

Tiểu ban
CÔNG NGHỆ THÔNG TIN

ỨNG DỤNG THUẬT TOÁN HỌC MÁY ĐỂ XÂY DỰNG CƠ SỞ TRI THỨC

Nguyễn Văn Thắm¹, Nguyễn Đỗ Kiều Loan²

¹Trường Đại học Thủy lợi, email: thamnv@tlu.edu.vn

²Học viện Tài chính

1. GIỚI THIỆU CHUNG

Trong những năm gần đây, công nghệ hợp nhất thông tin đa cảm biến ngày càng đóng vai trò quan trọng trong các lĩnh vực như quân sự, y tế và tài chính. Lý thuyết Dempster-Shafer là một phần mở rộng của lý thuyết xác suất Bayes, được sử dụng để xử lý sự không chắc chắn trong trường hợp các xác suất tiên nghiệm chưa được biết, điều mà phương pháp Bayes chủ quan không thể giải quyết hiệu quả, được ứng dụng rộng rãi trong phân tích rủi ro, nhận dạng mẫu, chẩn đoán lỗi và phân lớp. Hiệu quả của luật Dempster phụ thuộc vào cách xây dựng hàm gán xác suất cơ bản (BPA- Basic Probability Assignment) hợp lý. Tuy nhiên, việc tìm giải pháp tổng quát cho xác định BPA vẫn là vấn đề mở, với phần lớn nghiên cứu dựa trên lý thuyết phân phối xác suất hoặc lý thuyết mờ.

Bài báo nghiên cứu lý thuyết Dempster-Shafer và đề xuất mô hình xây dựng cơ sở tri thức dưới dạng các hàm BPA, với thuật toán AdaBoost đóng vai trò nền tảng. Mô hình được triển khai trên một số bộ dữ liệu và được đánh giá thông qua việc so sánh một số độ đo giữa AdaBoost và các phương pháp xây dựng BPA khác.

2. CƠ SỞ LÝ THUYẾT

Đặt $\mathcal{F} = \{E_1, \dots, E_n\}$ là một khung phân biệt gồm một tập hữu hạn không rỗng chứa n biến cố loại trừ lẫn nhau. Đặt $h = 2^n$. Tập lũy thừa (Power Set) của $\mathcal{F}$ là một tập gồm h phần tử $\mathcal{P}(\mathcal{F}) = \{\emptyset, E_1, \dots, E_n, E_1E_2, \dots, E_1\dots E_n\}$, trong đó $E_1\dots E_k, \forall k = \overline{1, n}$ được gọi là thế giới có thể (tập tiêu điểm).

Định nghĩa 1. [1,2] Hàm $m: \mathcal{P}(\mathcal{F}) \rightarrow \mathbb{R}_{[0,1]}$ được gọi là một BPA nếu thỏa mãn các tính chất:

$$(i) \quad m(\emptyset) = 0 \quad (ii) \quad \sum_{\Theta \in \mathcal{P}(\mathcal{F})} m(\Theta) = 1$$

Nếu $A \subset \mathcal{P}(\mathcal{F})$ và $A \neq \emptyset$ thì $m(A)$ biểu diễn cho khả năng A ủng hộ tuyên bố. Nếu $m(A) > 0$ thì tập con A được gọi là phần tử tiêu điểm và $m(A)$ được gọi là giá trị hàm khối lượng của A .

Định nghĩa 2. [1,2] Đặt m_1 và m_2 là hai BPA, $\Delta = \sum_{\Theta_i \cap \Theta_j = \emptyset} m_1(\Theta_i)m_2(\Theta_j)$ và

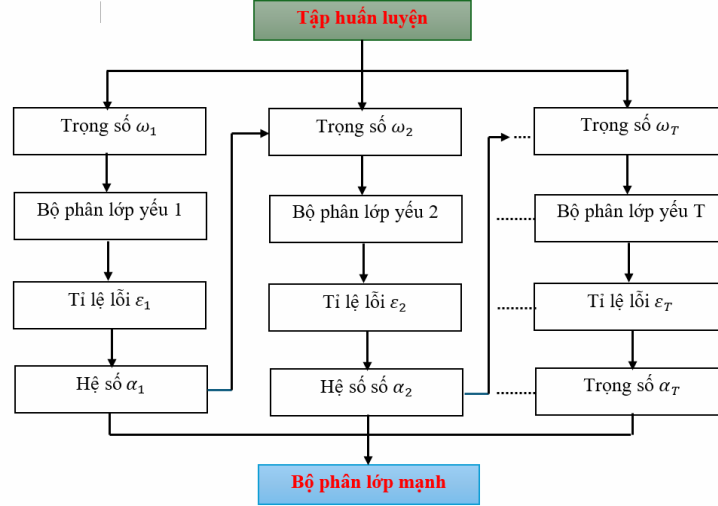
$\Omega = \sum_{\Theta_i \cap \Theta_j \neq \emptyset} m_1(\Theta_i)m_2(\Theta_j) \quad \forall i, j = \overline{1, h}$, luật Dempster (DSR) được định nghĩa như sau:

$$m(\Theta) = \begin{cases} 0 & \text{Nếu } \Theta = \emptyset \\ \frac{\Delta}{1 - \Omega} & \text{Ngược lại} \end{cases} \quad (1)$$

3. MÔ HÌNH ĐỀ XUẤT

3.1. Thuật toán Adaboost

Adaboost [3,4] là một thuật toán Boosting tiêu biểu, hoạt động bằng cách tập trung vào các mẫu bị phân loại sai ở mỗi vòng lặp để dần cải thiện độ chính xác và tạo thành mô hình mạnh. Thuật toán có ưu điểm về cấu trúc đơn giản, hiệu quả cao và không đòi hỏi chọn lọc đặc trưng đầu vào. Adaboost thường kết hợp nhiều bộ phân lớp yếu, phổ biến nhất là cây quyết định nông, để tạo thành một bộ phân lớp mạnh. Hình 1 minh họa trực quan quy trình chính của thuật toán này.



Hình 1. Sơ đồ huấn luyện lớp phân lớp mạnh trong thuật toán Adaboost

3.2. Xác định hàm BPA cho mệnh đề đơn

Dựa trên thuật toán Adaboost [3,4], các mẫu của bất kỳ hai thuộc tính nào trong tập huấn luyện được chọn làm tập huấn luyện mới. Giả sử các mẫu chứa N_c lớp ($N_c > 0$) thì có thể thu được $C_{N_c}^2$ phân lớp mạnh, trong đó phân lớp mạnh thứ t ($t \in [1, N_c]$) được tạo thành T_t phân lớp yếu. Giả sử A và B lần lượt biểu diễn cho lớp thứ nhất và lớp thứ hai của bộ dữ liệu. Nếu bộ phân lớp mạnh thứ t chấp nhận mẫu kiểm thử x thì kết quả chấp nhận mẫu thuộc lớp A là:

$$m_T(A) = \frac{\sum_{k=1, h_{t,k} > 0}^{T_t} \alpha_{t,k}}{\sum_{k=1}^{T_t} \alpha_{t,k}} \quad (2)$$

Trong đó $h_{t,k}$ là bộ phân lớp yếu của bộ phân lớp mạnh thứ t và $\alpha_{t,k}$ là trọng số của $h_{t,k}$. Do mẫu x chỉ có thể là thuộc lớp A hoặc lớp B nên kết quả chấp nhận mẫu thuộc lớp B là:

$$m_T(B) = 1 - m_T(A) \quad (3)$$

Khi các kết quả chấp nhận của tất cả các bộ phân lớp mạnh trên hai thuộc tính hiện tại được ghi lại, BPA của mẫu thử thuộc lớp A được tính như sau:

$$m(A) = \frac{1}{C_{N_c}^2} \sum_{t=1}^{C_{N_c}^2} m_t(A) \quad (4)$$

Các kết quả chấp nhận ở trên được xác định dựa trên hai thuộc tính bất kỳ trong các mẫu đã chọn. Nếu có n ($n > 1$) thuộc tính trong các mẫu thì có thể xác định C_n^2 giá trị BPA cho các thuộc tính đó. Tương tự, các giá trị BPA của các lớp khác trong các mẫu thử cũng có thể được xác định bằng cách sử dụng các tính công thức $m_T(A)$, $m_T(B)$, $m(A)$.

3.3. Xác định hàm BPA cho mệnh đề phức

Hàm BPA của mệnh đề đơn được chuyển thành mệnh đề phức thông qua tỷ lệ diện tích giữa các vùng tương ứng. Với mỗi $\Theta \in \mathcal{P}(\mathcal{F})$ và $\Theta \neq \emptyset$ mà có số phần tử là l ($1 \leq l \leq N_c$) thì được ký hiệu bằng $X(l)$. Khi $l=1$, $\text{area}(X(l))$ biểu diễn diện tích của vùng hình chữ nhật tương ứng với các mẫu của mệnh đề đơn. Khi $l > 1$, $\text{area}(X(l))$ biểu diễn diện tích của vùng giao nhau của hình chữ nhật tương ứng với các mẫu thuộc về ít nhất một trong các mệnh đề đơn khác nhau. Tỷ lệ diện tích của $X(l-1)$ và $X(l)$ được xác định như sau:

$$S(X(l), X(l-1)) = \frac{\text{area}(X(l))}{\text{area}(X(l-1))} \quad (5)$$

Trong quá trình phân bổ lại BPA, hàm BPA của mệnh đề phức được tính một cách đệ quy, tức $m(X(l))$ được tính dựa trên $m(X(l-1))$ với $2 \leq l \leq N_c$. Công thức đệ quy:

$$\begin{cases} m_{AR}(X(l)) = \sum_{\substack{X(l-1) \subset X(l) \\ X(l-1), X(l) \in 2^\Theta}} S(X(l), X(l-1)) m_{AR}(X(l-1)) \\ m_{AR}(X(l-1)) = 1 - S(X(l), X(l-1)) m_{AR}(X(l-1)) \end{cases} \quad 2 \leq l \leq N_c \quad (6)$$

với bước khởi tạo:

$$\begin{cases} m_{AR}(X(2)) = \sum_{\substack{X(1) \subset X(2) \\ X(1), X(2) \in 2^\Theta}} S((X_2, X_1)) m(X(1)) \\ m_{AR}(X(1)) = [1 - S((X_2, X_1))] m(X(1)) \end{cases} \quad (7)$$

Các hàm BPA sau khi phân bổ lại phải thỏa mãn: $\sum_{\substack{X(l) \in 2^\Theta \\ 1 \leq l \leq N_c}} m_{AR}(X(l)) = 1$

3.4. Mô hình xác định các hàm BPA dựa trên Adaboost

Đầu vào: D - Tập huấn luyện, N_c - số lớp của bộ dữ liệu, n - số thuộc tính của bộ dữ liệu, T - số vòng lặp, P - Tập kiểm tra, T_s - Số mẫu của tập kiểm tra

Đầu ra: Cơ sở tri thức biểu diễn bằng các hàm BPA.

Tiến trình:

Bước 1 - Huấn luyện bộ phân lớp mạnh:

Các mẫu tương ứng với hai thuộc tính bất kỳ trong tập huấn luyện được trích xuất để tạo thành một tập huấn luyện mới $D = \{(x_{i(1)}, x_{i(2)}, y_1), \dots, (x_{N(1)}, x_{N_s(1)}, y_N)\}$ với $x_{i(1)}$ và $x_{i(2)}$ lần lượt biểu diễn dữ liệu thuộc tính thứ nhất và thứ hai của mẫu thứ i trong tập kiểm tra, và $y_i \in \{1, 2, \dots, N_c\}$ là nhãn lớp của x_i . Thuật toán Adaboost được sử dụng để huấn luyện $C_{N_c}^2$ bộ phân lớp mạnh.

Bước 2 - Xác định BPA bằng cách huấn luyện bộ phân lớp: Các mẫu tương ứng với hai thuộc tính bất kỳ trong tập kiểm tra được trích xuất để tạo thành một tập kiểm tra mới $P = \{(p_{1(1)}, p_{1(2)}, q_1), \dots, (p_{T_s(1)}, p_{T_s(1)}, q_{T_s})\}$ với $p_{i(1)}$ và $p_{i(2)}$ lần lượt biểu diễn dữ liệu thuộc tính thứ nhất và thứ hai của mẫu thứ i trong tập kiểm tra, và $q_i \in \{1, 2, \dots, N_c\}$ là nhãn lớp của p_i . Các hàm BPA của các mẫu kiểm tra được xác định bằng cách sử dụng công thức (2)-(4).

Bước 3 - Phân bố lại BPA: Nếu mẫu $(p_{i(1)}, p_{i(2)})$ nằm trong vùng giao nhau, sử dụng công thức (6) để phân bổ lại các BPA được xác định trong Bước 2.

Bước 4 - Hợp nhất các BPA: Hợp nhất C_n^2 hàm BPA bằng cách sử dụng công thức (1).

4. THỰC NGHIỆM VÀ ĐÁNH GIÁ

4.1. Các thuật toán và các bộ dữ liệu thực nghiệm

Phần này so sánh phương pháp xác định hàm BPA dựa trên AdaBoost với các thuật toán phân lớp truyền thống gồm SVM (Support Vector Machine), DT (Decision Tree), KNN (K-Nearest Neighbors), và NB (Naïve Bayes). Đồng thời, các thuật toán cũng được thực nghiệm trên các bộ dữ liệu khác nhau. Thông tin chi tiết được trình bày tại Bảng 1.

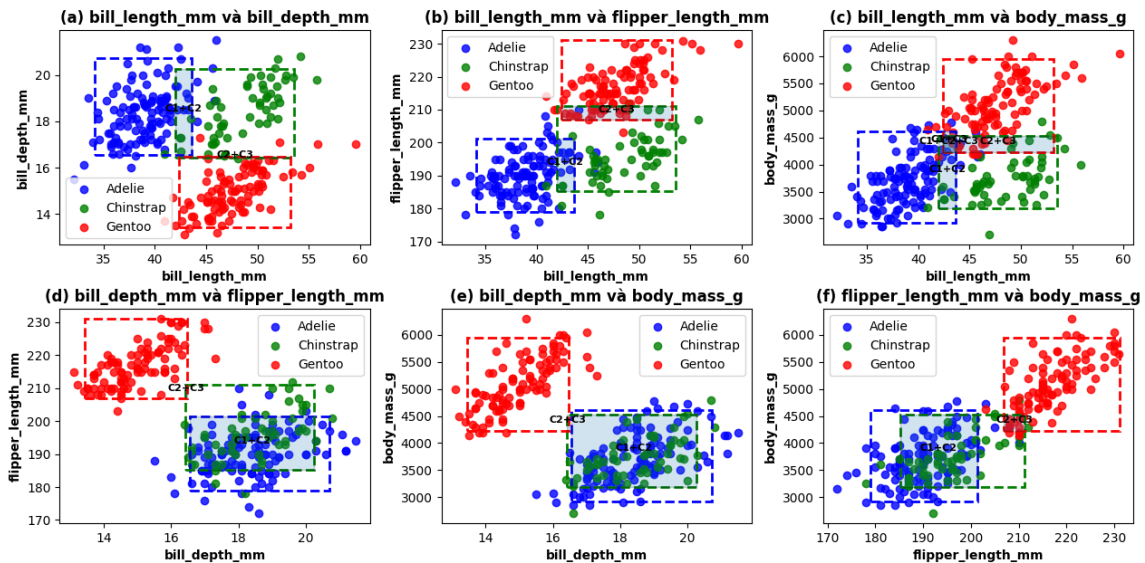
Bảng 1. Các bộ dữ liệu thực nghiệm

Bộ dữ liệu	Đặc trưng	Mẫu	Nhãn	Nguồn
Balance Scale	4	625	3	UCI Machine Learning Repository (https://archive.ics.uci.edu/)
Breast Cancer	30	569	2	sklearn
Diabetes	9	768	2	Kaggle (https://www.kaggle.com/)
Heart	13	1025	2	
Penguins	7	344	3	Seaborn

Bộ dữ liệu Penguins gồm 344 mẫu chim cánh cụt thuộc ba loài: Adelie (A), Chinstrap (C) và Gentoo (G). Các thuộc tính gồm: đảo cư trú (island), chiều dài và chiều sâu mỏ (bill_length_mm-BL, bill_depth_mm-BD), chiều dài vây (flipper_length_mm-FL), khối lượng cơ thể (body_mass_g-BM), và giới tính (sex: Male hoặc Female).

4.2. Kết quả xác định hàm BPA trên bộ dữ liệu Penguins

Bài báo sử dụng bốn thuộc tính sinh học của chim cánh cụt để xây dựng mô hình cơ sở tri thức: BL, BD, FL và BM. Với bốn thuộc tính, có thể tạo ra sáu tổ hợp cặp: (BL, BD), (BL, FL), (BL, BM), (BD, FL), (BD, BM), và (FL, BM), tương ứng với sáu hàm BPA cho mỗi mẫu thử. Phân phối mẫu theo từng cặp thuộc tính được minh họa trong Hình 2. Sáu biểu đồ phân phối dựa trên hai thuộc tính bất kỳ trong bộ dữ liệu Penguins. (a) Phân phối mẫu dựa trên các thuộc tính BL và BD; (b) Phân phối mẫu dựa trên các thuộc tính BL và FL; (b) Phân phối mẫu dựa trên các thuộc tính BL và BM; (d) Phân phối mẫu dựa trên các thuộc tính BD và FL; (e) Phân phối mẫu dựa trên các thuộc tính BD và BM; (f) Phân phối mẫu dựa trên các thuộc tính FL và BM;



Hình 2. Biểu đồ phân lớp của hai thuộc tính bất kỳ trong bộ dữ liệu penguins

Bảng 2 trình bày các giá trị BPA của sáu cặp thuộc tính đối với các tập tiêu điểm A, C, G, AC, AG, CG và ACG, được tính toán dựa trên công thức (5)-(7). Dòng cuối cùng thể hiện kết quả hợp nhất theo nguyên lý Dempster, trong đó giả thuyết A (Adelie) chiếm ưu thế tuyệt đối với $m(A) = 0.8616$, cao nhất trong tất cả các giá trị BPA. Giả thuyết C (Chinstrap), mặc dù thường giữ vị trí thứ hai ở từng cặp thuộc tính, sau khi hợp nhất chỉ còn $m(C) = 0.1384$. Các giả thuyết còn lại như G và các tổ hợp đều có giá trị BPA gần bằng 0. Như vậy, có thể kết luận rằng giả thuyết A nhận được mức độ tin tưởng cao nhất, tiếp theo là C, trong khi G và các tổ hợp không có sự hỗ trợ rõ ràng từ bằng chứng.

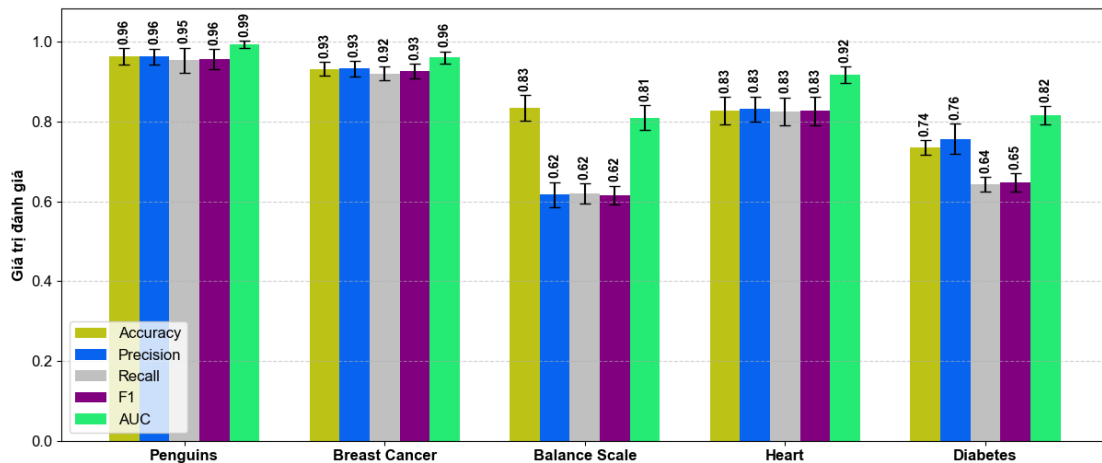
Bảng 2. Các hàm BPA và hàm BPA hợp nhất

Cặp thuộc tính	m(A)	m(C)	m(G)	m(AC)	m(AG)	m(CG)	m(ACG)
BL, BD	0.4447	0.2224	0.0	0.1224	0.0996	0.1109	0.0
BL, FL	0.5634	0.2480	0.0	0.0656	0.0376	0.0853	0.0
BL, BM	0.5210	0.2274	0.0	0.1019	0.0438	0.1059	0.0
BD, FL	0.4994	0.2555	0.0	0.1109	0.0564	0.0779	0.0
BD, BM	0.2636	0.4396	0.0	0.1567	0.0615	0.0786	0.0
FL, BM	0.5672	0.2991	0.0	0.0961	0.0033	0.0342	0.0
BPA hợp nhất	0.8616	0.1384	0.0	0.0000	0.0000	0.0000	0.0

4.3. Phân tích và đánh giá

4.3.1. Kết quả thực nghiệm của mô hình đề xuất

Chương trình thực nghiệm sử dụng `train_test_split` để chia dữ liệu thành 80% để huấn luyện và 20% để kiểm tra, trong đó bộ dữ liệu test luôn tách biệt chỉ dùng cho đánh giá với tham số `random seed = 42` nhằm đảm bảo khả năng tái lập. Mô hình được thiết lập với 10 bộ phân loại yếu, và hiệu năng được đánh giá thông qua `cross-validation 5-fold` sử dụng `StratifiedKFold`, phù hợp cho cả trường hợp dữ liệu cân bằng và mất cân bằng.



Hình 3. Các độ đo và độ lệch chuẩn của mô hình Adaboost trên các bộ dữ liệu

Hình 3 minh họa rõ sự khác biệt về hiệu năng của mô hình trên các tập dữ liệu. Với bộ dữ liệu Penguins và Breast Cancer, mô hình đạt mức rất cao và ổn định với tất cả các độ đo đều vượt 0.92, đồng thời độ lệch chuẩn nhỏ, cho thấy tính tin cậy và khả năng khái quát tốt. Đối với bộ dữ liệu Heart, hiệu năng duy trì khá cân bằng, dao động trong khoảng 0.82-0.92, cùng độ lệch chuẩn ở mức chấp nhận được, phản ánh sự ổn định ở mức khá. Ngược lại, trên bộ dữ liệu Diabetes, hiệu năng chỉ ở mức trung bình, trong đó độ đo Recall và F1 thấp nhất (0.64–0.65) và độ lệch chuẩn cao hơn, cho thấy mô hình còn hạn chế khi xử lý dữ liệu mất cân bằng hoặc phức tạp. Đặc biệt, mô hình chưa hiệu quả với bộ dữ liệu Balance Scale, khi các độ đo Precision, Recall và F1 chỉ khoảng 0.61 và độ lệch chuẩn lớn nhất so với các bộ dữ liệu khác. Nhìn chung, kết quả thực nghiệm chứng minh phương pháp đề xuất phù hợp và ổn định trên đa số tập dữ liệu.

