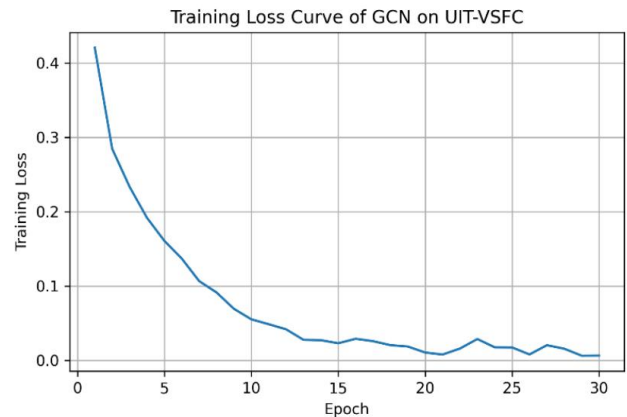
Nghiên cứu thực hiện thí nghiệm cắt tĩa thành phần (Ablation Study) để đánh giá đóng góp của các thành phần: chỉ dùng Word Embedding (không GCN): F1-score giảm xuống còn 65,2 %. GCN với 1 tầng: F1-score đạt 68,7 %. GCN với 2 tầng (đề xuất): F1-score đạt mức cao nhất 71,3 %. Kết quả này khẳng định việc sử dụng các tầng tích chập đồ thị là yếu tố then chốt giúp mô hình bắt được các phụ thuộc ngữ nghĩa tầm xa trong câu phản hồi của sinh viên.

4.4 Thảo luận

Hiệu năng vượt trội về macro F1-score: điểm đáng chú ý nhất là mô hình GCN đạt chỉ số macro F1-score là 71,34 %, cao hơn cả mô hình mạng nơ-ron hồi quy mạnh nhất trong bài báo gốc là Bi-GRU (70,60 %) và vượt xa SVM (68,20 %) cho thấy cấu trúc đồ thị giúp mô hình nhận diện tốt hơn các sắc thái tình cảm không chỉ dựa trên thứ tự từ mà còn dựa trên mối quan hệ không gian giữa các từ khóa. Chỉ số macro F1-score cao hơn chứng tỏ GCN có khả năng cân bằng tốt hơn giữa các lớp tình cảm, đặc biệt là khả năng nhận diện các lớp có ít dữ liệu (như Neutral) hiệu quả hơn các phương pháp truyền thống. Độ chính xác tương đương: Về mặt Accuracy, GCN đạt 88,50 %, cao hơn SVM (88,10 %) và chỉ thấp hơn Bi-GRU một khoảng nhỏ (~ 1,9 %). Tuy nhiên, cần lưu ý rằng GCN đạt được kết quả này với cấu trúc đồ thị đơn giản và số lượng tham số ít hơn đáng kể so với kiến trúc Bi-GRU hai chiều phức tạp. Trong khi các mô hình cơ sở (baseline) trong nghiên cứu [13] dựa trên chuỗi từ (sequences) hoặc túi từ (bag-of-words), mô hình GCN đề xuất khai thác quan hệ không gian giữa các từ thông qua cấu trúc đồ thị. Việc vượt qua các baseline này khẳng định rằng: Cấu trúc đồ thị có

thể bắt lấy các đặc trưng ngữ nghĩa của tiếng Việt tốt hơn các phương pháp thống kê truyền thống. Quan hệ lân cận trong “cửa sổ trượt” (sliding window) là đủ để đại diện cho sắc thái biểu cảm trong các câu phản hồi ngắn của sinh viên.

Phân tích quá trình huấn luyện (Training Convergence) Dựa trên biểu đồ đường cong hàm mất mát (Hình 2), mô hình GCN cho thấy khả năng hội tụ ổn định trên tập dữ liệu UIT-VSFC.



Hình 2 Biểu đồ đường cong hàm mất mát

- Tốc độ học: hàm mất mát giảm mạnh từ mức 0,4207 ở epoch 1 xuống dưới 0,1 chỉ sau 10 vòng lặp.
- Tính ổn định: từ epoch 20 trở đi, giá trị loss dao động ở mức cực thấp (< 0,02), cho thấy mô hình đã học được các đặc trưng ngữ nghĩa quan trọng từ cấu trúc đồ thị. Sự biến động nhẹ ở các epoch cuối có thể là dấu hiệu của việc mô hình đang cố gắng tối ưu hóa trên các mẫu nhiễu hoặc các mẫu thuộc lớp thiểu số (Neutral).