4.3.2. Kết quả thực nghiệm khi sử dụng các thuật toán

Dựa trên Hình 3 và Bảng 3, có thể thấy mô hình Adaboost luôn thể hiện hiệu quả vượt trội và ổn định hơn so với các mô hình SVM, NB, KNN và DT. Thuật toán này đạt độ đo Accuracy cao trên hầu hết các tập dữ liệu, chẳng hạn ở Penguins đạt 0.96 ± 0.02 , vượt rõ rệt so với SVM (0.85 ± 0.02). Độ đo Precision và Recall của Adaboost cũng duy trì ở mức cao và cân bằng, tiêu biểu như Breast Cancer có Recall 0.92 ± 0.02 , nhỉnh hơn KNN (0.89 ± 0.01) và NB (0.93 ± 0.01) nhưng ổn định hơn. Do đó, độ đo F1 của Adaboost thường đứng đầu, điển hình là 0.83 ± 0.03 ở Heart, cao hơn SVM (0.80 ± 0.02) và DT (0.81 ± 0.03). Đặc biệt, AUC của Adaboost đạt 0.99 ± 0.01 trên Penguins, chứng minh khả năng phân biệt lớp vượt trội, đồng thời duy trì mức cao trên các bộ dữ liệu khác như Breast Cancer (0.96 ± 0.02) và Heart (0.92 ± 0.02). Nhìn chung, Adaboost là mô hình cho hiệu năng toàn diện, cân bằng giữa độ chính xác vượt trội so ngay cả trong trường hợp bộ dữ liệu có số mẫu đào tạo ít.

Bảng 3. Các độ đo và độ lệch chuẩn của các thuật toán trên các bộ dữ liệu

Thuật toán	Độ đo	Bộ dữ liệu				
		Penguins	Breast Cancer	Balance Scale	Heart	Diabetes
SVM	Accuracy	0.85 ± 0.02	0.91 ± 0.03	0.81 ± 0.01	0.80 ± 0.02	0.66 ± 0.02
	Precision	0.91 ± 0.01	0.92 ± 0.04	0.54 ± 0.01	0.80 ± 0.02	0.64 ± 0.12
	Recall	0.77 ± 0.04	0.89 ± 0.03	0.58 ± 0.01	0.80 ± 0.02	0.53 ± 0.03
	F1	0.78 ± 0.05	0.90 ± 0.04	0.56 ± 0.01	0.80 ± 0.02	0.48 ± 0.04
	AUC	0.96 ± 0.02	0.94 ± 0.03	0.81 ± 0.03	0.89 ± 0.01	0.77 ± 0.04
NB	Accuracy	0.96 ± 0.02	0.93 ± 0.01	0.83 ± 0.03	0.82 ± 0.03	0.74 ± 0.02
	Precision	0.96 ± 0.02	0.94 ± 0.02	0.61 ± 0.03	0.83 ± 0.03	0.76 ± 0.04
	Recall	0.95 ± 0.03	0.93 ± 0.01	0.62 ± 0.02	0.82 ± 0.04	0.64 ± 0.02
	F1	0.95 ± 0.02	0.93 ± 0.01	0.61 ± 0.02	0.82 ± 0.04	0.65 ± 0.02
	AUC	0.99 ± 0.01	0.96 ± 0.02	0.81 ± 0.03	0.92 ± 0.02	0.82 ± 0.02
KNN	Accuracy	0.93 ± 0.02	0.90 ± 0.02	0.81 ± 0.03	0.86 ± 0.01	0.73 ± 0.03
	Precision	0.93 ± 0.03	0.90 ± 0.02	0.54 ± 0.02	0.86 ± 0.01	0.72 ± 0.05
	Recall	0.91 ± 0.04	0.89 ± 0.01	0.58 ± 0.02	0.86 ± 0.01	0.65 ± 0.03
	F1	0.92 ± 0.03	0.90 ± 0.02	0.56 ± 0.02	0.86 ± 0.01	0.65 ± 0.04
	AUC	0.98 ± 0.01	0.93 ± 0.02	0.77 ± 0.02	0.93 ± 0.01	0.79 ± 0.03
DT	Accuracy	0.96 ± 0.03	0.78 ± 0.03	0.83 ± 0.01	0.81 ± 0.03	0.69 ± 0.03
	Precision	0.96 ± 0.03	0.78 ± 0.03	0.56 ± 0.01	0.81 ± 0.03	0.65 ± 0.03
	Recall	0.95 ± 0.03	0.80 ± 0.03	0.60 ± 0.01	0.81 ± 0.03	0.65 ± 0.03
	F1	0.95 ± 0.03	0.78 ± 0.03	0.58 ± 0.01	0.81 ± 0.03	0.65 ± 0.03
	AUC	0.98 ± 0.02	0.85 ± 0.03	0.77 ± 0.02	0.92 ± 0.02	0.71 ± 0.03

5. KẾT LUẬN

Bài báo nghiên cứu một số thuật toán phân lớp dữ liệu, trong đó tập trung vào việc áp dụng thuật toán Adaboost để xây dựng mô hình tìm các cơ sở tri thức được biểu diễn bằng các BPA từ các bộ dữ liệu. Mô hình gồm bốn bước cơ bản: (1) huấn luyện bộ phân lớp mạnh, (2) xác định BPA bằng cách huấn luyện bộ phân lớp, (3) phân bố lại BPA, (4) hợp nhất các BPA. Chúng tôi đã triển khai thực nghiệm mô hình trên một số bộ dữ liệu với các tỷ lệ huấn luyện khác nhau. Kết quả cho thấy khi tăng tỷ lệ dữ liệu huấn luyện, độ chính xác của mô hình có xu hướng được cải thiện rõ rệt. Đồng thời, khi so sánh với các mô hình phân lớp khác, phương pháp xây dựng cơ sở tri thức dựa trên Adaboost thường cho kết quả chính xác cao hơn.

6. TÀI LIỆU THAM KHẢO

- [1] Gan D, Yang B, Tang Y. 2020. An Extended Base Belief Function in Dempster-Shafer Evidence Theory and Its Application in Conflict Data Fusion. *Mathematics*.; 8(12):2137.
- [2] Fu, W., Yu, S., Wang, X. 2021. A novel method to determine basic probability assignment based on Adaboost and its application in classification. *Entropy*, 23(7), 812.
- [3] He, Y., Xiao, F. 2022. A new base function in basic probability assignment for conflict management. *Appl Intell* 52, 4473-4487.
- [4] Tang, D.; Tang, L.; Dai, R.; Chen, J.W.; Li, X.; Rodrigues, J.J.P.C. 2020. MF-Adaboost: LDoS attack detection based on multi-features and improved adaboost. *Future Gener. Comput. Syst.*

GLOBAL WELL-POSEDNESS AND EXPOTENTIAL STABILITY OF BOUNDED MILD SOLUTIONS TO THE HEAT EQUATION

Nguyen Thi Van

Thuyloi Universirty, email: van@tlu.edu.vn

1. INTRODUCTION

In 1959, Serrin [6] established a fundamental result demonstrating that the exponential stability of strong solutions to Navier-Stokes equations implies the existence of time-periodic solutions to Navier-Stokes equations in bounded domains. In the late 20th century, Kozono and Nakao [2] introduced a new notion of mild solutions and showed the existence of such a solution to Navier-Stokes equations on the whole time-axis $\mathbb{R}$. More recently, Huy *et.al.* [4,5] investigated the exponential stability of mild solutions to Navier-Stokes equations on non-compact Riemannian manifolds with negative curvature. The result has been extended further to Boussinesq system by Ngoc *et.al.* [3]. Their approach is based on $L^p - L^q$ -dispersive and smoothing estimates of Stokes semigroups, combined with the Massera-type principle. Besides, the analytic semigroup $T(t)$ associated with the heat equation, as discussed in [7], satisfies a similar exponential-type estimate in [4]. This motivates the current study on the global well-posedness and stability properties of bounded mild solutions to the heat equation.

2. RESEARCH METHODOLOGY

Our approach employs dispersive and smoothing estimates of the semigroup generated by the linearized heat equation, combined with the fixed point argument and Cone Inequality Theorem. We first establish existence and uniqueness of bounded mild solutions to the linear equation. Then, via the fixed point argument, we extend the result to the semilinear case. Exponential stability is obtained using Cone Inequality Theorem.

3. RESEARCH RESULTS

a. Well-posedness: In this section, we establish the well-posedness of the following linear equation

$$\begin{cases} \frac{\partial u}{\partial t}(t, x) = \frac{\partial^2 u}{\partial x^2}(t, x) + \gamma G(t)v(t, x), & t \in \mathbb{R}_+, x \in [0, \pi], \\ u(t, 0) = u(t, \pi) = 0, & t \in \mathbb{R}_+, \\ u(0, x) = u_0(x), & x \in [0, \pi]. \end{cases} \quad (3.1)$$

where $G(t) = \sin t + \sin(\sqrt{2}t) + e^{-|t|}$, $t \in \mathbb{R}$, $\gamma > 0$ is a constant.

Suppose that $D(A) = \{u \in L^2[0, \pi], u'' \in L^2[0, \pi], u(0) = u(\pi) = 0\}$, $Au(\cdot) = -\Delta u = u''(\cdot)$, $\forall u(\cdot) \in D(A)$. Because $[0, \pi]$ is compact, $-A$ is the infinitesimal generator of an analytic semigroup e^{-tA} on $L^2[0, \pi]$, satisfying $\|e^{-tA}\|_{L^2[0, \pi]} \leq e^{-t}$. Next, we provide the definition of bounded

functions as $C_b(\mathbb{R}_+, L^2[0, \pi]) = \left\{ v : \mathbb{R}_+ \rightarrow L^2[0, \pi] : v \text{ is continuous on } \mathbb{R}_+ \text{ and } \sup_{t \in \mathbb{R}_+} \|v(t)\|_{L^2[0, \pi]} < \infty \right\}$

with the norm $\|v\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} = \sup_{t \in \mathbb{R}_+} \|v(t)\|_{L^2[0, \pi]}$.

From [2], by *mild solution* to (3.1) we mean the function $u(t)$ satisfying the integral equation

$$u(t) = e^{-tA}u_0 + \gamma \int_0^t e^{-(t-\tau)A} G(\tau)v(\tau) d\tau. \quad (3.2)$$

The following lemma gives us the well-posedness (in time) of mild solutions to (3.1) for each bounded function $v(t, x)$.

Lemma 3.1. *Suppose that $v \in C_b(\mathbb{R}_+, L^2[0, \pi])$. Then the problem (3.1) has one and only one mild solution $u \in C_b(\mathbb{R}_+, L^2[0, \pi])$. Moreover,*

$$\|u\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} \leq \|u_0\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} + 3\gamma \|v\|_{C_b(\mathbb{R}_+, L^2[0, \pi])}.$$

Proof

Utilizing the smooth estimate of the semigroup e^{-tA} on $L^2[0, \pi]$, we have

$$\begin{aligned} \|u(t)\|_{L^2[0, \pi]} &\leq \|e^{-tA}u_0\|_{L^2[0, \pi]} + \gamma \int_0^t \|e^{-(t-\tau)A} G(\tau)v(\tau)\|_{L^2[0, \pi]} d\tau \\ &\leq \|u_0\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} + 3\gamma \int_0^t e^{-(t-\tau)} d\tau \|v\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} \\ &\leq \|u_0\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} + 3\gamma \|v\|_{C_b(\mathbb{R}_+, L^2[0, \pi])}. \end{aligned}$$

b. Exponential stability: In this section, we establish the well-posedness and exponential stability of the bounded mild solution to the following semi-linear equation

$$\begin{cases} \frac{\partial u}{\partial t}(t, x) = \frac{\partial^2 u}{\partial x^2}(t, x) + \gamma G(t)u(t, x), & t \in \mathbb{R}_+, x \in [0, \pi] \\ u(t, 0) = u(t, \pi) = 0, & t \in \mathbb{R}_+, \\ u(0, x) = u_0(x), & x \in [0, \pi]. \end{cases} \quad (3.3)$$

The *mild solution* to (3.3) is the function $u(t)$ satisfying the integral equation

$$u(t) = e^{-tA}u_0 + \gamma \int_0^t e^{-(t-\tau)A} G(\tau)u(\tau) d\tau. \quad (3.4)$$

We now introduce the concept of a cone in a Banach space. A closed subset $\mathfrak{S}$ of a Banach space $\mathbf{W}$ is called a *cone* if and only if it satisfies the following conditions:

- (i) $x_0 \in \mathfrak{S} \Rightarrow \lambda x_0 \in \mathfrak{S}, \lambda \geq 0, \lambda \in \mathbb{R}$,
- (ii) $x_1, x_2 \in \mathfrak{S} \Rightarrow x_1 + x_2 \in \mathfrak{S}$,
- (iii) $\pm x_0 \in \mathfrak{S} \Rightarrow x_0 = 0$.

Then, we define the ordinal ($\leq$) on a cone $\mathfrak{S}$ be given in the Banach space $\mathbf{W}$ as follows: For $x, y \in \mathfrak{S}$, we then write $x \leq y$ if $y - x \in \mathfrak{S}$.

If the cone $\mathfrak{S}$ is *invariant* under a linear operator A , then A preserves the inequality, i.e., $x \leq y$ implies $Ax \leq Ay$. We now present the following Cone Inequality Theorem, adapted from [1, Theorem I.9.3].

Theorem 3.2 (Cone Inequality) *Suppose that $\mathfrak{S}$ is a cone in a Banach space $\mathbf{W}$ such that $\mathfrak{S}$ is invariant under a bounded linear operator $A \in \mathcal{L}(\mathbf{W})$ having spectral radius $r_A < 1$. For a vector $x \in \mathbf{W}$ satisfying*

$$x \leq z + Ax \text{ for some given } z \in \mathbf{W},$$

we have that it also satisfies the estimate $x \leq y$, where $y \in \mathbf{W}$ is a solution of the equation

$$y = z + Ay.$$

The main theorem is stated and proved as follows:

Theorem 3.3. Suppose that $\gamma < \frac{1}{3}$. Then the following assertions hold

a) The problem (3.3) has one and only one mild solution u on a small ball of $C_b(\mathbb{R}_+, L^2[0, \pi])$.

b) The bounded mild solution u of Equation (3.3) is exponentially stable in the sense that for any other mild solution $v \in C_b(\mathbb{R}_+, L^2[0, \pi])$ to Equation (3.3) such that $\|u(0) - v(0)\|_{C_b(\mathbb{R}_+, L^2[0, \pi])}$ is small enough, we have

$$\|u(t) - v(t)\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} \leq Ce^{-\alpha t} \|u(0) - v(0)\|_{C_b(\mathbb{R}_+, L^2[0, \pi])},$$

for all $t \geq 0$ and $\alpha < 1 - 3\gamma$.

Proof

a) Let $B_\delta := \left\{ v \in C_b(\mathbb{R}_+, L^2[0, \pi]) : \|v\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} \leq \delta \right\}$. For $v \in B_\delta$, we define

$$\psi(v)(t) = u(t),$$

where $u(t) := e^{-At}u_0 + \gamma \int_{-\infty}^t e^{-(t-\tau)A} G(\tau)v(\tau)d\tau$ for $t \geq 0$.

From Lemma 3.1, we have $\|u\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} \leq \|u_0\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} + 3\gamma\|v\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} < \delta$ if $\|u_0\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} < \alpha\delta$. Therefore, the transformation ψ acts from B_δ into itself. Moreover, for $v_1, v_2 \in B_\delta$, $\gamma < \frac{1}{3}$, we have $\|\psi(v_1) - \psi(v_2)\|_{C_b(\mathbb{R}_+, L^2[0, \pi])} \leq 3\gamma\|v_1 - v_2\|_{C_b(\mathbb{R}_+, L^2[0, \pi])}$. It implies that $\psi : B_\delta \rightarrow B_\delta$ is a contraction. There exists a unique fixed point u of ψ , and by definition of ψ , this function u is the unique bounded mild solution to Equation (3.3).

b) Let $v \in C_b(\mathbb{R}_+, L^2[0, \pi])$ be any solution of Equation (3.2) corresponding to initial value $v_0(x) = v(0, x)$. We have

$$u(t) - v(t) = e^{-tA}(u_0 - v_0) + \gamma \int_0^t e^{-(t-\tau)A} G(\tau)(u(\tau) - v(\tau))d\tau,$$

for $t \geq 0$. Then

$$\begin{aligned} \|u(t)\|_{L^2[0, \pi]} &\leq \|e^{-tA}(u_0 - v_0)\|_{L^2[0, \pi]} + \gamma \int_0^t \|e^{-(t-\tau)A} G(\tau)(u(\tau) - v(\tau))\|_{L^2[0, \pi]} d\tau \\ &\leq e^{-t} \|u(0) - v(0)\|_{L^2[0, \pi]} + 3\gamma \int_0^t e^{-(t-\tau)} \|u(\tau) - v(\tau)\|_{L^2[0, \pi]} d\tau. \end{aligned}$$

Let $\varphi(t) = u(t) - v(t)$. We obtain that $\sup_{t \geq 0} \varphi(t) < \infty$ and

$$\varphi(t) \leq e^{-t} \|u(0) - v(0)\|_{L^2[0, \pi]} + 3\gamma \int_0^t e^{-(t-\tau)} \varphi(\tau) d\tau \quad (3.5)$$

Cone Inequality Theorem will be used for $\mathbf{W} = L^\infty([0, \infty))$ which is the space of real-valued functions defined and essentially bounded on $[0, \infty)$ (endowed with the sup-norm denoted by $\|\cdot\|_\infty$) with the cone $\mathfrak{S}$ being the set of all (a.e.) nonnegative functions. Consider the linear operator A defined for $f \in \mathbf{W}$,

$$(Af)(t) = 3\gamma \int_0^t e^{-(t-\tau)} f(\tau) d\tau$$

Now,

$$\begin{aligned} \sup_{t \geq 0} |(Af)(t)| &= \sup_{t \geq 0} 3\gamma \int_0^t e^{-(t-\tau)} |f(\tau)| d\tau \\ &\leq 3\gamma \int_0^t e^{-(t-\tau)} d\tau \|f\|_\infty \\ &\leq 3\gamma \|f\|_\infty \end{aligned}$$

It implies that $A \in \mathcal{L}(\mathbf{W})$ and $\|A\| < 1$.

Besides, A leaves the cone $\mathfrak{S}$ invariant. The inequality (3.5) can now be expressed as

$$\varphi \leq y + A\varphi$$

where $y(t) = e^{-t} \|u(0) - v(0)\|_{L^2[0, \pi]}$, $t \geq 0$.

Hence, by Cone Inequality Theorem we obtain that $\varphi \leq \phi$, where ϕ is a solution in $\mathbf{W}$ of the equation $\phi = y + A\phi$ which can be rewritten as

$$\phi(t) = e^{-t} \|u(0) - v(0)\|_{L^2[0, \pi]} + 3\gamma \int_0^t e^{-(t-\tau)} \phi(\tau) d\tau.$$

Setting $\mathcal{G}(t) = e^{\alpha t} \phi(t)$, $t \geq 0$. Then

$$\mathcal{G}(t) = e^{-(1-\alpha)t} \|u(0) - v(0)\|_{L^2[0, \pi]} + 3\gamma \int_0^t e^{-(1-\alpha)(t-\tau)} \mathcal{G}(\tau) d\tau. \quad (3.6)$$

We define the linear operator $\wp$ defined for $\varphi \in \mathbf{W}$ by

$$(\wp \varphi)(t) = 3\gamma \int_0^t e^{-(1-\alpha)(t-\tau)} \varphi(\tau) d\tau$$

Equation (3.6) can be written as $\mathcal{G} = z + \wp \mathcal{G}$, where $z(t) = e^{-(1-\alpha)t} \|u(0) - v(0)\|_{L^2[0, \pi]}$.

Again, we can estimate

$$\begin{aligned} \sup_{t \geq 0} |(\wp \mathcal{G})(t)| &= \sup_{t \geq 0} 3\gamma \int_0^t e^{-(1-\alpha)(t-\tau)} |\mathcal{G}(\tau)| d\tau \\ &\leq \frac{3\gamma}{1-\alpha} \|\mathcal{G}\|_\infty. \end{aligned}$$

It yields $\wp \in \mathcal{L}(\mathbf{W})$ and if $\alpha < 1 - 3\gamma$ then $\|\wp\| \leq \frac{3\gamma}{1-\alpha} < 1$. Therefore, the equation $\mathcal{G} = z + \wp \mathcal{G}$ is uniquely solvable in $\mathbf{W}$ and its solution is $\mathcal{G} = (I - \wp)^{-1} z$. Hence,

$$\begin{aligned}\|\mathcal{G}\|_{C_b(\mathbb{R}_+; L^2[0, \pi])} &= \|(I - \wp)^{-1} z\|_{C_b(\mathbb{R}_+; L^2[0, \pi])} \leq \|(I - \wp)^{-1}\| \|z\|_{C_b(\mathbb{R}_+; L^2[0, \pi])} \\ &\leq \frac{1}{1 - \|\wp\|} \|u(0) - v(0)\|_{L^2[0, \pi]} \leq \frac{1 - \alpha}{1 - \alpha - 3\gamma} \|u(0) - v(0)\|_{L^2[0, \pi]} \\ &\leq C \|u(0) - v(0)\|_{L^2[0, \pi]}.\end{aligned}$$

Consequently, $\mathcal{G}(t) \leq C \|u(0) - v(0)\|_{L^2[0, \pi]}$. Because of $\mathcal{G}(t) = e^{\alpha t} \phi(t)$ and $\varphi \leq \phi$, we obtain that

$$\|u(t) - v(t)\|_{C_b(\mathbb{R}_+; L^2[0, \pi])} \leq C e^{-\alpha t} \|u(0) - v(0)\|_{C_b(\mathbb{R}_+; L^2[0, \pi])}, \text{ for } t \geq 0.$$

4. CONCLUSION

In the present paper, we prove the existence, uniqueness and exponential stability of bounded mild solutions to the heat equation by Cone Inequality Theorem. The exponential stability of these solutions is also investigated by using Gronwall's inequality. The result of exponential stability can be extended to prove the existence of time-periodic solutions to such the heat equation by Serrin-type theorem.

5. REFERENCES

- [1] Ju. L. Daleckii, M. G. Krein (1974), Stability of Solutions of Differential Equations in Banach Spaces. Transl. Amer. Math. Soc. Providence RI.
- [2] Kozono. H., Nakao. M (1996), Periodic solution of the Navier-Stokes equations in unbounded domains, Tôhoku Math. J. 48, 33–50.
- [3] T.T. Ngoc, H.Trung, N.T.Van, PT.Xuan (2025), On periodic solution for the Boussinesq system on real hyperbolic manifolds, Dynamical Systems, 40:1, 102-127, <https://doi.org/10.1080/00036811.2025.2459605>.
- [4] T.H. Nguyen, T.X. Pham, T.N.H.Vu, T.M.Vu (2021), Periodic Solutions of Navier-Stokes Equations on Non-compact Einstein Manifolds with Negative Ricci Curvature Tensor, Analysis and Mathematical Physics, Vol. 11, pp. 60.
- [5] N.T.Huy, V. T.N. Ha, N.T. Van (2022), Stability and periodicity of solutions to Navier-Stokes equations on non-compact Riemannian manifolds with negative curvature, Analysis and Mathematical Physics, Vol. 12, No. 4, pp. 89.
- [6] Serrin, J. (1959) A note on the existence of periodic solutions of the Navier-Stokes equations. Arch. Ration. Mech. Anal. 3, 120–122.
- [7] Toka Diagana (2008), Weighted pseudo-almost periodic solutions to some differential equations, Nonlinear Analysis: Theory, Methods and Applications, Vol. 68, No. 8, pp 2250-2260, <https://doi.org/10.1016/j.na.2007.01.054>.

NGHIỆM HẦU TUẦN HOÀN TIỆM CẬN CỦA PHƯƠNG TRÌNH VI PHÂN TUYẾN TÍNH KHÔNG Ô-TÔ-NÔM TRONG KHÔNG GIAN CÁC HÀM BỊ CHẶN

Nguyễn Ngọc Huy

Trường Đại học Thủy lợi, email: huynn@thu.edu.vn

1. GIỚI THIỆU CHUNG

Cho $X = (X, \|\cdot\|)$ là một không gian Banach.

Xét phương trình vi phân tuyến tính không ô-tô-nôm

$$\frac{dv}{dt} = A(t)v(t) + Bf(t), \quad t \in \mathbb{R}, v \in X \quad (1)$$

trong đó họ toán tử $(A(t))_{t \geq 0}$ là tuần hoàn theo chu kỳ T , f là hàm liên tục thuộc không gian các hàm bị chặn

$C_b(\mathbb{R}, X) := \{v: \mathbb{R} \rightarrow X \mid \sup_{t \in \mathbb{R}} \|v(t)\|_X < \infty\}$ trang bị bởi chuẩn $\|v\|_{C_b(\mathbb{R}, X)} := \sup_{t \in \mathbb{R}} \|v(t)\|_X$. B là

toán tử tuyến tính ánh xạ từ không gian $C_b(\mathbb{R}, X)$ vào X .

Bài toán nghiên cứu đáng điều của nghiệm hầu tuần hoàn, hầu tuần hoàn tiệm cận của các phương trình vi phân có nhiều ý nghĩa và được các tác giả quan tâm trong thời gian gần đây như Nguyen, N.H., Nguyen, T.H. & Vu, T.N.H.[3]; Xuan, P.T., Van, N.T.[5]. Đối với các công trình trước đây, Cheban [1] nghiên cứu nghiệm hầu tuần hoàn tiệm cận của các phương trình vi phân trong không gian Euclid n chiều. Gần đây Nguyen, Nguyen & Vu [3] thu được kết quả về sự tồn tại nghiệm hầu tuần hoàn tiệm cận cho lớp phương trình ứng dụng cụ thể Oseen-Navier-Stokes mô tả luồng chất lỏng chảy qua vật cản xoay và dịch chuyển với vận tốc phụ thuộc thời gian. Bên cạnh đó, Xuan & Van [5] nghiên cứu nghiệm hầu tuần hoàn tiệm cận của phương trình Navier-Stokes trên đa tạp hyperbolic.

Trong bài báo này chúng tôi tổng quát hóa kết quả đạt được trong [3] cho lớp phương trình tổng quát, cụ thể chúng tôi nghiên cứu sự tồn tại nghiệm hầu tuần hoàn tiệm cận của phương trình (1) trong không gian các hàm bị chặn. Bài toán tổng quát hóa các kết quả đạt được về đáng điều nghiệm từ phương trình cụ thể Navier-Stokes cho lớp phương trình tổng quát vẫn là bài toán mở, chúng tôi vẫn đang tiếp tục nghiên cứu các bài toán này để thu được các kết quả mới.

Kết quả chính đạt được trong bài báo này là nếu giả thiết họ toán tử $(A(t))_{t \in \mathbb{R}}$ là tuần hoàn theo chu kỳ T , f là hàm hầu tuần hoàn tiệm cận thì phương trình (1) có nghiệm hầu tuần hoàn tiệm cận. Kết quả thu được sau đó mở rộng cho phương trình nửa tuyến tính. Các kết quả này được nêu trong Định lý 2.1 và Định lý 2.2.

2. NỘI DUNG CHÍNH

Trước khi đưa ra kết quả chính, chúng tôi nhắc lại một số khái niệm và tính chất (trong tài liệu tham khảo [1], [2] và [4]).

Định nghĩa 2.1. Họ toán tử tuyến tính $U = (U(t, \tau))_{t \geq \tau}$ trong không gian Banach X là một họ tiến hóa liên tục mạnh nếu:

- 1) $U(t, t) = Id$ và $U(t, r)U(r, \tau) = U(t, \tau)$ với $t, r, \tau \in \mathbb{R}$.
- 2) Ánh xạ $(t, \tau) \mapsto U(t, \tau)x$ là liên tục với mọi $x \in X$.

Trong bài toán Cauchy

$$\begin{cases} \frac{du(t)}{dt} = A(t)u(t), t \geq \tau, \\ u(\tau) = x_\tau \in X, \end{cases}$$

với $A(t)$ (trong trường hợp tổng quát) là một toán tử tuyến tính không bị chặn trên X , thì $u(t) := U(t, \tau)u(\tau)$ là nghiệm của bài toán trên (xem Pazy [4]).

Sự tồn tại của họ tiến hóa $U = (U(t, \tau))_{t \geq \tau}$ cho ta định nghĩa nghiệm đủ tốt (mild solution) của phương trình (1) là hàm $v: \mathbb{R} \rightarrow X$ thỏa mãn phương trình tích phân

$$v(t) = \int_{-\infty}^t U(t, \tau) Bf(\tau) d\tau, t \in \mathbb{R}. \quad (2)$$

Tiếp theo chúng tôi nêu ra khái niệm về hàm hầu tuần hoàn, hầu tuần hoàn tiệm cận (xem [1] và [2]).

Định nghĩa 2.2. Cho X là không gian Banach. Một hàm $f \in C_b(\mathbb{R}, X)$ được gọi là hàm hầu tuần hoàn nếu với mọi $\epsilon > 0$ tồn tại một số thực $L_\epsilon > 0$ sao cho với mỗi $a \in \mathbb{R}$, ta có thể tìm được $\mathcal{T}_{a, \epsilon} \in [a, a + L_\epsilon]$ sao cho

$$\|f(t + \mathcal{T}_{a, \epsilon}) - f(t)\| < \epsilon \text{ với mọi } t \in \mathbb{R}.$$

Tập tất cả các hàm hầu tuần hoàn, ký hiệu bởi $AP(\mathbb{R}, X)$ là một không gian Banach trang bị bởi chuẩn $\|f\|_{AP(\mathbb{R}, X)} := \sup_{t \in \mathbb{R}} \|f(t)\|_X$.

Để đưa ra khái niệm hàm hầu tuần hoàn tiệm cận, ta cần sử dụng không gian $C_0(\mathbb{R}, X)$ là tập tất cả các hàm liên tục $\varphi: \mathbb{R} \rightarrow X$ thỏa mãn $\lim_{t \rightarrow \infty} \|\varphi(t)\|_X = 0$. Khi đó $C_0(\mathbb{R}, X)$ là một không gian Banach với chuẩn $\|\varphi\|_{C_0(\mathbb{R}, X)} := \sup_{t \in \mathbb{R}} \|\varphi(t)\|_X$.

Định nghĩa 2.3. Cho X là không gian Banach. Một hàm $f \in C_b(\mathbb{R}, X)$ được gọi là hàm hầu tuần hoàn tiệm cận nếu tồn tại các hàm $h \in AP(\mathbb{R}, X)$ và $\varphi \in C_0(\mathbb{R}, X)$ thỏa mãn $f(t) = h(t) + \varphi(t)$ với mọi $t \in \mathbb{R}$.

Không gian các hàm hầu tuần hoàn tiệm cận $AAP(\mathbb{R}, X)$ là một không gian Banach trang bị bởi chuẩn $\|f\|_{AAP(\mathbb{R}, X)} := \|h\|_{AP(\mathbb{R}, X)} + \|\varphi\|_{C_0(\mathbb{R}, X)}$.

Ta sử dụng Giả thiết sau đây được dùng trong phần tiếp theo của bài báo.

Giả thiết 2.1. Giả thiết rằng $A(t)$ là tuần hoàn theo chu kỳ T , nghĩa là $A(t + T) = A(t)$ với $T > 0$ là một hằng số cố định với mọi $t \in \mathbb{R}$.

Nhận xét 2.1. Với Giả thiết 2.1 được thỏa mãn, $A(t)$ là tuần hoàn theo chu kỳ T . Khi đó theo Pazy ([4, chương 5, Định lý 6.1]) ta có $(U(t, \tau))_{t \geq \tau}$ cũng tuần hoàn theo chu kỳ T theo nghĩa

$$U(t + T, \tau + T) = U(t, \tau) \text{ với mọi } t \geq \tau.$$

Ta thu được kết quả chính như sau:

Định lý 2.1. Giả sử rằng Giả thiết 2.1 được thỏa mãn. Giả thiết f là hàm hầu tuần hoàn tiệm cận và nghiệm đủ tốt của phương trình (1) thỏa mãn

$$\|v\|_{C_b(\mathbb{R}, X)} \leq M \|f\|_{C_b(\mathbb{R}, X)} \quad (3)$$

với M là một hằng số dương thì phương trình (1) có nghiệm đủ tốt hầu tuần hoàn tiệm cận trong không gian $C_b(\mathbb{R}, X)$.

Chứng minh. Với

$$\begin{aligned} v(t) &= \int_{-\infty}^t U(t, \tau) Bf(\tau) d\tau \\ &= \int_0^\infty U(t, t-\tau) Bf(t-\tau) d\tau, \quad t \in \mathbb{R} \end{aligned}$$

là nghiệm đủ tốt của phương trình (1) xác định trong (2). Theo giả thiết hàm f là hàm tuần hoàn tiệm cận nên tồn tại các hàm $h \in AP(\mathbb{R}, X)$ và $\varphi \in C_0(\mathbb{R}, X)$ thỏa mãn $f(t) = h(t) + \varphi(t)$, $t \in \mathbb{R}$.

Ta có

$$\begin{aligned} v(t) &= \int_{-\infty}^t U(t, \tau) Bf(\tau) d\tau \\ &= \int_{-\infty}^t U(t, \tau) Bh(\tau) d\tau + \int_{-\infty}^t U(t, \tau) B\varphi(\tau) d\tau. \end{aligned}$$

Đặt

$$\begin{aligned} S_1(h)(t) &= \int_{-\infty}^t U(t, \tau) Bh(\tau) d\tau, \\ S_2(\varphi)(t) &= \int_{-\infty}^t U(t, \tau) B\varphi(\tau) d\tau, \end{aligned}$$

Ta được

$$v(t) = S_1(h)(t) + S_2(\varphi)(t), \quad t \in \mathbb{R}.$$

Ta chứng minh được

$$S_1(h) \in AP(\mathbb{R}, X) \text{ và } S_2(\varphi) \in C_0(\mathbb{R}, X).$$

Thật vậy, để chứng minh $S_1(h) \in AP(\mathbb{R}, X)$. Ta có

$$\begin{aligned} S_1(h)(t) &= \int_{-\infty}^t U(t, \tau) Bh(\tau) d\tau \\ &= \int_0^\infty U(t, t-\tau) Bh(t-\tau) d\tau, \quad t \in \mathbb{R} \end{aligned}$$

Do hàm $h \in AP(\mathbb{R}, X)$ nên với mọi $\epsilon > 0$ tồn tại một số thực $L_\epsilon > 0$ sao cho với mỗi $a \in \mathbb{R}$, ta có thể tìm được $\mathcal{T}_{a, \epsilon} \in [a, a + L_\epsilon]$ sao cho $\|h(t + \mathcal{T}_{a, \epsilon}) - h(t)\| < \frac{\epsilon}{2M}$ với mọi $t \in \mathbb{R}$. Với T là chu kỳ của $(A(t))_{t \in \mathbb{R}}$ trong Giả thiết 2.1

Ta có

$$\begin{aligned} &S_1(h)(t + T) \\ &= \int_0^\infty U(t + T, t + T - \tau) Bh(t + T - \tau) d\tau \\ &= \int_0^\infty U(t, t - \tau) Bh(t + T - \tau) d\tau \end{aligned}$$

do $U(t + T, t + T - \tau) = U(t, t - \tau)$ (theo Giả thiết 2.1).

Tương tự ta có

$$\begin{aligned} &S_1(h)(t + nT) \\ &= \int_0^\infty U(t + nT, t + nT - \tau) Bh(t + nT - \tau) d\tau \\ &= \int_0^\infty U(t, t - \tau) Bh(t + nT - \tau) d\tau, \quad n \in \mathbb{Z}. \end{aligned}$$

Khi đó với $\epsilon > 0$ bất kỳ ở trên, ta đặt $l_\epsilon = L_\epsilon + T$. Với $a \in \mathbb{R}$, ta chọn $n_a \in \mathbb{Z}$ sao cho $n_a T \in [a, a + l_\epsilon]$, thì $|n_a T - \mathcal{T}_{a,\epsilon}| \leq l_\epsilon$. Theo (3) ta có

$$\begin{aligned} & \|S_1(h)(t + n_a T) - S_1(h)(t)\| \\ &= \left\| \int_0^\infty U(t, t - \tau) B[h(t + n_a T - \tau) - h(t - \tau)] d\tau \right\| \\ &= \left\| \int_0^\infty U(t, t - \tau) B[h(t + n_a T - \tau) \right. \\ &\quad \left. - h(t + \mathcal{T}_{a,\epsilon} - \tau) + h(t + \mathcal{T}_{a,\epsilon} - \tau) - h(t - \tau)] d\tau \right\| \\ &\leq M \left(\left\| h(\cdot + n_a T) - h(\cdot + \mathcal{T}_{a,\epsilon}) \right\| \right. \\ &\quad \left. + \left\| h(\cdot + \mathcal{T}_{a,\epsilon}) - h(\cdot) \right\| \right) < \epsilon, \quad t \in \mathbb{R}. \end{aligned}$$

Do đó $S_1(h) \in AP(\mathbb{R}, X)$.

Mặt khác, theo (3) ta có

$$\|S_2(\varphi)(t)\| \leq M \|\varphi(t)\|.$$

Do $\lim_{t \rightarrow \infty} \|\varphi(t)\| = 0$ nên $\lim_{t \rightarrow \infty} \|S_2(\varphi)(t)\| = 0$.

Hay $S_2(\varphi) \in C_0(\mathbb{R}, X)$.

Do đó $v(t)$ là nghiệm đủ tốt hầu tuần hoàn tiệm cận của phương trình (1).

Tiếp theo ta xét phương trình nửa tuyến tính

$$\frac{dv}{dt} = A(t)v(t) + Bg(v)(t), \quad t \in \mathbb{R}, v \in X \quad (4)$$

trong đó họ toán tử $A(t)$ thỏa mãn giả thiết của phương trình tuyến tính (1) và toán tử phi tuyến $g : C_b(\mathbb{R}, X) \rightarrow C_b(\mathbb{R}, X)$ thỏa mãn các điều kiện sau:

- (a) $\|g(0)\|_{C_b(\mathbb{R}, X)} \leq \gamma$ trong đó γ là một hằng số không âm.
- (b) g ánh xạ một hàm hầu tuần hoàn tiệm cận thành một hàm hầu tuần hoàn tiệm cận.
- (c) Tồn tại các hằng số dương ρ và L thỏa mãn

$$\|g(v_1) - g(v_2)\|_{C_b(\mathbb{R}, X)} \leq L \|v_1 - v_2\|_{C_b(\mathbb{R}, X)}$$

với mọi $v_1, v_2 \in C_b(\mathbb{R}, X)$

$$\text{và} \quad \|v_1\|_{C_b(\mathbb{R}, X)}, \|v_2\|_{C_b(\mathbb{R}, X)} \leq \rho. \quad (5)$$

Khi đó nghiệm đủ tốt của phương trình (4) là hàm $v : \mathbb{R} \rightarrow X$ thỏa mãn phương trình tích phân

$$v(t) = \int_{-\infty}^t U(t, \tau) Bg(v)(\tau) d\tau, \quad t \in \mathbb{R}. \quad (6)$$

Ta được kết quả về sự tồn tại và duy nhất nghiệm hầu tuần hoàn tiệm cận của phương trình (4).

Định lý 2.2. Giả thiết các điều kiện của Định lý 2.1 được thỏa mãn và hàm g thỏa mãn các điều kiện trong (5). Khi đó nếu L và γ đủ nhỏ thì phương trình (4) có duy nhất nghiệm đủ tốt hầu tuần hoàn tiệm cận $\hat{v}$ thuộc hình cầu bán kính ρ trong không gian $C_b(\mathbb{R}, X)$.

Chứng minh. Xét hình cầu

$B_{aap,\rho}$ xác định bởi

$B_{aap,\rho} := \{x \in C_b(\mathbb{R}, X) : \|x\|_{C_b(\mathbb{R}, X)} \leq \rho\}$ trong đó x là hàm hầu tuần hoàn tiệm cận. Với $x \in B_{aap,\rho}$ ta định nghĩa ánh xạ Φ như sau

$$\Phi(x) := v$$

với $v \in C_b(\mathbb{R}, X)$ là nghiệm hầu tuần hoàn tiệm cận của phương trình

$$\frac{dv}{dt} = A(t)v(t) + Bg(x)(t), \quad t \in \mathbb{R}, v \in X. \quad (7)$$

Chứng minh định lý gồm 2 bước chính như sau:

+ Bước 1: Ta sẽ chứng minh với L và γ đủ nhỏ thì phép biến đổi $\Phi(x)$ ánh xạ $B_{aap,\rho}$ vào chính nó. Thật vậy, do với $x \in B_{aap,\rho}$, theo tính chất của g thỏa mãn (5) ta có

$$\begin{aligned} \|g(x)\|_{C_b(\mathbb{R}, X)} &\leq \|g(x) - g(0)\|_{C_b(\mathbb{R}, X)} + \|g(0)\|_{C_b(\mathbb{R}, X)} \\ &\leq L\rho + \gamma. \end{aligned}$$

Áp dụng Định lý 2.1 cho phương trình (7) và từ (3) ta được $v(t)$ là nghiệm hầu tuần hoàn tiệm cận của phương trình (7) thỏa mãn

$$\|v\|_{C_b(\mathbb{R}, X)} \leq M \|g(x)\|_{C_b(\mathbb{R}, X)} \leq M(L\rho + \gamma).$$

Do đó với L và γ đủ nhỏ thì phép biến đổi $\Phi(x)$ ánh xạ $B_{aap,\rho}$ vào chính nó.

+ Bước 2: Ta chứng minh với L và γ đủ nhỏ, ánh xạ $\Phi : B_{aap,\rho} \rightarrow B_{aap,\rho}$ là ánh xạ co. Thật vậy, theo công thức (6) ta có nghiệm đủ tốt của phương trình (7) được xác định bởi

$$\begin{aligned} \Phi(x)(t) &= v(t) \\ &= \int_{-\infty}^t U(t, \tau) Bg(x)(\tau) d\tau, \quad t \in \mathbb{R}. \end{aligned} \quad (8)$$

Với $v_1, v_2 \in C_b(\mathbb{R}, X)$, từ (8) và áp dụng lần nữa Định lý 2.1 ta có hàm

$v := \Phi(v_1) - \Phi(v_2)$ là nghiệm đủ tốt hầu tuần hoàn tiệm cận của phương trình

$$v'(t) = A(t)v(t) + B(g(v_1) - g(v_2))(t)$$

và ta được

$$\begin{aligned} \|\Phi(v_1) - \Phi(v_2)\|_{C_b(\mathbb{R}, X)} &\leq M \|g(v_1) - g(v_2)\|_{C_b(\mathbb{R}, X)} \\ &\leq ML \|v_1 - v_2\|_{C_b(\mathbb{R}, X)}. \end{aligned}$$

Do đó với L và γ đủ nhỏ, ánh xạ

$\Phi : B_{aap,\rho} \rightarrow B_{aap,\rho}$ là ánh xạ co nên tồn tại duy nhất điểm bất động $\hat{v}$ của Φ và theo cách xác định của ánh xạ Φ thì $\hat{v}$ là nghiệm hầu tuần hoàn tiệm cận duy nhất của phương trình (4).

3. KẾT LUẬN

Bài báo nghiên cứu sự tồn tại và duy nhất nghiệm hầu tuần hoàn tiệm cận cho phương trình vi phân tuyến tính không ô-tô-nôm trong không gian Banach các hàm bị chặn. Tác giả xây dựng khung lý thuyết dựa trên khái niệm họ tiến hóa tuyến tính, hàm hầu tuần hoàn, hàm hầu tuần hoàn tiệm cận và áp dụng vào phương trình tuyến tính, cũng như phương trình nửa tuyến tính. Các kết quả chính thu được bằng cách sử dụng phương pháp tích phân, tính chất của họ toán tử tuần hoàn và định lý ánh xạ co.

4. TÀI LIỆU THAM KHẢO

- [1] D.N. Cheban, Asymptotically Almost Periodic Solutions of Differential Equations, Hindawi Publishing Corporation, New York, 2009.
- [2] B.M. Levitan, V.V. Zhikov, Almost Periodic Functions and Differential Equations, Cambridge University Press, 1982.
- [3] Nguyen, N.H., Nguyen, T.H. & Vu, T.N.H. Asymptotically Almost Periodic and Almost Automorphic Solutions to the Non-autonomous Oseen-Navier-Stokes Equations. *Acta Math Vietnam* 50, 79–99 (2025). <https://doi.org/10.1007/s40306-024-00557-1>.
- [4] A. Pazy, Semigroup of Linear Operators and Application to Partial Differential Equations, Springer-Verlag, Berlin, 1983.
- [5] Xuan, P.T., Van, N.T. On asymptotically almost periodic solutions to the Navier–Stokes equations in hyperbolic manifolds. *J. Fixed Point Theory Appl.* 25, 71 (2023). <https://doi.org/10.1007/s11784-023-01074-8>.

MỘT SỐ TÍNH CHẤT CỦA TOÁN TỬ HESSIAN THƯỜNG

Nguyễn Hữu Thọ

Trường Đại học Thủy lợi, email: nhtho@tlu.edu.vn

1. GIỚI THIỆU CHUNG

Đa thức đối xứng cơ bản đã được giảng dạy và có nhiều ứng dụng ở bậc phổ thông. Dạng tổng quát của đa thức đối xứng cơ bản chính là các hàm đối xứng cơ bản, các tính chất của các hàm đối xứng cơ bản đã được sử dụng nhiều trong nghiên cứu phương trình đạo hàm riêng phi tuyến, chẳng hạn như các phương trình đạo hàm riêng liên quan đến các bài toán về độ cong (xem [5], [6]), liên quan đến sự tồn tại và tính duy nhất nghiệm của bài toán Dirichlet đối với phương trình kiểu k – Hessian trong miền bị chặn với dữ kiện không nhất thiết trơn (xem [1], [3]).