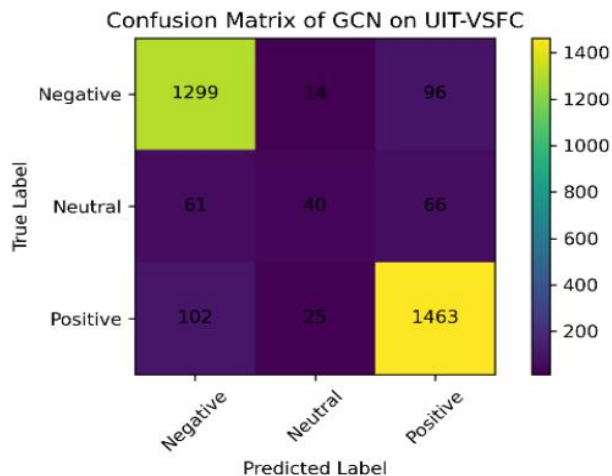
Hiệu năng tổng thể và sự mất cân bằng lớp

Kết quả đạt được với Accuracy là 88,50 % nhưng macro F1-score chỉ đạt 71,34 %. Sự chênh lệch này được giải thích thông qua việc phân tích báo cáo phân loại:

- Lớp tích cực và tiêu cực: mô hình đạt hiệu suất rất cao với F1-score lần lượt là 0,91 và 0,90. Điều này chứng tỏ kiến trúc GCN kết hợp với cửa sổ trượt (sliding window) kích thước $k = 2$ đã nắm bắt hiệu quả các từ khóa mang tính biểu cảm mạnh (ví dụ: “tốt”, “nhiệt tình”, “tệ”, “khó hiểu”).

- Lớp trung tính (neutral): đây là điểm yếu nhất của mô hình với F1-score chỉ đạt 0,33. Tuy nhiên, hiệu suất thấp ở lớp Trung tính (F1-score 0,33) là một hạn chế đáng kể. Điều này có thể do tập dữ liệu UIT-VSFC có sự mất cân bằng lớn (lớp trung tính chỉ chiếm tỷ lệ nhỏ) và các câu trung tính thường có cấu trúc ngữ pháp tương đồng với câu tích cực/tiêu cực nhưng thiếu các tính từ mang tính định hướng mạnh. Phân tích ma trận nhầm lẫn (Confusion Matrix Analysis).

Quan sát Hình 3 (Confusion Matrix), có thể thấy rõ các sai số của mô hình:



Hình 3 Ma trận nhầm lẫn

Độ chính xác trên lớp đa số: trong số 1590 mẫu tích cực, mô hình dự đoán đúng 1463 mẫu (92 %). Tương tự, 1299/1409 mẫu tiêu cực được phân loại chính xác.

Sự nhầm lẫn của lớp trung tính: chỉ có 40/167 mẫu trung tính được dự đoán đúng. Số mẫu trung tính bị nhầm sang tiêu cực (61 mẫu) và tích cực (66 mẫu) chiếm phần lớn. Điều này xảy ra do đặc thù của tiếng Việt trong phản hồi sinh viên; các câu trung tính thường chứa các từ ngữ khách quan hoặc mang tính góp ý nhẹ nhàng, đôi khi sử dụng cấu trúc tương tự câu tích cực/tiêu cực nhưng thiếu các tính từ mạnh, khiến bộ lọc của GCN khó phân biệt khi kích thước tập huấn luyện của lớp này quá nhỏ. Mô hình cho thấy khả năng hội tụ nhanh và ổn định, chứng tỏ cấu trúc GCN đơn giản với cửa sổ trượt phù hợp cho các bài toán phân tích tình cảm trong ngữ cảnh giáo dục với các câu phản hồi ngắn.

5 Kết luận

Nghiên cứu đã tiếp cận dựa trên đồ thị sử dụng mạng GCN để phân tích tình cảm phản hồi của sinh viên Việt Nam. Đóng góp cốt lõi của nghiên cứu là chứng minh tính hiệu quả của cấu trúc đồ thị đồng xuất hiện từ trong việc nắm bắt các quan hệ ngữ nghĩa không gian, giúp mô hình đạt chỉ số F1-score vượt trội hơn so với các mô hình tuần tự truyền thống trên cùng một bộ dữ liệu UIT-VSFC. Nghiên cứu cũng khẳng định rằng GCN là một lựa chọn tối ưu về mặt chi phí tính toán nhưng vẫn đảm bảo độ chính xác cho các bài toán phân loại văn bản tiếng Việt có cấu trúc câu ngắn và đặc thù giáo dục. Về hướng nghiên cứu tương lai, việc tích hợp các đặc trưng cú pháp sâu hơn như cây phụ thuộc (dependency tree) vào cấu trúc cạnh của đồ thị để mô hình hóa chính xác hơn các quan hệ ngữ pháp phức tạp trong tiếng Việt là cần thiết. Ngoài ra, nghiên cứu các kỹ thuật lấy mẫu hoặc các hàm mất mát (loss function) chuyên biệt để cải thiện khả năng nhận diện lớp trung tính – một thách thức lớn do sự mất cân bằng dữ liệu đã được chỉ ra trong kết quả thực nghiệm của bài báo này.