Nhằm mở rộng các kết quả trong [1] và [2], chúng tôi đang quan tâm tới lớp các bài toán về phương trình Hessian thương. Như một lẽ tự nhiên, chúng tôi cần nghiên cứu tới các toán tử dạng $\frac{\sigma_k}{\sigma_l}(\eta)$ với $1 \leq l < k \leq n$. Trong báo cáo này chúng tôi sẽ tập trung nghiên cứu một số tính chất cơ bản của toán tử Hessian thương, đây là bước rất quan trọng để đạt được những kết quả trọn vẹn trong hướng nghiên cứu tiếp theo.

2. NỘI DUNG BÁO CÁO

Hàm đối xứng cơ bản bậc k của n biến, ký hiệu là σ_k , và được xác định bởi

$$\sigma_k(\lambda) = \sum_{1 \leq i_1 < i_2 < \dots < i_k \leq n} \lambda_{i_1} \lambda_{i_2} \dots \lambda_{i_k}, \quad (1)$$

trong đó $1 \leq k \leq n$ và $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_n) \in \mathbb{R}^n$. Hàm σ_k thường được xét trong nón tương ứng trong $\mathbb{R}^n$ đó là Γ_k , nón đó được xác định như sau:

$$\Gamma_k = \left\{ \lambda \in \mathbb{R}^n \mid \sigma_j(\lambda) > 0, j = 1, 2, \dots, k \right\}. \quad (2)$$

Dễ thấy Γ_k là một nón lồi có đỉnh là gốc tọa độ, $\Gamma_k \subseteq \Gamma_j$ với $k \geq j$, và Γ_n chính là nón dương $\left\{ \lambda \in \mathbb{R}^n \mid \lambda_i > 0, i = 1, 2, \dots, n \right\}$. Đặt $\sigma_{k-1}(\lambda | i) = \frac{\partial \sigma_k}{\partial \lambda_i}$ và $\sigma_{k-2}(\lambda | ij) = \frac{\partial^2 \sigma_k}{\partial \lambda_i \partial \lambda_j}$, sau đây là một số tính chất cơ bản của σ_k (xem trong [5]), các tính chất này sẽ được sử dụng trong quá trình chứng minh một số kết quả về $\frac{\sigma_k}{\sigma_l}(\eta)$.

Định lý 2.1 ([4]) Xét $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_n) \in \mathbb{R}^n$ với $\lambda \in \Gamma_k$ và $1 \leq k \leq n$, khi đó ta có

- (1) $\sigma_{k-1}(\lambda | i) > 0$ với $1 \leq i \leq n$;
- (2) $\sigma_k(\lambda) = \sigma_k(\lambda | i) + \lambda_i \sigma_{k-1}(\lambda | i)$ với $1 \leq i \leq n$;
- (3) Nếu $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$ thì $\sigma_{k-1}(\lambda | 1) \leq \sigma_{k-1}(\lambda | 2) \leq \dots \leq \sigma_{k-1}(\lambda | n)$;
- (4) $\sum_{i=1}^n \sigma_{k-1}(\lambda | i) = (n - k + 1) \sigma_{k-1}(\lambda)$;

(5) Nếu $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$ thì ta có $\sigma_{k-1}(\lambda | k) \geq C_{n,k} \sum_i \sigma_{k-1}(\lambda | i)$, trong đó $C_{n,k}$ là hằng số dương chỉ phụ thuộc vào n và k ;

(6) Nếu $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$ thì $\sigma_{k-1}(\lambda | k) \geq C_{n,k} \sigma_{k-1}(\lambda)$;

(7) Nếu $n \geq k > l \geq 0, n \geq r > s \geq 0, k \geq r, l \geq s$ thì ta có

$$\left[\frac{\sigma_k(\lambda)}{C_n^k} \right]^{\frac{1}{k-l}} \leq \left[\frac{\sigma_r(\lambda)}{C_n^r} \right]^{\frac{1}{r-s}}.$$

(8) Nếu $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_n) \in \Gamma_k$ thì $\left[\frac{\sigma_k(\lambda)}{\sigma_l(\lambda)} \right]^{\frac{1}{k-l}}, (0 \leq l < k \leq n)$ là lõm theo λ . Do đó với $(\varsigma_1, \varsigma_2, \dots, \varsigma_n)$ ta có

$$\sum_{i,j} \frac{\partial^2 \left[\frac{\sigma_k(\lambda)}{\sigma_l(\lambda)} \right]}{\partial \lambda_i \partial \lambda_j} \varsigma_i \varsigma_j \leq \left(1 - \frac{1}{k-l} \right) \frac{\left[\sum_i \frac{\partial \left[\frac{\sigma_k(\lambda)}{\sigma_l(\lambda)} \right]}{\partial \lambda_i} \varsigma_i \right]^2}{\frac{\sigma_k(\lambda)}{\sigma_l(\lambda)}}.$$

Với $z = (z_1, z_2, \dots, z_n) \in \mathbb{C}^n$, ký hiệu $u_{i\bar{j}} = \partial_{i\bar{j}} u = \frac{\partial^2 u}{\partial z_i \partial \bar{z}_j}$, $\eta_{ij} = \tau \delta_{ij} + u_{ij}$ với $\tau \geq 0$. Gọi

$\lambda = \lambda(u_{i\bar{j}}) = (\lambda_1, \lambda_2, \dots, \lambda_n)$ là các giá trị riêng của $\{u_{i\bar{j}}\}$ và $\eta = \lambda(\eta(u_{i\bar{j}})) = (\eta_1, \eta_2, \dots, \eta_n)$ là các giá trị riêng của $\{\eta_{i\bar{j}}\}$; khi đó ta có $\eta_i = \lambda_i + \tau$. Xét phương trình

$$F(u) = \left(\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right)^{\frac{1}{k-l}} \quad (3)$$

trong miền Ω , với Ω là miền C^4 trong $\mathbb{C}^n$. Nếu $\tau = 0$ phương trình (3) trở thành phương trình Hessian thương cổ điển. Khi $\tau = 0, k = n, l = 0$ phương trình (3) trở thành phương trình Monge – Ampere.

Hàm u được gọi là một nghiệm (η, k) -chấp nhận được của phương trình (3) nếu $\lambda = \lambda\{u_{i\bar{j}}\} = (\lambda_1, \lambda_2, \dots, \lambda_n) \in \Gamma_k$ với $z \in \Omega$. Khi đó $\{u_{i\bar{j}}\}$ là ma trận đường chéo tại mỗi điểm $z \in \Omega$.

Để thuận tiện ta đặt $F^{i\bar{j}} = \frac{\partial F}{\partial u_{i\bar{j}}}$.

Giả sử $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$ khi đó $\eta_1 \geq \eta_2 \geq \dots \geq \eta_n$ và $\sigma_l(\eta | 1) \leq \dots \leq \sigma_l(\eta | n)$ với $1 \leq l \leq k-1$. Sau đây chúng ta nhận được một số kết quả quan trọng về toán tử $\frac{\sigma_k}{\sigma_l}(\eta)$.

Định lý 2.2 Nếu $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_n) \in \Gamma_k$ thì $\left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]^{\frac{1}{k-l}}, (0 \leq l < k \leq n)$ là lõm theo λ .

Chứng minh. Ta sẽ chứng minh rằng, với $(\varsigma_1, \varsigma_2, \dots, \varsigma_n)$ thì ta có

$$\sum_{i,j} \frac{\partial^2 \left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]}{\partial \lambda_i \partial \lambda_j} \varsigma_i \varsigma_j \leq \left(1 - \frac{1}{k-l} \right) \frac{\left[\sum_i \frac{\partial \left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]}{\partial \lambda_i} \varsigma_i \right]^2}{\frac{\sigma_k(\eta)}{\sigma_l(\eta)}}.$$

Chúng ta lưu ý rằng: $\lambda \in \Gamma_k$ sẽ suy ra $\sigma_j(\lambda) > 0$ với mọi $1 \leq j \leq k$ do đó $\sigma_j(\eta) > 0$ và các đạo hàm xác định. Theo Định lý 2.1 phần (8), ta có

$$\begin{aligned} \sum_{i,j} \frac{\partial^2 \left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]}{\partial \lambda_i \partial \lambda_j} \varsigma_i \varsigma_j &= \sum_{i,j,a,b} \frac{\partial^2 \left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]}{\partial \eta_a \partial \eta_b} \frac{\partial \eta_a}{\partial \lambda_i} \frac{\partial \eta_b}{\partial \lambda_j} \varsigma_i \varsigma_j \leq \left(1 - \frac{1}{k-l} \right) \frac{\left[\sum_a \frac{\partial \left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]}{\partial \eta_a} \left(\sum_i \frac{\partial \eta_a}{\partial \lambda_i} \varsigma_i \right) \right]^2}{\frac{\sigma_k(\eta)}{\sigma_l(\eta)}} \\ &\leq \left(1 - \frac{1}{k-l} \right) \frac{\left[\sum_{a,i} \frac{\partial \left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]}{\partial \eta_a} \frac{\partial \eta_a}{\partial \lambda_i} \varsigma_i \right]^2}{\frac{\sigma_k(\eta)}{\sigma_l(\eta)}} \leq \left(1 - \frac{1}{k-l} \right) \frac{\left[\sum_{a,i} \frac{\partial \left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]}{\partial \lambda_i} \varsigma_i \right]^2}{\frac{\sigma_k(\eta)}{\sigma_l(\eta)}} \end{aligned}$$

và từ đây chúng ta nhận được tính lõm của $\left[\frac{\sigma_k(\eta)}{\sigma_l(\eta)} \right]^{\frac{1}{k-l}}$, $(0 \leq l < k \leq n)$.

Định lý 2.3 Với $0 \leq l < k \leq n$, r là ma trận Hermit cấp $n \times n$, $\lambda(r)$ là các giá trị riêng của r với $\eta = \tau \sum_{j=1}^n \lambda_j \cdot \mathbf{1} + \lambda \in \Gamma_k$. Khi đó

$$f_1 \geq \frac{\tau}{n\tau + 1} C_{n,k,l} \sum_{i=1}^n f_i,$$

ở đó $0 \leq f_1 \leq \dots \leq f_n$ là các giá trị riêng của $\{F^{ij}(r)\}$, $C_{n,k,l}$ là các hằng số phụ thuộc vào n, k, l .

Chứng minh. Khi $l = 0$ kết quả nhận được theo Bổ đề 9 trong [2]. Do đó ta chỉ chứng minh khi $l \geq 1$. Không mất tính tổng quát giả sử rằng ma trận r là ma trận đường chéo với $\lambda_1(r) \geq \dots \geq \lambda_n(r)$. Từ đó $\eta_1(r) \geq \dots \geq \eta_n(r)$. Bằng cách tính toán trực tiếp ta có

$$\begin{aligned} f_i &= \frac{\partial \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}}(\eta)}{\partial \eta_p} \frac{\partial \eta_p}{\partial \lambda_i} = \frac{1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} + \sum_p \frac{\sigma_{k-1}(\eta | p)(\tau + \delta_{ip})\sigma_l - \sigma_k \sigma_{l-1}(\eta | p)(\tau + \delta_{ip})}{\sigma_l^2(\eta)} \\ &= \frac{1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{(\tau + 1)\sigma_{k-1}(\eta | i)\sigma_l(\eta | i)(1 - \alpha_i) + \tau \sum_{p \neq i} \sigma_{k-1}(\eta | p)\sigma_l(\eta | p)(1 - \alpha_p)}{\sigma_l^2(\eta)} \end{aligned}$$

ở đó α_p xác định bởi

$$\alpha_p = \frac{\sigma_k(\eta | p) \sigma_{l-1}(\eta | p)}{\sigma_{k-1}(\eta | p) \sigma_l(\eta | p)}, p = 1, \dots, n.$$

Chú ý rằng

$$\sigma_{k-1}(\eta | p) \geq C_{n,k} \sum_{i=1}^n \sigma_{k-1}(\eta | i), \sigma_l(\eta | p) \geq C_{n,l} \sum_{i=1}^n \sigma_l(\eta | i), \forall p \geq k. \quad (4)$$

Ta sẽ chia thành hai trường hợp:

Trường hợp 1: $\sigma_k(\eta | p) > 0$. Do $\eta \in \Gamma_k$ nên $\eta | p \in \Gamma_k$, $\eta | p$ là véc tơ $n -$ chiều với thành phần thứ p là bằng 0. Theo bất đẳng thức Newton – MacLaurin với $1 \leq p \leq n$ ta có

$$\alpha_p \leq \frac{C_{n-1}^{l-1} C_{n-1}^k}{C_{n-1}^l C_{n-1}^{k-1}} = \frac{l(n-k)}{k(n-l)},$$

$$\text{và từ đó } 1 - \alpha_p \geq 1 - \frac{l(n-k)}{k(n-l)} = \frac{n(k-l)}{k(n-l)}.$$

Trường hợp 2: $\sigma_k(\eta | p) \leq 0$. Khi đó từ công thức xác định α_p như ở trên, ta thấy $\alpha_p \leq 0$ và $1 - \alpha_p \geq 1 \geq \frac{n(k-l)}{k(n-l)}$. Bây giờ sử dụng đánh giá (4) ta nhận được

$$\begin{aligned} f_1 &\geq \frac{1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \sum_{p \geq k} \frac{\tau \sigma_{k-1}(\eta | p) \sigma_l(\eta | p) (1 - \alpha_p)}{\sigma_l^2(\eta)} \\ &\geq \frac{n(k-l)}{k(n-l)} \frac{\tau(n-k+1)}{k-l} C_{n,k} C_{n,l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\sigma_{k-1}(\eta) \sigma_l(\eta)}{\sigma_l^2(\eta)} \\ &= \tau C_{n,k,l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\sigma_{k-1}(\eta)}{\sigma_l(\eta)} \geq C_{n,k,l} \frac{\tau}{n\tau+1} \sum_{i=1}^n f_i \end{aligned}$$

ở đây bất đẳng cuối suy ra từ

$$\begin{aligned} \sum_{i=1}^n f_i &= \frac{1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \left[(\tau+1) \frac{\sum_{i=1}^n \sigma_{k-1}(\eta | i) \sigma_l(\eta) - \sum_{i=1}^n \sigma_{l-1}(\eta | i) \sigma_k(\eta)}{\sigma_l^2(\eta)} + \right. \\ &\quad \left. + \tau \frac{\sum_{i=1}^n \sum_{p \neq i} \sigma_{k-1}(\eta | i) \sigma_l(\eta) - \sum_{i=1}^n \sum_{p \neq i} \sigma_{l-1}(\eta | i) \sigma_k(\eta)}{\sigma_l^2(\eta)} \right] \\ &= \frac{n\tau+1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\sum_{i=1}^n \sigma_{k-1}(\eta | i) \sigma_l(\eta) - \sigma_k(\eta) \sigma_{l-1}(\eta | i)}{\sigma_l^2(\eta)} \\ &\leq \frac{n\tau+1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{1}{\sigma_l} \sum_{i=1}^n \sigma_{k-1}(\eta | i) \\ &= \frac{(n\tau+1)(n-k+1)}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\sigma_{k-1}(\eta)}{\sigma_l(\eta)}. \end{aligned}$$

Như vậy, định lý được chứng minh.

Chú ý. Theo bất đẳng thức Newton – MacLaurin, nếu $0 \leq l < k$ ta nhận thấy

$$\begin{aligned}
 \sum_{i=1}^n f_i(\lambda) &= \frac{n\tau+1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\sum_{i=1}^n \sigma_{k-1}(\eta|i) \sigma_l(\eta) - \sigma_k(\eta) \sigma_{l-1}(\eta|i)}{\sigma_l^2(\eta)} \\
 &= \frac{n\tau+1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{(n-k+1)\sigma_{k-1}\sigma_l - (n-l+1)\sigma_{l-1}\sigma_k}{\sigma_l^2} \\
 &= \frac{(n\tau+1)(n-k+1)}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\sigma_{k-1}}{\sigma_k} - \frac{(n\tau+1)(n-l+1)}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\sigma_{l-1}}{\sigma_l} \\
 &= (n\tau+1) \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \left[\frac{kC_n^k \sigma_{k-1}}{(k-l)C_n^{k-1} \sigma_k} - \frac{lC_n^l \sigma_{l-1}}{(k-l)C_n^{l-1} \sigma_l} \right] = \frac{n\tau+1}{k-l} \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \left[k \frac{C_n^k \sigma_{k-1}}{C_n^{k-1} \sigma_k} - l \frac{C_n^l \sigma_{l-1}}{C_n^{l-1} \sigma_l} \right] \\
 &\geq (n\tau+1) \left(\frac{\sigma_k}{\sigma_l} \right)^{\frac{1}{k-l}-1} \frac{\frac{\sigma_{k-1}}{C_n^{k-1}}}{\frac{\sigma_k}{C_n^k}} = (n\tau+1) \left(\frac{\frac{\sigma_k}{C_n^k}}{\frac{\sigma_l}{C_n^l}} \right)^{\frac{1}{k-l}-1} \left[\frac{\frac{\sigma_{k-1}}{C_n^{k-1}}}{\frac{\sigma_k}{C_n^k}} \right]^{\frac{1}{k-l}} \\
 &\geq (n\tau+1) \left(\frac{C_n^k}{C_n^l} \right)^{\frac{1}{k-l}-1} := (n\tau+1) C_{n,k,l}
 \end{aligned}$$

khí $\lambda \in \mathbb{R}^n$ với $\eta \in \Gamma_k$.

Từ đó chúng ta nhận được hệ quả quan trọng sau đây về tính elliptic của $\frac{\sigma_k}{\sigma_l}(\eta)$, tính chất này đóng vai trò rất quan trọng khi nghiên cứu về sự tồn tại nghiệm chấp nhận được của phương trình (3). Ta nhớ lại rằng: Toán tử A được gọi là elliptic chặt nếu tồn tại hằng số $c > 0$ sao cho: $a^{ij}(x)\zeta_i\zeta_j \geq c|\zeta|^2$ với mọi x thuộc tập compact M , mọi $\zeta \in \mathbb{R}^n$.

Hệ quả. Nếu $\lambda \in \mathbb{R}^n$ với $\eta \in \Gamma_k$, toán tử $\frac{\sigma_k}{\sigma_l}(\eta)$ là elliptic chặt theo λ .

3. KẾT LUẬN

Báo cáo trình bày về một số tính chất cơ bản của của toán tử k -Hessian thương $\frac{\sigma_k}{\sigma_l}(\eta)$ với $1 \leq l < k \leq n$, các tính chất này đóng vai trò quan trọng để đạt được những kết quả trọn vẹn trong hướng nghiên cứu tiếp theo về phương trình (3).

4. TÀI LIỆU THAM KHẢO

- [1] A. Colesanti, P. Salani. 1999. Hessian equations in non-smooth domain. Nonlinear Anal., Vol. 38, No. 6, Ser.A: Theory Methods, pp. 803-812.
- [2] W. Dong, W.Wei. 2022. The Neumann problem for a type of fully nonlinear complex equations. J. Differential equations, Vol. 306, pp. 525-546.
- [3] F. Jiang, N.S. Trudinger, X.P. Yang. 2015. On the Dirichlet problem for a class of augmented Hessian equations. J. Differ. Equ. Vol. 258, pp. 1548-1576.
- [4] N.M. Ivovichikina. 1991. The Dirichlet problem for the equation of curvature equation of order m. (English traslation). Leningrad Math. J. No. 2, pp. 631-654.
- [5] N.S. Trudinger. 1990. The Dirichlet problem for the prescribed curvature equations. Arch. Rational Mech. Anal, No. 111, pp. 153-179.

MÔ HÌNH PHÂN LỚP DỮ LIỆU KHÔNG CHẮC CHẮN

Nguyễn Văn Thắm, Nguyễn Quỳnh Diệp

Trường Đại học Thủy lợi, email: thamnv@tlu.edu.vn

1. GIỚI THIỆU CHUNG

Dữ liệu không chắc chắn (Uncertain Data) là dữ liệu chứa sai số, thiếu hụt hoặc không rõ ràng, bắt nguồn từ giới hạn kỹ thuật, đo lường hoặc thông tin không đầy đủ [1-4]. Dữ liệu không chắc chắn phổ biến trong mạng cảm biến (IoT, robot, xe tự hành), văn bản nhiễu trên mạng xã hội, dữ liệu doanh nghiệp lỗi thời, dữ liệu y tế mang tính xác suất và mô hình toán học chỉ mang tính xấp xỉ. Khi lưu trữ các dữ liệu này, cần lựa chọn mô hình cơ sở dữ liệu phù hợp để xử lý tính không chắc chắn. Tuy nhiên, các nghiên cứu đã công bố phần lớn sử dụng các tập dữ liệu Iris, Wine, Glass và Ecoli từ UCI Machine Learning Repository để thực nghiệm. Đây đều là những tập dữ liệu chắc chắn, do đó các nghiên cứu thường phải giả định thêm yếu tố không chắc chắn khi thử nghiệm, đồng thời chưa xem xét đến tính hội tụ trong quá trình tính toán lại medoid của từng lớp.

Bài báo tập trung nghiên cứu dữ liệu không chắc chắn và đề xuất một mô hình phân lớp phù hợp cho loại dữ liệu này. Mô hình được triển khai trên tập dữ liệu không chắc chắn và được đánh giá thông qua so sánh chất lượng phân lớp với một số thuật toán phân lớp khác.

2. CƠ SỞ LÝ THUYẾT

2.1. Các khái niệm cơ bản

Để biểu diễn dữ liệu không chắc chắn ở cấp độ thuộc tính, hai mô hình được sử dụng phổ biến là mô hình không chắc chắn đơn biến (Univariate Uncertainty Model) và mô hình không chắc chắn đa biến (Multivariate Uncertainty Model).

Trong mô hình không chắc chắn đa biến, một điểm dữ liệu không chắc chắn m -chiều được biểu diễn bởi một vùng m -chiều cùng với một hàm mật độ xác suất đa biến, trong đó xác suất được phân bố sao cho điểm dữ liệu có thể trùng khớp với bất kỳ điểm nào trong vùng này. Ngược lại, trong mô hình không chắc chắn đơn biến, mỗi thuộc tính của điểm dữ liệu m -chiều được gán với một khoảng giá trị và một hàm mật độ xác suất đơn biến, dùng để gán giá trị xác cho bất kỳ điểm nào trong khoảng tương ứng.

Định nghĩa 1. Cho o điểm dữ liệu không chắc chắn đa biến, $R \subseteq \mathfrak{R}^m$ là vùng mà o được xác định, $f: \mathfrak{R}^m \rightarrow \mathfrak{R}_0^+$ là hàm mật độ xác suất của o tại mỗi điểm $z \in R$. Điểm o được định nghĩa là một cặp (R, f) .

Định nghĩa 2. Cho $f^{(h)}: \mathfrak{R} \rightarrow \mathfrak{R}_0^+$ là hàm mật độ xác suất gán một giá trị xác suất cho mỗi $z \in I^{(h)}$, với $h \in [1 \dots m]$. Mỗi thuộc tính $a^{(h)}$ là một cặp $(I^{(h)}, f^{(h)})$, trong đó $I^{(h)} = [l^{(h)}, u^{(h)}]$ là khoảng xác định của $a^{(h)}$. Một điểm dữ liệu không chắc chắn đơn biến o là một bộ $(a^{(1)}, \dots, a^{(m)})$.

Hàm mật độ xác suất liên quan đến biểu diễn của điểm dữ liệu không chắc chắn đa biến có thể là liên tục hoặc rời rạc. Một hàm mật độ xác suất đa biến m -chiều liên tục được xác định trên một vùng $\mathfrak{R} \subseteq \mathfrak{R}^m$ là một hàm $f: \mathfrak{R} \rightarrow \mathfrak{R}_0^+$ sao cho: $\int_{z \in R} f(z) dz = 1$ và $\int_{z \in \mathfrak{R}^m \setminus R} f(z) dz = 0$.

Đặt $z_u \in \mathfrak{R}^m$ với $u \in [1 \dots v]$. Một hàm mật độ xác suất đa biến rời rạc m -chiều được xác định trên một tập các điểm $S = \{z_1, \dots, z_v\}$ là một hàm $f: \mathfrak{R}^m \rightarrow \mathfrak{R}_0^+$ sao cho:

$$\sum_{z \in S} f(z) = 1 \quad \text{và} \quad \sum_{z \in \mathbb{R}^m \setminus S} f(z) dz = 0$$

Đối với mô hình đơn biến, một hàm mật độ xác suất đơn biến liên tục hoặc rời rạc có thể được định nghĩa một cách đơn giản theo một hàm mật độ xác suất đa biến liên tục hoặc rời rạc, trong đó tập xác định của định nghĩa là một tập hợp con của $\mathfrak{R}$, tức $m = 1$.

2.2. Độ đo không chắc chắn

Để đo khoảng cách giữa các điểm dữ liệu không chắc chắn, cần xây dựng một khái niệm khoảng cách phù hợp trong điều kiện không chắc chắn gắn liền với thuật toán phân lớp [1-4]. Khái niệm này được cụ thể hóa thông qua hàm khoảng cách không chắc chắn, cho phép định lượng mức độ khác biệt giữa các điểm dữ liệu chưa chính xác.

Định nghĩa 3. [1,3] Cho một tập các điểm không chắc chắn $\mathcal{D} = \{o_1, \dots, o_n\}$. Hàm khoảng cách không chắc chắn được định nghĩa trên $\mathcal{D}$ là một hàm $\Delta: \mathcal{D} \times \mathcal{D} \times \mathfrak{R} \rightarrow \mathfrak{R}_0^+$ thỏa mãn các điều kiện:

$$\int_{z \in \mathfrak{R}} \Delta(o_i, o_j, z) dz = 1 \quad \forall o_i, o_j \in \mathcal{D} \quad \text{và} \quad \Delta(o_i, o_j, z) = \begin{cases} 1 & \text{nếu } i = j, z = 0 \\ 0 & \text{nếu } i = j, z \neq 0 \end{cases}$$

Với mỗi cặp điểm dữ liệu không chắc chắn $o_i, o_j, i \neq j$, Δ có thể được suy ra từ các hàm mật độ xác suất liên quan đến các đối tượng không chắc chắn. Định nghĩa của Δ phụ thuộc vào mô hình không chắc chắn được sử dụng để biểu diễn o_i và o_j . Hàm đặc trưng của sự kiện A là $I[A]$. Nếu sự kiện A xảy ra thì $I[A] = 1$, ngược lại $I[A] = 0$

Định nghĩa 4. [1,3] Đặt $o_i = (R_i, f_i)$ và $o_j = (R_j, f_j)$ là điểm dữ liệu dữ liệu không chắc chắn đa biến, $d(x, y)$ là khoảng cách Euclide giữa x và y , hàm khoảng cách không chắc chắn đa biến (Multivariate uncertain distance function) $\Delta(o_i, o_j, z)$ được định nghĩa như sau:

$$\Delta(o_i, o_j, z) = \int_{x \in R_i} \int_{y \in R_j} I[d(x, y) = z] f_i(x) f_j(y) dx dy$$

Đặt $\psi^{(h)}: \mathcal{D} \times \mathcal{D} \times \mathfrak{R} \rightarrow \mathfrak{R}$ với $\psi^{(h)}(o_i, o_j, x_h) = \int_{u \in I_i^{(h)}} \int_{v \in I_j^{(h)}} I[|u - v| = x_h] f_i^{(h)}(u) f_j^{(h)}(v) du dv$ và

$h \in [1 \dots m]$. Hàm vô hướng $f_{dist}: \mathfrak{R}^m \rightarrow \mathfrak{R}$ là một hàm ánh xạ từ không gian vector $\mathfrak{R}^m$ vào tập số

thực: $f_{dist} = \sqrt{\frac{1}{m} \sum_{h=1}^m x_h^2}$.

Định nghĩa 5. [1,3] Đặt $o_i = ((I_i^{(1)}, f_i^{(1)}), \dots, (I_i^{(m)}, f_i^{(m)}))$ và $o_j = ((I_j^{(1)}, f_j^{(1)}), \dots, (I_j^{(m)}, f_j^{(m)}))$ là các điểm dữ liệu dữ liệu không chắc chắn đơn biến, hàm khoảng cách không chắc chắn đơn biến $\Delta(o_i, o_j, z)$ được định nghĩa như sau:

$$\Delta(o_i, o_j, z) = \int_{x_1 \in \mathfrak{R}} \dots \int_{x_m \in \mathfrak{R}} I[f_{dist}(x_1, \dots, x_m) = z] \prod_{h=1}^m \psi^{(h)}(o_i, o_j, x_h) dx_1 \dots dx_m$$

Định nghĩa 6. [1,3] Cho một tập các điểm không chắc chắn $\mathcal{D} = \{o_1, \dots, o_n\}$, Δ là hàm khoảng cách không chắc chắn được định nghĩa trên $\mathcal{D}$. Khoảng cách không chắc chắn là một hàm $\delta: \mathcal{D} \times \mathcal{D} \rightarrow \mathfrak{R}_0^+$ được định nghĩa như sau: $\delta(o_i, o_j) = \int_{z \in \mathfrak{R}} z \Delta(o_i, o_j, z) dz$

Nếu o_i, o_j là hai điểm dữ liệu dữ liệu không chắc chắn đa biến thì

$$\delta(o_i, o_j) = \int_{x \in R_i} \int_{y \in R_j} d(x, y) f_i(x) f_j(y) dx dy$$

Đặt $\psi^{(h)}(o_i, o_j) = \int_{x \in I_i^{(h)}} \int_{y \in I_j^{(h)}} |x - y| f_i^{(h)}(x) f_j^{(h)}(y) dx dy$. Nếu o_i, o_j là hai điểm dữ liệu không

chắc chắn đơn biến thì $\delta(o_i, o_j) = f_{\text{dist}}(\psi^{(1)}(o_i, o_j), \dots, \psi^{(m)}(o_i, o_j))$

Ví dụ: Giả sử có hai điểm dữ liệu không chắc chắn o_i, o_j , mỗi điểm được biểu diễn bởi một phân phối xác suất một chiều. Điểm o_i phân phối đều trên đoạn $[2, 4]$. Điểm o_j phân phối đều trên đoạn $[5, 6]$. Theo định nghĩa 6, khoảng cách không chắc chắn $\delta(o_i, o_j)$ được tính bằng kỳ vọng của khoảng cách tuyệt đối giữa hai biến ngẫu nhiên $X \sim o_i$ và $Y \sim o_j$. Mật độ xác suất của X là

$f_i(x) = \frac{1}{4-2} = \frac{1}{2}$ với $x \in [2, 4]$ và mật độ xác suất của Y là $f_j(t) = 1$ với $y \in [5, 6]$. Khi đó,

$$\delta(o_i, o_j) = \int_2^4 \int_5^6 |x - y| \cdot \frac{1}{2} \cdot 1 dy dx. \quad \text{Vì } y > x \quad \text{nên} \quad \int_5^6 |x - y| dy = \int_5^6 (y - x) dy = \frac{11}{2} - x. \quad \text{Ta có,}$$

$$\delta(o_i, o_j) = \int_2^4 \frac{1}{2} \left(\frac{11}{2} - x \right) dx = 2.5. \quad \text{Khoảng cách không chắc chắn giữa hai điểm dữ liệu } o_i, o_j \text{ là } 2.5.$$

3. MÔ HÌNH ĐỀ XUẤT

3.1. Mô hình phân lớp UK-Medoids

Thuật toán UK-Medoids (Uncertain K-medoids) là một phần mở rộng của phương pháp K-Medoids [1,3], được thiết kế để xử lý dữ liệu không chắc chắn. Trong UK-Medoids, mỗi điểm dữ liệu không còn được biểu diễn bằng một giá trị cố định, mà bằng một phân phối xác suất. Khoảng cách giữa các điểm không được tính trực tiếp bằng khoảng cách Euclid, mà dựa trên kỳ vọng của khoảng cách giữa các phân phối xác suất tương ứng.

Thuật toán UK-Medoids

Đầu vào: Tập các điểm dữ liệu không chắc chắn $\mathcal{D} = \{o_1, \dots, o_n\}$, số lượng lớp k

Đầu ra: Tập k lớp $\{\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_k\}$

Bước 1: Tính các khoảng cách không chắc chắn theo từng cặp

$$d = \{(o_i, o_j) : \delta(o_i, o_j) \text{ for } o_i \text{ in } \mathcal{D} \text{ for } o_j \text{ in } \mathcal{D}\}$$

Bước 2: Khởi tạo ngẫu nhiên k medoids

$S = \text{initialize_medoids}(\mathcal{D}, k)$

while True:

$S_{\text{Prime}} = S.\text{copy}()$

$S = \text{set}()$

$\mathcal{C} = \{i : \text{set}() \text{ for } i \text{ in range}(k)\}$

Bước 3: Gán mỗi đối tượng cho medoid gần nhất

for o in $\mathcal{D}$:

$\text{minDist} = \text{float}('inf')$

$\text{assignedCluster} = -1$

 for j , medoid in enumerate(S_{Prime}):

$\Delta = d[(o, \text{medoid})]$

 if $\Delta < \text{minDist}$:

$\text{minDist} = \Delta$

$\text{assignedCluster} = j$

$\mathcal{C}[\text{assignedCluster}].\text{add}(o)$

```

# Bước 4: Tính lại medoids cho mỗi lớp
for i in range(k):
    cluster = C[i]
    bestMedoid = min(cluster, key=lambda o: sum(d[(o, o2)] for o2 in cluster))
    S.add(bestMedoid)
# Bước 5: Kiểm tra sự hội tụ
if S == SPrime: break
return List(C.values())

```

Định lý 1: Cho một tập các điểm không chắc chắn $\mathcal{D} = \{o_1, \dots, o_n\}$ và số lần lặp tối đa I , độ phức tạp của thuật toán UK-Medoids là $O(n^2 I)$.

4. THỰC NGHIỆM VÀ ĐÁNH GIÁ

4.1. Bộ dữ liệu thực nghiệm

EmoBank (A Corpus of Multi-Level and Multi-Dimensional Emotion Annotations) là một tập dữ liệu cảm xúc đa tầng (<https://github.com/JULIELab/EmoBank>) tiêu biểu, đồng thời được xem là một tập dữ liệu có yếu tố không chắc chắn trong lĩnh vực học máy và xử lý cảm xúc. Tập dữ liệu bao gồm hơn 10.000 câu ở cấp độ câu, được trích xuất từ nhiều nguồn khác nhau như Yelp, Reddit, tin tức, blog,... Mục tiêu chính của EmoBank là hỗ trợ việc đo lường và phân tích cảm xúc theo hai mô hình:

- Mô hình cảm xúc đa chiều bao gồm ba chiều cảm xúc cơ bản:
- + Valence (V): Mức độ tích cực, dao động từ buồn đến vui.
- + Arousal (A): Mức độ kích thích, từ bình tĩnh đến hưng phấn.
- + Dominance (D): Mức độ kiểm soát, từ bị động đến làm chủ.

Mỗi chiều được đánh giá trên thang điểm từ 1 đến 5, dựa trên trung bình ý kiến của nhiều người chủ thích (Annotators), phản ánh cảm xúc từ cả góc nhìn của người viết và người đọc. Điều này dẫn đến sự không nhất quán (Inconsistency) trong gán nhãn cảm xúc - một đặc điểm quan trọng đối với nghiên cứu biểu diễn cảm xúc không chắc chắn.

- Mô hình cảm xúc phân loại: Phân loại các câu theo một số cảm xúc cơ bản như joy, anger, sadness, fear, surprise và các trạng thái cảm xúc khác.

4.2. Độ đo đánh giá thực nghiệm

Để đánh giá kết quả phân lớp cho dữ liệu không chắc chắn, bài báo tìm hiểu các độ đo Silhouette Score (SIS), Davies-Bouldin Index (DBI) và Calinski-Harabasz Index (CHI).

Chỉ số SIS là một độ đo đánh giá chất lượng phân lớp, phản ánh mức độ gắn kết nội lớp và tách biệt liên lớp. Giá trị SI nằm trong khoảng $[-1, 1]$: giá trị gần 1 cho thấy điểm dữ liệu được phân lớp tốt; gần 0 cho thấy điểm nằm gần ranh giới lớp và gần -1 cho thấy khả năng bị phân sai lớp.

Chỉ số DBI sử dụng để đánh giá mức độ tương đồng giữa các lớp trong phân lớp. Giá trị DBI càng nhỏ phản ánh chất lượng phân lớp càng cao, do các lớp được phân tách rõ và có tính gắn kết nội bộ tốt. Ngược lại, DBI lớn cho thấy sự chồng lấn giữa các lớp hoặc phân lớp không hiệu quả.

Chỉ số CHI đánh giá chất lượng phân lớp dựa trên tỷ lệ giữa phân tán giữa các lớp và phân tán trong lớp. Giá trị CHI càng lớn cho thấy phân lớp càng tốt, với các lớp được tách biệt rõ ràng và tập trung nội lớp cao.

4.3. Phân tích và đánh giá

Các hàm mật độ xác suất liên tục (PDF - Probability Density Function) như phân phối đồng đều (Uniform Distribution) và phân phối chuẩn (Normal Distribution), cùng với hàm khối xác suất rời rạc (PMF - Probability Mass Function) như phân phối nhị thức (Binomial Distribution), là những

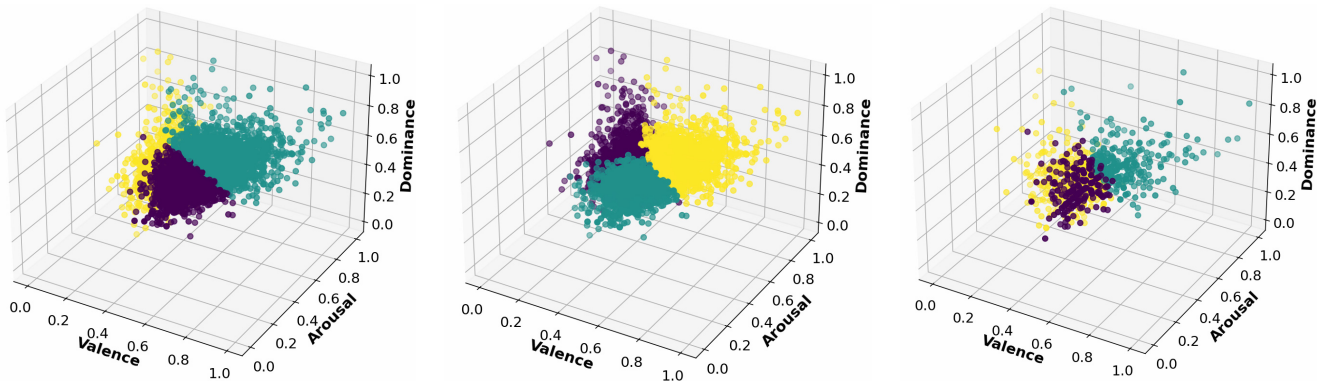
thành phần cốt lõi trong xác suất thống kê và học máy. Trong thực nghiệm, các hàm này được sử dụng để đánh giá ảnh hưởng của việc chọn phân phối xác suất khi triển khai các mô hình phân lớp.

Hình 1 thể hiện kết quả phân lớp dữ liệu sử dụng thuật toán UK-Medoids theo ba mô hình phân phối xác suất phân phối đồng đều, phân phối chuẩn và phân phối nhị thức. Dữ liệu được phân lớp trên không gian đặc trưng gồm ba trục: Valence, Arousal, Dominance tương ứng với ba thuộc tính trong tập dữ liệu. Thuật toán UK-Medoids cho thấy khả năng phân lớp tốt trên nhiều loại phân phối xác suất khác nhau. Khi sử dụng phân phối đồng đều, các lớp được phân tách tương đối rõ ràng, có hình dạng đồng đều và kích thước không quá chênh lệch. Tuy nhiên, ranh giới giữa các lớp vẫn còn sự chồng lấn nhẹ, đặc biệt ở vùng tiếp giáp. Với phân phối chuẩn, lớp màu vàng được tách biệt rõ ràng và chiếm ưu thế ở vùng có giá trị Dominance cao. Trong khi đó, lớp xanh và tím có xu hướng giao nhau nhẹ ở vùng biên, dẫn đến khả năng bị nhầm lẫn tại các vùng chuyển tiếp giữa lớp. Khi sử dụng phân phối nhị thức, các lớp có xu hướng tập trung với mật độ điểm cao tại vùng trung tâm. Một số điểm rải rác có thể là nhiễu hoặc không thuộc về lớp rõ ràng. Cụm màu vàng mở rộng đáng kể về phía trục Valence và Dominance, tuy nhiên phân lớp có dấu hiệu không đều, với một số lớp bị co lại hoặc lệch tâm so với cấu trúc chung.

UK-Medoids sử dụng phân phối đồng đều

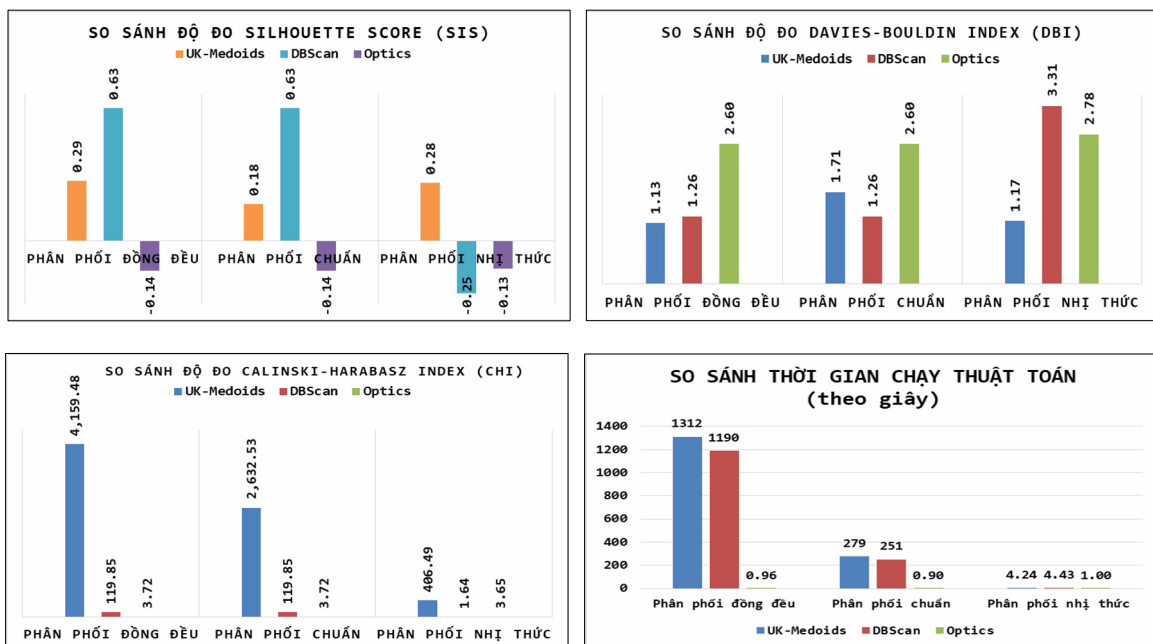
UK-Medoids sử dụng phân phối chuẩn

UK-Medoids sử dụng phân phối nhị thức



Hình 1. Kết quả phân lớp dữ liệu không chắc chắn sử dụng thuật toán UK-Medoids

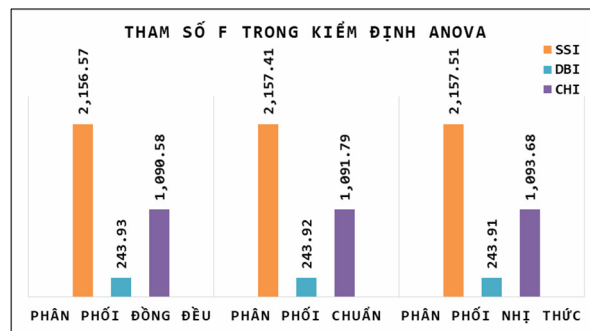
Hình 2 minh họa kết quả thực nghiệm so sánh hiệu quả của thuật toán UK-Medoids với DBScan và Optics, dựa trên hai tiêu chí: chất lượng phân lớp và hiệu suất tính toán.



Hình 2. So sánh độ đo chất lượng phân lớp và hiệu quả của thuật toán

Thuật toán DBScan đạt chỉ số SIS cao nhất với các phân phối đồng đều và chuẩn, trong khi đó UK-Medoids thể hiện hiệu quả vượt trội với phân phối nhị thức, nơi mà DBScan và Optics cho kết quả SIS âm, phản ánh khả năng phân lớp kém. Với chỉ số DBI, UK-Medoids thường đạt giá trị thấp nhất hoặc đứng thứ hai, cho thấy các lớp được phân biệt rõ, đặc biệt là trong phân phối nhị thức. Với chỉ số CHI, UK-Medoids vượt trội tuyệt đối trên cả ba loại phân phối, là minh chứng cho khả năng tạo lớp rõ ràng và chặt chẽ của thuật toán này. Xét về mặt hiệu suất, Optics luôn có tốc độ nhanh nhất, trong khi UK-Medoids chậm hơn đáng kể, đặc biệt ở phân phối đồng đều và chuẩn. Với phân phối nhị thức, UK-Medoids và DBScan có thời gian tương đương nhưng vẫn kém Optics. Nhìn chung, UK-Medoids thích hợp cho các bài toán cần chất lượng phân lớp cao, còn Optics phù hợp khi ưu tiên tốc độ.

Kết quả kiểm định ANOVA trong Hình 3 cho thấy tham số F ở cả ba độ đo SIS, DBI và CHI đều đạt giá trị rất lớn, xấp xỉ lần lượt là 2156, 244, và 1990 trên cả ba dạng phân phối dữ liệu (đồng đều, chuẩn, nhị thức), đồng thời giá trị p xấp xỉ 0. Điều này khẳng định sự khác biệt về hiệu quả phân cụm giữa các thuật toán UK-Medoids, DBScan và Optics là có ý nghĩa thống kê, cho thấy lựa chọn thuật toán có tác động quan trọng đến chất lượng phân lớp dữ liệu không chắc chắn.