Tài liệu tham khảo

1. Dinh Minh Hoa, Nguyen Tran Thanh Nha, Nguyen Van Xanh, Bui Thi Thanh Tu, Tran Khai Thien. Exploring the potential of graph neural networks for Vietnamese sentiment analysis. *Tạp chí Khoa học Huflit*, 9(2), 11-11.
2. Han, H., Wang, S., Qiao, B., Dang, L., Zou, X., Xue, H., & Wang, Y. (2025). Aspect-Based Sentiment Analysis Through Graph Convolutional Networks and Joint Task Learning. *Information*, 16(3), 201.
3. Mei, A., Mei, J., Hao, T., Rao, M., Zhang, Z., & Li, D. (2025). MSA-GCN: A Graph Convolutional Network Approach for Internal Feature Fusion in Multimodal Sentiment Analysis. In *Proceedings of the 2nd Guangdong-Hong Kong-Macao Greater Bay Area International Conference on Digital Economy and Artificial Intelligence* (pp. 273-279).
4. Zaki, M., Youssef, B., El-gamal, S., & Abd-Elrahem, M. (2026). LISA: language independent sentiment analysis using graph neural networks. *Complex & Intelligent Systems*, 12(1), 45.
5. Huang, X., Peng, H., Sun, S., Hao, Z., Lin, H., & Wang, S. (2025). Multi-view attention syntactic enhanced graph convolutional network for aspect-based sentiment analysis. In *International Conference on Database Systems for Advanced Applications* (pp. 289-304). Singapore: Springer Nature Singapore.
6. Dang Van Thin, Duong Ngoc Hao, Ngan Luu-Thuy Nguyen. (2023). A systematic literature review on vietnamese aspect-based sentiment analysis. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 22(8), 1-28.
7. Huu-Thanh Duong, Tram-Anh Nguyen-Thi, Vinh Truong Hoang. (2022). Vietnamese sentiment analysis under limited training data based on deep neural networks. *Complexity*, 2022(1), 3188449.
8. Lac Si Le, Dang Van Thin, Ngan Luu-Thuy Nguyen, Son Quoc Trinh. (2020). A multi-filter BiLSTM-CNN architecture for Vietnamese sentiment analysis. In *International Conference on Computational Collective Intelligence* (pp. 752-763). Cham: Springer International Publishing.
9. Dang Van Thin, Duong Ngoc Hao, Ngan Luu-Thuy Nguyen. (2024). A Study of Vietnamese Sentiment Classification with Ensemble Pre-Trained Language Models. *Vietnam Journal of Computer Science*, 11(01), 137-165.
10. Wu, Z., Pan, S., Chen, F., Long, G., Zhang, C., & Yu, P. S. (2020). A comprehensive survey on graph neural networks. *IEEE transactions on neural networks and learning systems*, 32(1), 4-24.
11. Wang, K., Ding, Y., & Han, S. C. (2024). Graph neural networks for text classification: A survey. *Artificial intelligence review*, 57(8), 190.
12. Wang, K., Shen, W., Yang, Y., Quan, X., & Wang, R. (2020). Relational graph attention network for aspect-based sentiment analysis. *arXiv preprint arXiv:2004.12362*.
13. Kiet Van Nguyen, Vu Duc Nguyen, Phu X. V. Nguyen, Tham T.H. Truong, Ngan Luu-Thuy Nguyen. (2018). UIT-VSFC: Vietnamese students' feedback corpus for sentiment analysis. In *2018 10th International Conference on Knowledge and Systems Engineering (KSE)* (pp. 19-24). IEEE.
14. Tôn Nữ Thị Sáu, Đỗ Phước Sang, Phạm Thị Thu Trang. (2021). Phân tích ý kiến theo khía cạnh trên bình luận phản hồi của sinh viên cho Tiếng Việt. *TNU Journal of Science and Technology*, 226(18), 48-55.

Graph Convolutional Networks for Vietnamese Student Feedback Sentiment Analysis

Nguyen Thi Phong Dung

Faculty of Information Technology, Nguyen Tat Thanh University, Ho Chi Minh City, Viet Nam

ntpdung@ntt.edu.vn

Abstract Situational analysis plays a crucial role in understanding students' ideas and improving the quality of educational services. However, situational analysis in Vietnamese still suffers from many linguistic specification formulas and limitations in annotated data modes. In this paper, we proposed a deep learning method based on schema using a Geometric Schema Analysis Network (GCN - Graph Neural Networks) to classify the sentiment responses of Vietnamese students. Each input sentence was modeled like a map where words were represented as nodes and edges were constructed based on fixed window sliders. Embedding learning from beginning to end in the GCN was difficult, allowing for efficient language expression learning. Experiments were conducted on the UIT Vietnam Student Response Corpus (UIT-VSFC). The proposed model achieved an accuracy of 88.5% and F1 of 71.3%. Experimental results showed that neural network models based on schemas were a promising direction for analyzing Vietnamese sentiment, especially in the work of extracting educational feedback.

Keywords Sentiment analysis; Graph convolutional neural networks; Vietnamese language; Student feedback; Deep learning.