Hình 3. Kết quả kiểm định thống kê Anova

5. KẾT LUẬN

Bài báo nghiên cứu vấn đề phân lớp dữ liệu không chắc chắn thông qua việc tìm hiểu một số độ đo khoảng cách không chắc chắn giữa các điểm dữ liệu. Dựa trên đó, mô hình phân lớp UK-Medoids được đề xuất, bao gồm bốn bước chính: (1) Tính toán khoảng cách không chắc chắn giữa các cặp điểm dữ liệu; (2) Khởi tạo ngẫu nhiên k điểm medoid ban đầu; (3) Gán mỗi điểm dữ liệu vào cụm có medoid gần nhất theo khoảng cách đã tính; (4) Cập nhật lại medoid cho từng cụm. Mô hình được triển khai và đánh giá trên tập dữ liệu không chắc chắn EmoBank. Kết quả thực nghiệm cho thấy thuật toán UK-Medoids có khả năng phân lớp hiệu quả trên nhiều loại phân phối xác suất khác nhau. Chất lượng phân lớp được đo lường bằng ba độ đo SIS, DBI và CHI đều cho kết quả khả quan. Ngoài ra, thời gian thực thi của thuật toán vẫn nằm trong phạm vi chấp nhận được, qua đó khẳng định tính khả thi của mô hình trong ứng dụng thực tiễn. Trong tương lai, mô hình đề xuất sẽ được mở rộng thử nghiệm trên các bộ dữ liệu không chắc chắn thực tế với quy mô mẫu lớn hơn.

6. TÀI LIỆU THAM KHẢO

- [1] Gullo, F., Ponti, G., Tagarelli, A. 2008. Clustering Uncertain Data Via K-Medoids. In: Greco, S., Lukasiewicz, T. (eds) Scalable Uncertainty Management. SUM 2008. Lecture Notes in Computer Science, vol 5291. Springer, Berlin, Heidelberg.
- [2] Erich, S., Alexander, K., Tobias, E., Andreas, Z., Klaus Arthur, S., Arthur Z., 2015. A framework for clustering uncertain data. Proc. VLDB Endow. 8(12): 1976-1979.
- [3] Tavakkol, B., Son, Y. 2021. Fuzzy kernel K-medoids clustering algorithm for uncertain data objects. Pattern Anal Applic 24, 1287-1302.
- [4] Nicholls K, Kirk PDW, Wallace C. 2024. Bayesian clustering with uncertain data. PLoS Comput Biol 20(9).

PHÁT HIỆN TRẠNG THÁI LỖI MÁY CÔNG NGHIỆP, ỨNG DỤNG TRÊN TẬP DỮ LIỆU MIMII

Nguyễn Đắc Phương Thảo, Nguyễn Thị Đình, Hoàng Quang Vinh, Bùi Minh Đức

Trường Đại học Thủy lợi, email: ndpthao@thu.edu.vn

1. GIỚI THIỆU CHUNG

Trong sản xuất công nghiệp, hư hỏng thiết bị có thể gây gián đoạn lớn; giám sát bằng tín hiệu âm thanh là giải pháp hiệu quả nhờ giàu thông tin, dễ thu thập và không cần tiếp xúc [1]. Nghiên cứu này sử dụng MIMII Dataset [3] để phân loại đa lớp trạng thái hoạt động (bình thường/bất thường) của thiết bị (fan, pump, slider, valve) dựa trên đặc trưng âm thanh và huấn luyện bằng LightGBM [2], so sánh với SVM và XGBoost. SVM thích hợp cho dữ liệu nhỏ, LightGBM nhanh và gọn nhẹ phù hợp thiết bị biên, còn XGBoost ổn định và kháng nhiễu tốt nhưng huấn luyện chậm. Mục tiêu là xây dựng hệ thống phát hiện bất thường chính xác, linh hoạt và triển khai được trong bảo trì dự đoán.

Bên cạnh đó, học sâu (CNN, Autoencoder) đã chứng minh hiệu quả trong việc tự động trích xuất đặc trưng từ tín hiệu âm thanh trong khi các mô hình hiện đại như Transformer (AST) cùng với cách tiếp cận học tự giám sát tiếp tục mở ra tiềm năng lớn trong nâng cao độ chính xác và khả năng khái quát hóa của hệ thống [4][5][6][7]. Tuy nhiên, theo nhiều nghiên cứu các kỹ thuật học sâu thường đòi hỏi tài nguyên và dữ liệu lớn, gây khó khăn khi triển khai trong môi trường hạn chế. Việc khai thác các mô hình học máy truyền thống hiệu quả như LightGBM và XGBoost kết hợp trích xuất đặc trưng thủ công vẫn mang ý nghĩa thực tiễn cao, cân bằng giữa độ chính xác và khả năng triển khai.

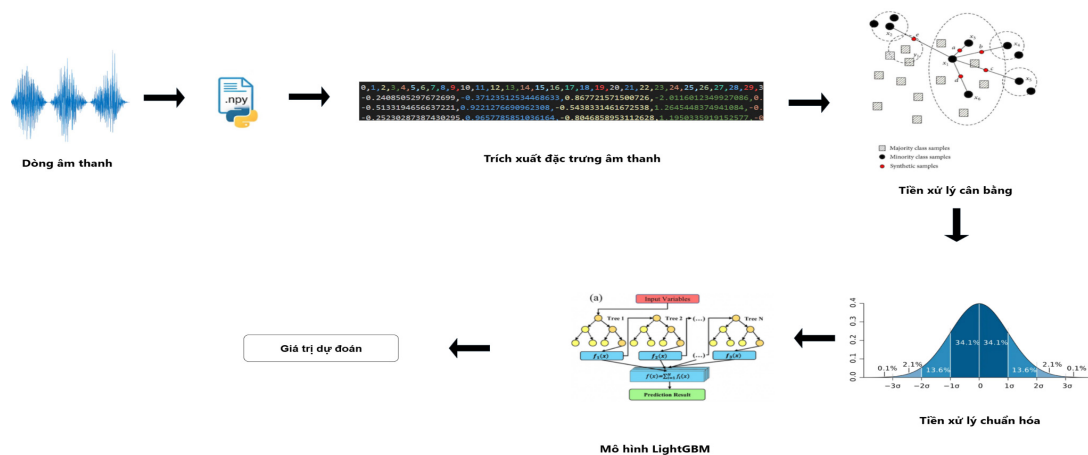
2. PHƯƠNG PHÁP NGHIÊN CỨU

2.1. Tiền xử lý

Âm thanh đầu vào được chuẩn hóa về độ dài cố định 9 giây với tần số lấy mẫu 16 kHz. Từ mỗi đoạn, trích xuất 80 đặc trưng bao gồm: hệ số MFCC (20 hệ số $\times$ {trung bình, độ lệch chuẩn} $\rightarrow$ 40 chiều), Chroma (12 hệ số $\times$ 2 $\rightarrow$ 24 chiều), Spectral Contrast (7 dải $\times$ 2 $\rightarrow$ 14 chiều), cùng hai chỉ số Spectral Centroid và Zero-Crossing Rate [3].

Để khắc phục mất cân bằng lớp, dữ liệu huấn luyện được xử lý bằng SMOTE. Sau đó, toàn bộ tập dữ liệu được chia 80% huấn luyện – 20% kiểm thử theo phương pháp phân tầng với `random_state=42`.

2.2. Light Gradient Boosting Machine – LightGBM

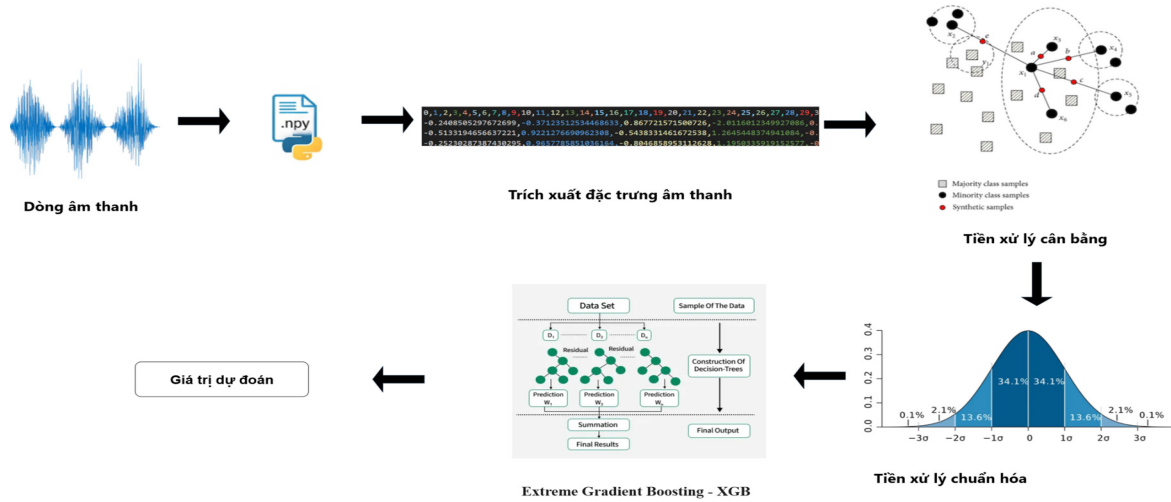


Hình 2.1. Pipeline của mô hình LightGBM

Mô hình tận dụng tốt mối quan hệ phi tuyến giữa các đặc trưng phổ, xử lý hiệu quả dữ liệu lớn, mất cân bằng và phù hợp triển khai thực tế.

LightGBM được huấn luyện với `objective='multiclass'`, `num_class=8`, `num_leaves=31`, `learning_rate=0.05`, `n_estimators=100`. Mô hình dự đoán nhãn và xác suất 8 lớp (4 máy $\times$ 2 trạng thái), đánh giá bằng AUC-ROC, AUC-PR, F1-score (macro) và MCC, đồng thời phân tích riêng từng loại máy theo dạng nhị phân (normal vs abnormal) [8].

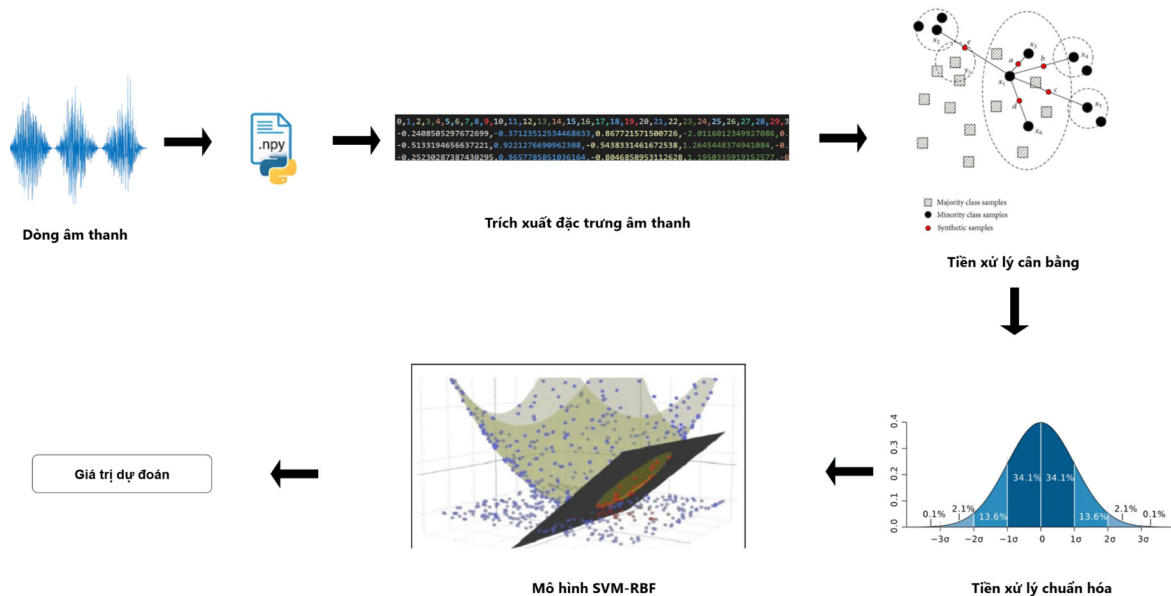
2.3. Extreme Gradient Boosting - XGB



Hình 2.2. Pipeline của mô hình XGB

Mô hình XGBoost được xây dựng dựa trên khả năng học phi tuyến của gradient boosting, phù hợp với tính chất phức tạp và đa dạng của dữ liệu âm thanh. Quá trình huấn luyện sử dụng tham số `eval_metric='logloss'` cho bài toán phân loại nhị phân (normal/abnormal), được thực hiện riêng lẻ cho từng loại máy (fan, pump, slider, valve) cũng như trên tập dữ liệu tổng hợp.

2.4. Support Vector Machine – SVM



Hình 2.3. Pipeline của mô hình SVM

Mô hình phân loại được xây dựng dựa trên Support Vector Machine (SVM) với hàm nhân Radial Basis Function (RBF). Các siêu tham số được lựa chọn thủ công với giá trị cố định: tham số điều chuẩn $C=10$ và $\gamma = \text{"scale"}$. Sau khi huấn luyện, mô hình có thể dự đoán nhãn và xác suất cho 8 lớp tương ứng với bốn loại máy (fan, pump, slider, valve) trong hai trạng thái hoạt động (normal/abnormal).[8].

3. KẾT QUẢ NGHIÊN CỨU

AUC-PR (Area Under the Precision–Recall Curve) đánh giá khả năng phân biệt lớp trong dữ liệu mất cân bằng, nhấn mạnh vào precision và recall thay vì độ chính xác tổng thể.

MCC (Matthews Correlation Coefficient) đo tương quan cân bằng giữa dự đoán và thực tế, được coi là thước đo đáng tin cậy hơn accuracy hay F1 trong trường hợp phân bố lớp không đều.

3.1. Kết quả của mô hình LightGBM

Bảng 3.1. Kết quả của mô hình LightGBM

	AUC-ROC	AUC-PR	F1-Score	MCC
Fan	0,9532	0,8860	0,7933	0,7224
Pump	0,9683	0,8861	0,8176	0,7969
Slider	0,9923	0,9839	0,9533	0,9414
Valve	0,9472	0,8450	0,7799	0,7565
All	0,9853	0,8989	0,8293	0,8583

Mô hình LightGBM đạt hiệu suất cao trên cả 4 loại thiết bị, với các chỉ số AUC-ROC > 0.94 và F1-score > 0.77. Đặc biệt, Slider có kết quả xuất sắc nhất (AUC-ROC: 0.9923, F1: 0.9533, MCC: 0.9414), cho thấy khả năng phân biệt rất tốt giữa trạng thái bình thường và bất thường. Kết quả tổng hợp toàn bộ dữ liệu (All) cũng rất ấn tượng: AUC-ROC: 0.9853, F1-score: 0.8293, MCC: 0.8583, chứng tỏ mô hình tổng quát hóa tốt trên nhiều loại máy.

3.2. Kết quả của mô hình XGB

Bảng 3.2. Kết quả của mô hình XGB

	AUC-ROC	AUC-PR	F1-Score	MCC
Fan	0,9882	0,9858	0,9476	0,8938
Pump	0,9994	0,9994	0,9873	0,9747
Slider	0,9984	0,9986	0,9837	0,9673
Valve	0,9988	0,9988	0,9846	0,9691
All	0,9879	0,9885	0,9465	0,8917

Mô hình XGBoost đạt hiệu suất xuất sắc trong phát hiện bất thường âm thanh máy móc, với AUC-ROC trung bình 0.9879 và F1-score 0.9465. Pump và Slider đạt kết quả tốt nhất (F1-score lần lượt 0.9873 và 0.9837), cho thấy mô hình xử lý hiệu quả đặc trưng âm thanh của hai thiết bị này. Fan và Valve cũng cho thấy độ chính xác cao, phản ánh khả năng tổng quát hóa mạnh của mô hình. Các chỉ số MCC (0.8917) và AUC-PR (0.9885) khẳng định tính ổn định trong bối cảnh mất cân bằng nhãn. Việc kết hợp đặc trưng phổ âm thanh với chuẩn hóa và SMOTE đã góp phần nâng cao hiệu năng. Hiệu suất có thể tiếp tục cải thiện qua tối ưu tham số và chiến lược tăng cường dữ liệu.

3.3. Kết quả của mô hình SVM

Bảng 3.3. Kết quả của mô hình SVM

	AUC-ROC	AUC-PR	F1-Score	MCC
Fan	0.9696	0.9305	0.8366	0.7846
Pump	0.9644	0.8652	0.7702	0.7542
Slider	0.9923	0.9804	0.9318	0.9165
Valve	0.9513	0.8575	0.7425	0.7246
All	0.9694	0.9084	0.8203	0.7950

Mô hình SVM đạt kết quả rất tốt trên cả 4 thiết bị, đặc biệt với Slider ($F1 = 0.929$, $MCC = 0.91$). Fan và Pump cũng có hiệu suất rất ổn định. Valve tuy vẫn thấp hơn các máy còn lại, nhưng vẫn đạt mức chấp nhận được ($F1 \approx 0.74$). So với Autoencoder, SVM vượt trội rõ rệt nếu có dữ liệu đã được gán nhãn và phân loại lỗi rõ ràng.

3.4. Phân tích sự khác biệt hiệu năng

XGBoost đạt kết quả tốt nhất nhờ cơ chế regularization (thêm tham số phạt để kiểm soát độ phức tạp của cây, tránh quá khớp) và tối ưu bậc hai, giúp mô hình ổn định hơn trên dữ liệu nhiễu và mất cân bằng. LightGBM dù huấn luyện nhanh nhưng dễ overfit nếu không tinh chỉnh tham số, nên kém ổn định hơn. SVM nhạy cảm với nhiễu và khó mở rộng khi dữ liệu lớn, dẫn tới hiệu quả thấp hơn hai mô hình boosting.

XGBoost cho thấy ưu thế nhờ khả năng khai thác thông tin từ dữ liệu mất cân bằng, kết hợp với regularization mạnh và thuật toán tối ưu bậc hai, giúp mô hình duy trì độ ổn định trên cả bốn loại máy. Ngược lại, LightGBM có tốc độ huấn luyện nhanh nhưng nhạy cảm với phân bố dữ liệu: hiệu suất rất cao trên Slider ($F1 > 0.95$) nhưng giảm mạnh ở Fan và Valve. Điều này phản ánh việc LightGBM dễ bị ảnh hưởng bởi đặc thù dữ liệu từng thiết bị và nguy cơ overfitting khi số lá quá nhiều. Trong khi đó, SVM hoạt động khá tốt khi ranh giới giữa normal/abnormal rõ ràng, nhưng kém ổn định ở dữ liệu nhiễu cao và khó mở rộng cho tập dữ liệu lớn, do chi phí tính toán tăng nhanh theo số mẫu và số chiều đặc trưng.

4. KẾT LUẬN

So với nghiên cứu trong bài báo [8], nghiên cứu này có điểm nhấn ở việc tập trung khai thác các mô hình học máy truyền thống (XGBoost, LightGBM, SVM) kết hợp đặc trưng thủ công để đạt hiệu quả cao trên dữ liệu MIMII. Đóng góp đặc thù nằm ở việc chỉ ra rằng các mô hình này, khi được xử lý và tối ưu phù hợp, có thể đem lại kết quả tốt và dễ triển khai trong môi trường công nghiệp hạn chế tài nguyên.

Đóng góp chính của nghiên cứu là chứng minh rằng các mô hình học máy truyền thống kết hợp đặc trưng thủ công vẫn có thể đạt hiệu quả cao và dễ triển khai trong môi trường công nghiệp hạn chế tài nguyên. Trong khi đó, các phương pháp học sâu hiện đại như Autoencoder [4], IDNN [5] hay Audio Spectrogram Transformer (AST) [6] đã đạt độ chính xác rất cao trên các tập dữ liệu âm thanh công nghiệp, song thường đòi hỏi GPU và dữ liệu quy mô lớn. Do đó cách tiếp cận trong nghiên cứu này cung cấp một giải pháp thay thế thực tiễn.

Bên cạnh đó, nghiên cứu này còn sử dụng SMOTE để xử lý mất cân bằng dữ liệu. Tuy nhiên, SMOTE cũng có những hạn chế: kỹ thuật này tạo ra mẫu tổng hợp dựa trên nội suy tuyến tính, có thể làm giảm tính đa dạng và đôi khi sinh ra mẫu không thực sự đại diện cho phân bố dữ liệu gốc. Trong trường hợp dữ liệu nhiễu, SMOTE thậm chí có thể khuếch đại nhiễu, làm suy giảm hiệu quả mô hình.

Nghiên cứu mới dừng ở mức so sánh các mô hình với cấu hình mặc định, chưa có đánh giá hệ thống khi thay đổi tham số (số đặc trưng, số vòng lặp, mức độ nhiễu). Việc phân tích các yếu tố này trong nghiên cứu tiếp theo sẽ làm rõ tính ổn định, khả năng kháng nhiễu và hiệu quả triển khai của mô hình. Mặc dù MIMII cung cấp nền tảng quan trọng cho phát hiện bất thường, tập dữ liệu vẫn còn hạn chế: nhiễu chưa phản ánh đầy đủ điều kiện thực tế, chỉ bao gồm bốn loại thiết bị, dữ liệu bất thường mất cân bằng và thời lượng mẫu ngắn, dễ bỏ sót lỗi. Do đó, các hướng nghiên cứu tương lai cần tập trung vào nâng cao khả năng khái quát hóa thông qua học chuyển giao và tự giám sát, áp dụng tăng cường dữ liệu và lọc nhiễu, mở rộng tập dữ liệu với nhiều loại thiết bị, cũng như phát triển mô hình lai kết hợp phát hiện bất thường và phân loại lỗi. Bên cạnh đó, tối ưu hóa mô hình nhẹ và triển khai trên thiết bị biên (IoT, edge devices) là bước quan trọng để ứng dụng hệ thống trong môi trường công nghiệp thực tế.

5. TÀI LIỆU THAM KHẢO

- [1] J. Piotrowski, *Proactive Maintenance for Mechanical Systems*. Amsterdam, Netherlands: Elsevier, 2011. [Online]. Available: https://api.pageplace.de/preview/DT0400.9781483292595_A23885366/preview-9781483292595_A23885366.pdf.
- [2] G. Ke, Q. Meng, T. Finley, T. Wang, W. Chen, W. Ma, Q. Ye, and T. Liu, “LightGBM: A highly efficient gradient boosting decision tree,” in *Proc. Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2017, pp. 3146–3154. [Online]. Available: <https://papers.nips.cc/paper/2017/file/6449f44a102fde848669bdd9eb6b76fa-Paper.pdf>.
- [3] H. Purohit, R. Tanabe, K. Ichige, Y. Kawaguchi, T. Endo, Y. Nikaido, K. Suefusa, and N. Harada, “MIMII dataset: Sound dataset for malfunctioning industrial machine investigation and inspection,” in *Proc. Detection and Classification of Acoustic Scenes and Events Workshop (DCASE)*, 2019. [Online]. Available: <https://arxiv.org/abs/1909.09347>.
- [4] Y. Koizumi, S. Saito, H. Uematsu, Y. Kawachi, and N. Harada, “Unsupervised detection of anomalous sound based on deep learning and the Neyman–Pearson lemma,” *IEEE/ACM Trans. Audio, Speech, and Language Processing*, vol. 27, no. 1, pp. 212–224, Jan. 2019, doi: 10.1109/TASLP.2018.2877011. [Online]. Available: <https://arxiv.org/pdf/1810.09133>.
- [5] R. Tanabe, H. Purohit, K. Suefusa, T. Endo, T. Nishida, and Y. Kawaguchi, “Anomalous sound detection based on interpolation deep neural network,” in *Proc. DCASE Workshop*, 2020, pp. 151–155. [Online]. Available: <https://arxiv.org/abs/2005.09234>.
- [6] Y. Gong, Y.-A. Chung, and J. Glass, “AST: Audio spectrogram transformer,” in *Proc. Interspeech*, 2022, pp. 561–565, doi: 10.21437/Interspeech.2022-10543. [Online]. Available: <https://arxiv.org/pdf/2104.01778>.
- [7] S. Giri, D. K. Weber, M. E. P. Davies, and M. D. Plumbley, “Self-supervised learning for anomalous sound detection in machine condition monitoring,” in *Proc. IEEE Int. Conf. Acoustics, Speech, and Signal Processing (ICASSP)*, 2023, pp. 1–5, <https://ieeexplore.ieee.org/document/10095133>.
- [8] L. Ganterta, G. Zanin, and J. Silva, “Multiclass classification of faulty industrial machinery using sound samples,” *GTA Technical Reports*, 2024. [Online]. Available: <https://www.gta.ufrj.br/ftp/gta/TechReports/GZS24.pdf>.

PHÂN LOẠI TIẾNG THỞ TRONG GIÁC NGỦ ĐỂ PHÁT HIỆN NGỪNG THỞ KHI NGỦ

Nguyễn Đắc Phương Thảo, Nguyễn Quang Hiếu, Bùi Minh Đức, Đỗ Thị Hiền Lương

Trường Đại học Thủy lợi, email: ndpthao@tlu.edu.vn

1. GIỚI THIỆU CHUNG

Ngưng thở khi ngủ tắc nghẽn (Obstructive Sleep Apnea – OSA) là một rối loạn giấc ngủ phổ biến, đặc trưng bởi ngưng thở hoặc giảm thông khí kéo dài ≥ 10 giây, lặp lại nhiều lần trong suốt đêm. Tình trạng này gây gián đoạn giấc ngủ, giảm oxy máu, kích hoạt hệ thần kinh giao cảm, từ đó làm tăng nguy cơ tăng huyết áp, bệnh tim mạch, đột quỵ, đái tháo đường và suy giảm nhận thức [1].

Theo phân tích dịch tễ học của Benjafield và cộng sự, tỷ lệ mắc OSA ở Bắc Mỹ là 15–30% nam giới và 10–15% nữ giới; toàn cầu có khoảng 936 triệu người trưởng thành (30–69 tuổi) có nguy cơ mắc [2]. Tại Việt Nam, tình trạng này ngày càng được quan tâm nhiều hơn, đặc biệt ở những người trung niên, béo phì hoặc có bất thường giải phẫu vùng hầu họng [1].

Chiến lược điều trị hiện nay chủ yếu nhằm kiểm soát các yếu tố nguy cơ và cải thiện sự thông khí đường thở. Các phương pháp bao gồm sử dụng thiết bị thông khí áp lực dương liên tục (Continuous Positive Airway Pressure – CPAP), đặt khí cụ trong miệng nhằm duy trì độ mở của đường hô hấp trên, hoặc can thiệp phẫu thuật các cấu trúc như màn hầu, lưỡi gà và khẩu cái mềm. Trong một số trường hợp nặng hoặc kháng trị, liệu pháp kích thích thần kinh hạ thiết cũng có thể được xem xét như một lựa chọn thay thế [1].

Đa ký giấc ngủ (Polysomnography – PSG) là tiêu chuẩn vàng chẩn đoán OSA, nhưng bị hạn chế bởi thiết bị chuyên môn, chi phí cao, yêu cầu nhân lực kỹ thuật và gây khó chịu cho bệnh nhân khi thực hiện tại bệnh viện vào ban đêm. Điều này dẫn đến tỷ lệ chẩn đoán thấp, kể cả ở các nước phát triển. Tại Hoa Kỳ, chi phí y tế liên quan đến OSA ước tính khoảng 12,4 tỷ USD năm 2015 [2].

Các giải pháp sàng lọc thay thế có độ chính xác cao, chi phí thấp và dễ sử dụng tại nhà đang thu hút nhiều sự quan tâm. Trong đó, các phương pháp sử dụng trí tuệ nhân tạo (AI) để phân tích tiếng thở khi ngủ thu từ micro môi trường nổi lên như một hướng tiếp cận đầy tiềm năng. Không yêu cầu gắn cảm biến lên cơ thể, không gây khó chịu, dễ triển khai nhiều đêm liên tiếp – mô hình này đặc biệt phù hợp với các thiết bị hạn chế tài nguyên như điện thoại di động hoặc hệ thống nhúng [3]. Hướng đi này mở ra triển vọng cho việc sàng lọc và quản lý OSA tại cộng đồng với chi phí hợp lý và độ bao phủ cao.

2. XÂY DỰNG MÔ HÌNH

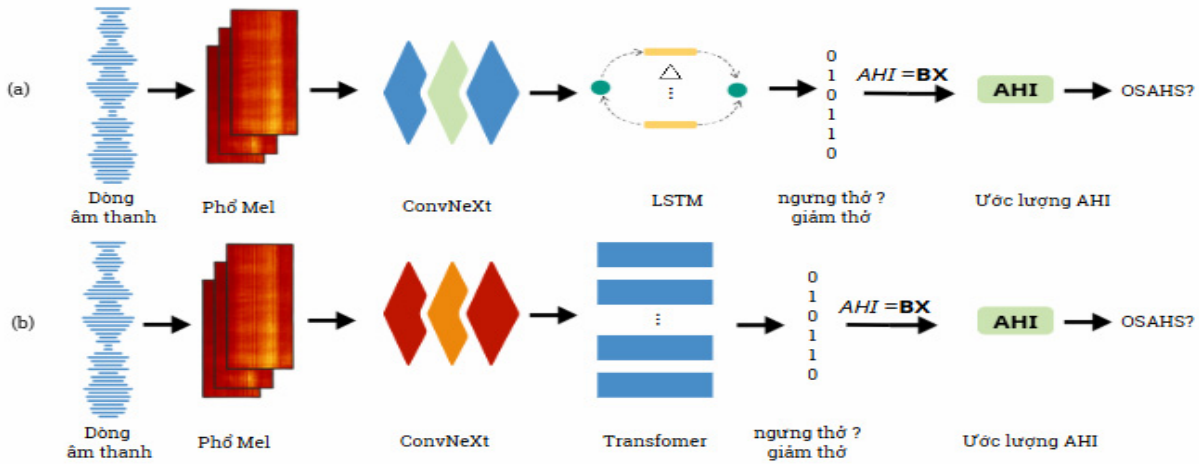
Nghiên cứu này xây dựng một hệ thống sàng lọc hội chứng ngưng thở khi ngủ tắc nghẽn (Obstructive Sleep Apnea–Hypopnea Syndrome – OSAHS) bằng cách nhận diện ngưng thở và giảm thông khí từ tiếng thở ban đêm, tập trung giảm tham số mô hình để đảm bảo tính gọn nhẹ và hiệu quả tính toán.

Mô hình chính sử dụng kiến trúc ConvNeXt kết hợp LSTM, được thiết kế theo hướng tiếp cận của Song và cộng sự (2025) [3], tập trung vào tối ưu hóa hiệu suất nhận diện các sự kiện hô hấp bất thường và đơn giản hóa quy trình xử lý tín hiệu. Ngoài ra, nghiên cứu còn phát triển thêm một biến thể sử dụng ConvNeXt kết hợp Transformer nhằm đánh giá và so sánh hiệu quả phân loại cũng như khả năng ước lượng chỉ số AHI (Apnea–Hypopnea Index).

Cấu trúc hệ thống bao gồm hai pipeline xử lý song song: (a) ConvNeXt + LSTM và (b) ConvNeXt + Transformer. Dữ liệu đầu vào sử dụng là tập PSG-Audio công khai, gồm 305 bản ghi giấc ngủ trọn đêm được thu tại Bệnh viện Đa khoa Sismanoglio–Amalia Fleming (Athens), với định dạng EDF, tần số lấy mẫu 48kHz và độ sâu 16 bit [4].

Các bản ghi được chia thành các đoạn âm thanh dài 9 giây, trượt khung 3 giây, và gán nhãn dựa trên mức độ chồng lấp với chu kỳ hô hấp trước – sau sự kiện ngưng thở hoặc giảm thông khí. Phổ Mel được trích xuất, chuẩn hóa làm đầu vào cho mô hình ConvNeXt kết hợp LSTM hoặc Transformer để phát hiện mẫu âm thanh đặc trưng. Kết quả phân loại được dùng trong mô hình hồi quy tuyến tính để ước lượng chỉ số AHI, đánh giá mức độ nghiêm trọng của OSAHS.

Cả hai biến thể đều duy trì tổng số tham số khoảng 2,1 triệu, giúp đảm bảo độ nhẹ của mô hình, tiết kiệm bộ nhớ và phù hợp với các ứng dụng có giới hạn tài nguyên. Toàn bộ quy trình xử lý tín hiệu được tinh giản gồm ba bước chính: (1) trích xuất phổ Mel từ âm thanh ban đêm, (2) phân loại các đoạn âm thanh chứa sự kiện hô hấp bất thường, và (3) ước lượng AHI bằng hồi quy tuyến tính.



Hình 2.1. Khung tổng thể hệ thống sàng lọc OSAHS với hai pipeline:
(a) ConvNeXt + LSTM và (b) ConvNeXt + Transformer.

2.1. Xử lý tín hiệu âm thanh

Toàn bộ tín hiệu từ bộ dữ liệu PSG-Audio được chia thành các đoạn dài 9 giây, trượt cửa sổ 3 giây để bao quát ít nhất hai chu kỳ hô hấp (2,7–6 s/chu kỳ). Phổ công suất tính bằng STFT (khung 1024 mẫu, bước trượt 499 mẫu, cửa sổ Blackman), chiếu qua 64 bộ lọc Mel (60 Hz–22,5 kHz), tạo ma trận Mel-spectrum kích thước 64×684. Ma trận được chuẩn hóa trung bình bằng 0 và độ lệch chuẩn bằng 1 trước khi đưa vào mô hình học sâu. Thay vì áp dụng các thuật toán khử nhiễu truyền thống như spectral subtraction, chúng tôi để mạng nơ-ron sâu tự học trực tiếp trên Mel-spectrum đã chuẩn hóa, vừa đơn giản hóa quy trình xử lý vừa phù hợp với các thiết bị có tài nguyên hạn chế [3].

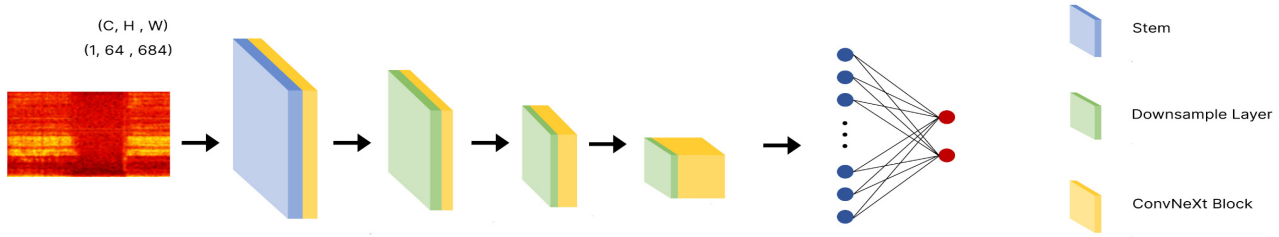
2.2. Chia dữ liệu

Nghiên cứu sử dụng chiến lược chia dữ liệu độc lập theo chủ thể (subject-independent split), phân bổ toàn bộ dữ liệu của một bệnh nhân vào duy nhất một tập (train, validation hoặc test) với tỷ lệ 80% train, 10% validation, 10% test. Phương pháp này tránh chồng lấp, đảm bảo đánh giá trên bệnh nhân mới, phản ánh thực tế triển khai và tăng khả năng tổng quát hóa mô hình.

2.3. Kiến trúc mô hình

2.3.1. Mạng trích đặc trưng ConvNeXt

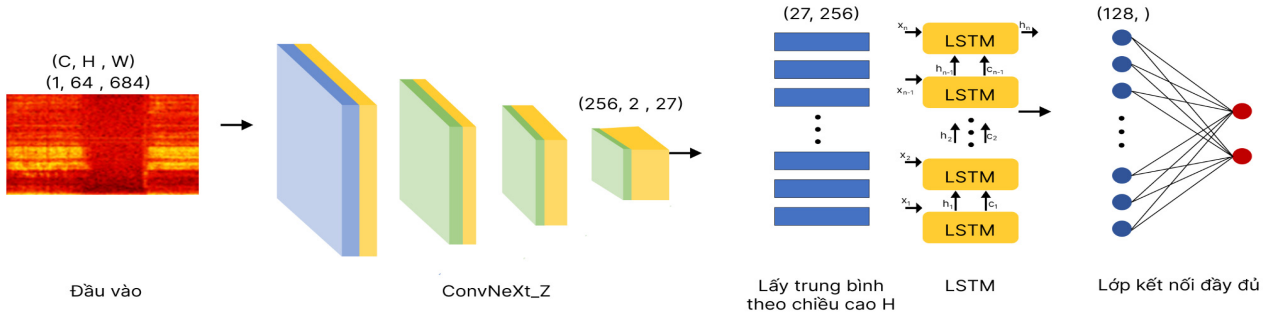
Mô hình ConvNeXt được đề xuất nhằm trích xuất đặc trưng cục bộ từ tín hiệu âm thanh. Đầu vào của mô hình là ma trận Mel-spectrum hai chiều đã được chuẩn hóa với kích thước (1, 64, 684). ConvNeXt_Z là một phiên bản rút gọn nhẹ của ConvNeXt được sử dụng làm backbone để trích xuất đặc trưng sâu, cho ra tensor kích thước (256, 2, 27). Cụ thể, Quy trình gồm: (1) Stem với convolution 4×4, stride 4 và LayerNorm giảm độ phân giải; (2) Downsample Layer với LayerNorm, convolution 2×2, stride 2, tăng kênh và giảm độ phân giải; (3) ConvNeXt Block với depthwise convolution 7×7, hai lớp pointwise convolution 1×1, kích hoạt GELU, layer scale, DropPath và residual connection. Đầu ra của backbone được flatten và qua một fully-connected layer tạo embedding 128 chiều, làm đầu vào cho lớp phân loại nhị phân. Cuối cùng, hàm SoftMax được áp dụng để dự đoán xác suất hai lớp [5].



Hình 2.2. Cấu trúc chi tiết của mô hình ConvNeXt

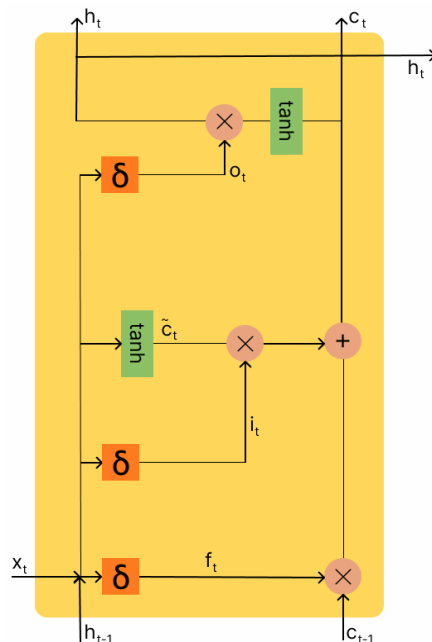
2.3.2. Mô hình ConvNeXt + LSTM

Mô hình ConvNeXt + LSTM được thiết kế để nhận diện các đặc trưng tạm thời trong âm thanh, giúp phát hiện ngưng thở – giảm thông khí. Đầu vào mô hình là phổ Mel hai chiều được chuẩn hóa với kích thước (64, 684). Tín hiệu này được đưa qua ConvNeXt_Z để trích xuất đặc trưng sâu, cho ra tensor có kích thước (256, 2, 27). Sau đó, đặc trưng được gộp trung bình theo chiều cao thành (256, 27), rồi hoán vị thành (27, 256) để phù hợp với mô hình tuần tự. Mạng LSTM một lớp (128 đơn vị ẩn) xử lý đầu vào để học thông tin thời gian. Trạng thái ẩn cuối cùng của LSTM được sử dụng làm đầu vào cho lớp kết nối hoàn toàn với đầu ra gồm 2 lớp. Cuối cùng, dùng SoftMax dự đoán xác suất có hay không có sự kiện ngưng thở – giảm thông khí trong đoạn âm thanh.



Hình 2.3. Cấu trúc tổng thể của mô hình ConvNeXt kết hợp LSTM

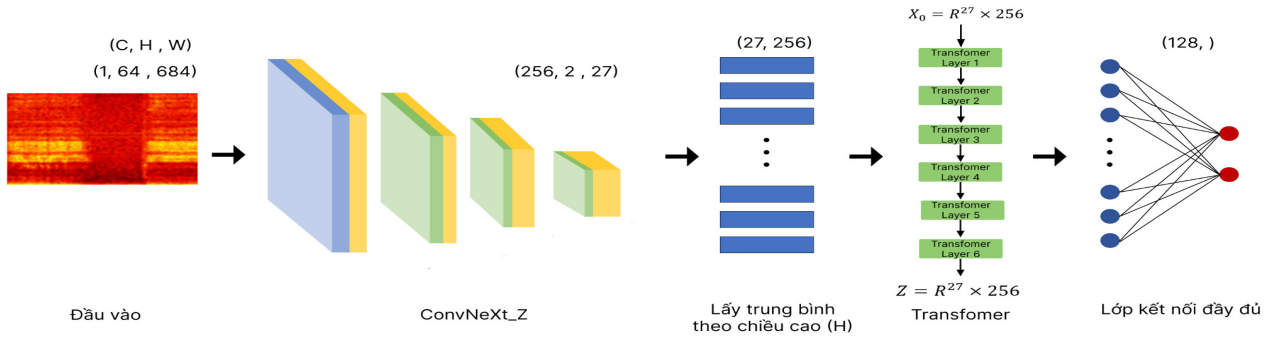
Cấu trúc của một ô LSTM bao gồm ba cổng chính: cổng quên f_t , cổng đầu vào i_t , và cổng đầu ra o_t . Các cổng này sử dụng hàm kích hoạt sigmoid (σ) và tanh để điều chỉnh quá trình cập nhật, duy trì hoặc loại bỏ thông tin trong trạng thái bộ nhớ c_t và trạng thái ẩn h_t . Trạng thái mới được truyền sang bước tiếp theo dựa trên dữ liệu đầu vào x_t cùng với các trạng thái trước đó h_{t-1} , c_{t-1} .



Hình 2.4. Cấu trúc chi tiết của một LSTM trong mô hình

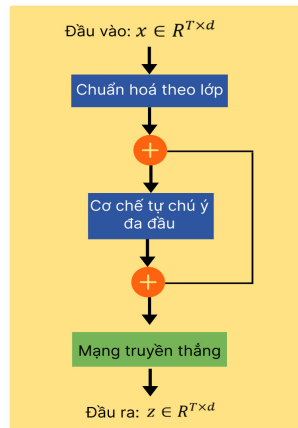
2.3.3. Mô hình ConvNeXt + Transformer

Mô hình ConvNeXt + Transformer được thiết kế nhằm phát hiện các sự kiện ngưng thở – giảm thông khí trong tín hiệu âm thanh. Đầu vào mô hình là phổ Mel hai chiều được chuẩn hóa với kích thước (64, 684). Tín hiệu này được đưa qua ConvNeXt_Z để trích xuất đặc trưng sâu, cho ra tensor có kích thước (256, 2, 27). Tiếp theo, áp dụng trung bình theo chiều cao, thu được ma trận đặc trưng (27, 256). Ma trận này đóng vai trò là đầu vào cho chuỗi 6 lớp Transformer, với đầu vào ban đầu $X_0 \in \mathbb{R}^{27 \times 256}$ và đầu ra cuối cùng $Z \in \mathbb{R}^{27 \times 256}$. Sau đó, kết quả được làm phẳng và đưa qua một lớp kết nối đầy đủ để tạo vector embedding 128 chiều, phục vụ cho bài toán phân loại nhị phân nhằm xác định có hay không sự kiện ngưng thở – giảm thông khí (Hình 2.5).



Hình 2.5. Cấu trúc tổng thể của mô hình ConvNeXt kết hợp Transformer

Một Transformer layer (Hình 2.6) gồm ba thành phần chính: (i) chuẩn hoá theo lớp, (ii) cơ chế tự chú ý đa đầu kèm residual connection, (iii) mạng truyền thẳng cũng kèm residual connection. Với cấu trúc này, đầu vào $x \in \mathbb{R}^{T \times d}$ được tuần tự biến đổi để cho ra đầu ra $z \in \mathbb{R}^{T \times d}$, giữ nguyên kích thước nhưng mang thông tin biểu diễn giàu đặc trưng hơn.



Hình 2.6. Cấu trúc chi tiết của một Transformer layer trong mô hình

Nghiên cứu trong bài báo “Screening for obstructive sleep apnea hypopnea using sleep breathing sounds based on the PSG-audio dataset” [3] mới chỉ dừng ở mô hình ConvNeXt + LSTM để trích xuất đặc trưng chuỗi từ tiếng thở. Trong nghiên cứu này, chúng tôi mở rộng với ConvNeXt + Transformer, thay thế LSTM bằng cơ chế Multi-Head Self-Attention. Ưu điểm kỳ vọng của hướng này là khai thác tốt hơn quan hệ dài hạn giữa các khung thời gian trong Mel-spectrogram, đồng thời giữ số tham số gọn nhẹ (~2.1 triệu), phù hợp cho thiết bị hạn chế tài nguyên

2.4. Ước lượng chỉ số AHI

Để ước lượng chỉ số AHI cho cả đêm, tín hiệu âm thanh được chia thành các đoạn dài 9 giây với mức chồng lấn 6 giây, sau đó chuyển đổi thành Mel-spectrogram với kích thước $1 \times 64 \times 684$ và đưa vào ConvNeXt_Z nhằm trích xuất đặc trưng ($256 \times 2 \times 27$). Với mô hình ConvNeXt + LSTM, đặc trưng sau gộp trung bình theo chiều tần số có dạng (27×256), được đưa vào LSTM một lớp 128 ẩn để học phụ thuộc dài hạn theo thời gian, sau đó qua fully-connected để dự đoán nhị phân sự kiện

ngưng thở – giảm thông khí: “có sự kiện” (label = 1) hoặc “không sự kiện” (label = 0). Với ConvNeXt + Transformer, đặc trưng sau gộp trung bình theo chiều tần số có dạng (27×256), qua 6 lớp Transformer (LayerNorm, Multi-Head Self-Attention, Feed-Forward, residual connection), flatten và fully-connected để dự đoán xác suất nhị phân tương tự [3].

3. KẾT QUẢ

3.1. Đánh giá hiệu suất phân loại trên tập dữ liệu không phụ thuộc chủ thể (Subject-Independent Dataset)

Hai mô hình ConvNeXt + LSTM và ConvNeXt + Transformer được huấn luyện và đánh giá trên tập validation và testing không trùng lặp đối tượng nhằm kiểm tra khả năng tổng quát hóa đối với dữ liệu mới. Bảng 1 trình bày chi tiết kết quả đánh giá dựa trên các chỉ số Accuracy và F1-score.

Bảng 1. So sánh hiệu suất hai mô hình trên tập dữ liệu không phụ thuộc chủ thể

Mô hình	Tham số (Triệu)	Tập dữ liệu	Accuracy (%)	F1-score (%)
ConvNeXt + LSTM	2.10	Validation	63.12	70.74
		Testing	64.75	71.63
ConvNeXt +Transformer	2.10	Validation	62.56	72.89
		Testing	65.99	74.10

So với kết quả trong bài báo “*Screening for obstructive sleep apnea hypopnea using sleep breathing sounds based on the PSG-audio dataset*” [3], đạt Accuracy và F1-score cao hơn trên tập dữ liệu không phụ thuộc chủ thể, cho thấy việc thay thế LSTM bằng Transformer đã cải thiện khả năng khái quát hóa – điểm yếu chính của phương pháp gốc.

Kết quả (Bảng 1) cho thấy cả hai mô hình đều có hiệu suất ổn định trên dữ liệu không trùng lặp chủ thể. ConvNeXt + Transformer đạt F1-score 72.89% trên validation so với 70.74% của ConvNeXt + LSTM, và 74.10% trên testing so với 71.63% của ConvNeXt + LSTM. Về Accuracy, ConvNeXt + LSTM cao hơn một chút trên validation (63.12% so với 62.56%), trong khi ConvNeXt + Transformer lại vượt trội trên testing (65.99% so với 64.75%).

Tổng hợp kết quả, ConvNeXt + Transformer vừa đảm bảo độ chính xác cao vừa đạt F1-score vượt trội, chứng tỏ ổn định và hiệu quả hơn trong môi trường thực tế có đa dạng chủ thể và nhãn không đồng đều, do đó là lựa chọn tối ưu cho dữ liệu chưa từng xuất hiện khi huấn luyện.

3.2. Hiệu suất ước lượng AHI

Để đánh giá khả năng mô hình trong việc ước lượng chỉ số AHI – một chỉ số quan trọng phản ánh mức độ nghiêm trọng của chứng ngưng thở khi ngủ – nghiên cứu sử dụng ba thước đo chính:

Hệ số tương quan Pearson (PCC): Đo mức độ tương quan tuyến tính giữa AHI thực tế và dự đoán, giá trị gần 1 cho thấy dự đoán chính xác và bám sát thực tế.

Sai số tuyệt đối trung bình (MAE): Đo chênh lệch trung bình tuyệt đối giữa giá trị dự đoán và thực tế, không phụ thuộc hướng sai số.

Căn bậc hai sai số bình phương trung bình (RMSE): nhấn mạnh các sai số lớn hơn bằng cách bình phương sai số trước khi lấy trung bình, do đó phản ánh mức độ ổn định và tin cậy của dự đoán.

Kết quả ước lượng trên tập dữ liệu không phụ thuộc chủ thể được trình bày trong Bảng 2.

Bảng 2. Hiệu suất ước lượng chỉ số AHI

	PCC	MAE (lần/giờ)	RMSE (lần/giờ)
Transformer_Independent	0.848	10.89	17.56
LSTM_Independent	0.824	11.07	17.87

Kết quả trong Bảng 2 cho thấy cả hai mô hình đều đạt hiệu suất tốt trên tập dữ liệu không phụ thuộc chủ thể, với PCC đều vượt mức 0.82. Trong đó, mô hình Transformer_Independent thể hiện hiệu suất nhỉnh hơn so với LSTM_Independent trên cả ba chỉ số đánh giá.

Mô hình Transformer_Independent vượt trội với $PCC = 0.848$, bám sát xu hướng thực tế hơn LSTM_Independent ($PCC = 0.824$). Transformer cũng đạt $MAE = 10.89$ (so với 11.07) và $RMSE = 17.56$ (so với 17.87 lần/giờ), thể hiện độ chính xác cao và kiểm soát sai số lớn tốt hơn.

Dù chênh lệch không lớn, Transformer_Independent vượt trội ở cả ba chỉ số (PCC , MAE , $RMSE$), thể hiện khả năng tổng quát hóa mạnh mẽ hơn LSTM_Independent, đặc biệt phù hợp với dữ liệu thực tế có sự đa dạng chủ thể.

4. KẾT LUẬN

Nghiên cứu đề xuất mô hình ConvNeXt + Transformer (~2.1M tham số) cho sàng lọc OSAHS từ tiếng thở ban đêm, vừa gọn nhẹ vừa hiệu quả, phù hợp cả trên thiết bị hạn chế tài nguyên. So với ConvNeXt + LSTM, mô hình không chỉ đạt F1-score và accuracy cao hơn mà còn cho kết quả ước lượng AHI chính xác, ổn định (PCC cao, MAE , $RMSE$ thấp). Điều này cho thấy Transformer có lợi thế rõ rệt trong việc khai thác quan hệ dài hạn và đặc trưng ngữ cảnh, giúp nâng cao khả năng khái quát hóa trên dữ liệu mới.

Đóng góp chính của nghiên cứu là mở rộng mô hình gốc ConvNeXt + LSTM bằng việc tích hợp Transformer, qua đó chứng minh rằng Transformer không chỉ có thể thay thế LSTM trong pipeline gọn nhẹ mà còn mang lại hiệu quả cao hơn trong phân loại tiếng thở và ước lượng AHI. Đây là một bước phát triển tiếp theo so với bài báo [3], mở ra hướng nghiên cứu kết hợp ConvNeXt và Transformer trong phân tích tín hiệu sinh lý.

Nghiên cứu đồng thời khẳng định tiềm năng ứng dụng của mô hình ConvNeXt + Transformer trong các hệ thống sàng lọc tự động ngoài môi trường lâm sàng. Tuy nhiên, hạn chế của nghiên cứu là mô hình chưa được kiểm chứng về độ ổn định trong điều kiện dữ liệu nhiễu và đa dạng môi trường (khác biệt thiết bị thu âm, tạp âm thực tế, sự khác biệt giữa các bệnh viện). Do đó, các nghiên cứu tương lai cần mở rộng thử nghiệm trên tập dữ liệu phong phú và thực tế hơn, nhằm nâng cao độ tin cậy và khả năng triển khai trong lâm sàng cũng như cộng đồng [6].

Ngoài ra, một hướng nghiên cứu tiếp theo là mở rộng so sánh với các mô hình học sâu khác ngoài ConvNeXt kết hợp LSTM và Transformer. Các kiến trúc như GRU, CNN thuần túy hoặc mô hình lai (CNN-GRU, CNN-BiLSTM) có thể mang lại kết quả bổ sung nhờ cách khai thác đặc trưng khác biệt. Việc đối chiếu các mô hình này sẽ giúp đánh giá toàn diện hơn, đảm bảo tính khách quan và hỗ trợ lựa chọn kiến trúc tối ưu cho phát hiện ngưng thở khi ngủ [7].

5. TÀI LIỆU THAM KHẢO

- [1] Bệnh viện Trung ương Huế. 2024. Ngưng thở khi ngủ (Obstructive Sleep Apnea - OSA) là sự rối loạn trong giấc ngủ. Truy cập: <https://www.facebook.com/bvtwhue/posts/5770896826289172>.
- [2] Benjafield, A. V., et al. 2019. Estimation of the global prevalence and burden of obstructive sleep apnoea: a literature-based analysis. *The Lancet Respiratory Medicine*, 7(8), 687-698.
- [3] Yujun Song, Li Ding, Jianxin Peng, Lijuan Song, Xiaowen Zhang. 2025. Screening for obstructive sleep apnea hypopnea using sleep breathing sounds based on the PSG-audio dataset. *Biomedical Signal Processing and Control*, 103, 107472.
- [4] Korompili, G., Pappa, M. D., Tzallas, A. T., Tsoukalas, G., & Fotiadis, D. I. (2021). PSG-Audio, a scored polysomnography dataset with simultaneous audio recordings for sleep apnea studies. *Scientific Data*, 8(1), 292. <https://doi.org/10.1038/s41597-021-00977-w>.
- [5] Liu, Z., Mao, H., Wu, C.-Y., Feichtenhofer, C., Darrell, T., & Xie, S. (2022). A ConvNet for the 2020s. arXiv preprint arXiv:2201.03545. <https://arxiv.org/pdf/2201.03545>.
- [6] A. Amiri, H. Mahdavi, and S. Shirmohammadi, "Automatic detection of obstructive sleep apnea using deep learning from snore sound," *Biomedical Signal Processing and Control*, vol. 68, p. 102666, 2021. <https://doi.org/10.1016/j.bspc.2021.102666>.
- [7] C. H. Chang, Y. H. Chen, and C. H. Lin, "Deep residual networks for screening obstructive sleep apnea from tracheal sound recordings," *Computers in Biology and Medicine*, vol. 155, p. 106672, 2023. <https://doi.org/10.1016/j.combiomed.2023.106672>.

ỨNG DỤNG HỌC SÂU TRONG DỰ BÁO CẤP ĐỘ MỰC NƯỚC ĐỂ CẢNH BÁO LŨ LỤT

Trần Thị Cẩm Giang, Phan Anh

Trường Đại học Thủy lợi, email: giangttc@tlu.edu.vn

1. GIỚI THIỆU CHUNG

Trước tình hình các hiện tượng thời tiết cực đoan như lũ lụt gia tăng cả về tần suất và cường độ gây thiệt hại nghiêm trọng, Việt Nam với hệ thống sông ngòi dày đặc, thường xuyên chịu ảnh hưởng nặng nề bởi lũ, đặc biệt ở Đồng bằng sông Cửu Long và miền Trung [1]. Dự báo mực nước sông chính xác và kịp thời trở thành yêu cầu cấp thiết nhằm hỗ trợ ứng phó hiệu quả. Tuy nhiên, các mô hình thủy lực truyền thống bị hạn chế bởi yêu cầu dữ liệu phức tạp và thời gian xử lý dài, đặc biệt trong điều kiện thiếu dữ liệu. Các mô hình như CNN, RNN, LSTM, Transformer [2][4] và mô hình lai RCNN đã cho thấy tiềm năng trong dự báo thủy văn [3]. Tuy nhiên, vẫn còn thách thức về dữ liệu nhiều, độ chính xác trong điều kiện thời tiết cực đoan và khả năng dự báo theo thời gian thực. Nghiên cứu này so sánh hiệu quả của các mô hình ML và DL trong dự báo mực nước sông Miami (Nam Florida) giai đoạn 2010–2020 nhằm xác định mô hình tối ưu cho cảnh báo lũ.

2. MÔ HÌNH BÀI TOÁN

Bài toán đặt ra là xây dựng và đánh giá hiệu quả các mô hình học máy và học sâu để dự đoán mực nước tại các vị trí hạ lưu sông Miami (trạm S1, S25A, S25B, S26) với các mốc thời gian dự báo 1 giờ, 7 giờ, 13 giờ, 19 giờ và 24 giờ trước khi lũ xảy ra. Mục tiêu là phát triển mô hình dự báo ngắn hạn chính xác, ổn định và đáp ứng yêu cầu thời gian thực, sử dụng dữ liệu mực nước, lưu lượng bơm, lượng mưa và thủy triều từ năm 2010–2020, bao gồm cả điều kiện thời tiết bình thường và thời tiết cực đoan (bão nhiệt đới Isaias, tháng 8/2020).

2.1. Khu vực nghiên cứu và tập dữ liệu

Khu vực nghiên cứu là hạ lưu sông Miami (Nam Florida), dài 5,6 dặm, với các trạm thủy văn S1, S4, S25A, S25B và S26. Bộ dữ liệu bao gồm mực nước, lưu lượng bơm, lượng mưa và thủy triều từ 1/1/2010 đến 31/12/2020, được chia thành tập huấn luyện (80%, 2010–2018) và tập kiểm tra (20%, 2018–2020). Dữ liệu thời tiết cực đoan từ bão Isaias (1–3/8/2020) được sử dụng để đánh giá hiệu suất trong điều kiện bất ổn. Bộ dữ liệu nghiên cứu chính được miêu tả trong bảng sau:

Bảng 1. Mô tả đặc trưng dữ liệu và cách sử dụng

Mô tả	Viết tắt	Đơn vị	Biến mục tiêu
Mực nước tại cống xả S26	Flow_S26	feet khối/giây	Không
Lưu lượng bơm tại S26	Pump_S26	feet khối/giây	Không
Mực nước hạ lưu S26	TWS_S26	Feet khối/giây	Không
Mực nước tại cống xả S25A	Flow_S25A	feet khối/giây	Không
Mực nước hạ lưu S25A	TWS_25A	feet	Có
Tổng lưu lượng cống S25B	Flow_S25B	feet khối/giây	Không
Lưu lượng bơm tại S25B	Pump_S25B	feet khối/giây	Không
Mực nước hạ lưu S25B	TWS_S25B	feet	Có
Mực nước tại trạm S1	WS_S1	feet	Có
Mực nước tại trạm S4	WS_S4	feet	Có
Cường độ mưa trung bình theo lưới	Grid_Rainfall	inch/giờ	Có

2.2. Tiền xử lý dữ liệu

Dữ liệu được chuẩn hóa theo phương pháp Min-max nhằm ngăn chặn hiện tượng gradient trở nên quá nhỏ hoặc quá lớn. Phương pháp này được thể hiện theo công thức:

$$x' = \frac{x - x_{\min}}{x_{\max} - x_{\min}} \quad (1)$$

trong đó: x' là giá trị sau khi chuẩn hóa, x là giá trị gốc và $x_{\min}$, $x_{\max}$ là giá trị nhỏ nhất và lớn nhất của tập dữ liệu.

2.3. Mô hình nghiên cứu

2.3.1. Ý tưởng và nguyên lý

Mạng nơ-ron hồi tiếp (RNN) là kiến trúc học sâu phù hợp để xử lý dữ liệu tuần tự, đặc biệt hiệu quả với các phụ thuộc ngắn hạn theo thời gian. Tuy nhiên, RNN gặp khó khăn trong việc học các quan hệ kéo dài do hạn chế bộ nhớ và hiện tượng mất dần gradient. Trong khi đó, mạng nơ-ron tích chập (CNN) có khả năng học các quan hệ không gian cục bộ, giúp rút ngắn chuỗi đầu vào và giảm áp lực lưu trữ cho RNN. Tận dụng sự bổ sung này, kiến trúc kết hợp RCNN (Recurrent-Convolutional Neural Network) đã được đề xuất, trong đó RNN trích xuất đặc trưng thời gian, còn CNN tiếp tục khai thác các mối quan hệ không gian. Các lớp gộp sau CNN giúp chọn lọc đặc trưng nổi bật, giảm chiều dữ liệu và tăng khả năng tổng quát hóa. Với cơ chế kết hợp thời gian – không gian, RCNN cho thấy tiềm năng cao trong các bài toán chuỗi thời gian dài hạn [3].

2.3.2. Kiến trúc RCNN

Tầng tích chập lặp là khối mô-đun cốt lõi của RCNN. Các trạng thái của các đơn vị trong tầng tích chập lặp phát triển qua các bước thời gian rời rạc. Hoạt động của một đơn vị tại vị trí (i, j) trên bản đồ đặc trưng k tại thời điểm t được xác định thông qua 2 bước:

Bước đầu tính toán đầu vào $z_{ijk}(t)$, $z_{ijk}(t)$ được tính bằng tổng hợp của hai thành phần chính, đầu vào truyền thẳng và đầu vào lặp lại:

$$z_{ijk}(t) = (w_k^f)^T u^{(i,j)}(t) + (w_k^r)^T x^{(i,j)}(t-1) + b_k \quad (2)$$

trong đó:

$u^{(i,j)}(t)$: đại diện cho đầu vào truyền thẳng từ các bản đồ đặc trưng của lớp trước.

$x^{(i,j)}(t-1)$: đầu vào lặp từ lớp hiện tại tại thời điểm trước $(t-1)$.

(w_k^f) , (w_k^r) : trọng số feed-forward và lặp.

b_k : hệ số bias.

Sau khi tính toán đầu vào $z_{ijk}(t)$, hoạt động (hay trạng thái) $x_{ijk}(t)$ của đơn vị được xác định bằng cách áp dụng một chuỗi các hàm phi tuyến tính và chuẩn hóa.

$$x_{ijk}(t) = g(f(z_{ijk}(t))) \quad (3)$$

$$f(z_{ijk}(t)) = \max(z_{ijk}(t), 0) \quad (4)$$

trong đó:

f : là hàm kích hoạt tuyến tính chỉnh lưu (ReLU). Hàm này giúp giải quyết vấn đề gradient biến mất và thường được định nghĩa là:

$$f(z_{ijk}(t)) = \max(z_{ijk}(t), 0) \quad (5)$$

g : là hàm chuẩn hóa phản ứng cục bộ.

Kiến trúc trong mã

RCNN bao gồm:

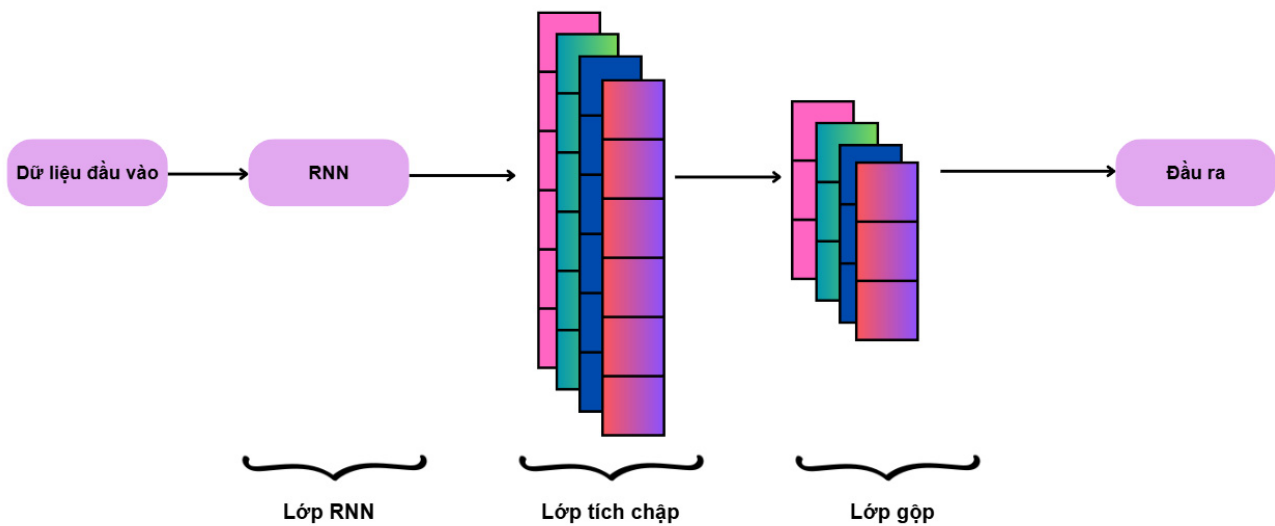
- Một tầng tích chập đầu tiên.
- Bốn tầng RCL, kết hợp kết nối feed-forward và lặp.
- Ba tầng max-pooling để giảm kích thước không gian.

Khi triển khai qua (T) bước thời gian, RCNN tạo ra mạng feed-forward với độ sâu tối đa (T+1), nhưng giữ số lượng tham số cố định.

Cấu trúc đa đường dẫn

RCNN tạo ra nhiều đường dẫn từ đầu vào đến đầu ra với độ dài khác nhau (từ 1 đến (T+1)). Điều này giúp:

- Giảm vấn đề biến mất gradient (vanishing gradient).
- Tăng khả năng học các đặc trưng phức tạp mà vẫn duy trì tổng quát hóa.



Hình 2. Mô hình hoạt động của RCNN

RCNN mang lại nhiều ưu điểm nổi bật trong dự báo mực nước, nhờ khả năng tích hợp ngữ cảnh tốt hơn khi sử dụng kết nối lặp để xử lý thông tin trong cùng tầng, khắc phục hạn chế của CNN truyền thống vốn chỉ xử lý ngữ cảnh ở tầng cao hơn. RCNN vừa tận dụng khả năng nắm bắt phụ thuộc dài hạn của RNN, vừa trích chọn đặc trưng cục bộ hiệu quả nhờ CNN, giúp mô hình thích hợp với dữ liệu có tính tuần hoàn và chịu tác động ngoại sinh phức tạp. So với các kiến trúc Transformer vốn mạnh về chuỗi dài nhưng đòi hỏi tập dữ liệu lớn và chi phí tính toán cao, RCNN cho thấy sự cân bằng giữa độ chính xác và tính khả thi trong điều kiện dữ liệu và hạ tầng thực tế còn hạn chế. Hiệu suất của RCNN được chứng minh qua kết quả dự báo chính xác trên các tập dữ liệu phức tạp.

3. KẾT QUẢ NGHIÊN CỨU

Bài báo đã sử dụng các độ đo MAE, RMSE để đánh giá hiệu suất mô hình RCNN với các thuật toán khác như Linear Regression (LR), Random Forest (RF), XGBoost (XGB), Multi-Layer Perceptron (MLP), Convolutional Neural Network (CNN), Recurrent Neural Network (RNN), Long Short-Term Memory (LSTM).

3.1. Phân tích hiệu suất đầu ra theo thời điểm dự đoán

Bảng 2. Sai số tuyệt đối trung bình (MAE) của các mô hình DL của tất cả các vị trí (S1, 25A, S25B, S26) tại các bước thời gian dự đoán khác nhau (giờ)

	1 giờ	7 giờ	13 giờ	19 giờ	Tất cả
RCNN	0.05536	0.05986	0.05619	0.06071	0.06298
RNN	0.08859	0.10208	0.09741	0.10734	0.10073
CNN	0.06595	0.06226	0.06804	0.06379	0.07254
LR	0.08094	0.08205	0.09149	0.08489	0.1011
LSTM	0.14473	0.14631	0.16057	0.15769	0.15381
MLP	0.13053	0.12095	0.12384	0.11263	0.116
RF	0.11328	0.12314	0.11199	0.1251	0.11939
XGBOOST	0.21802	0.19117	0.21916	0.19107	0.18506

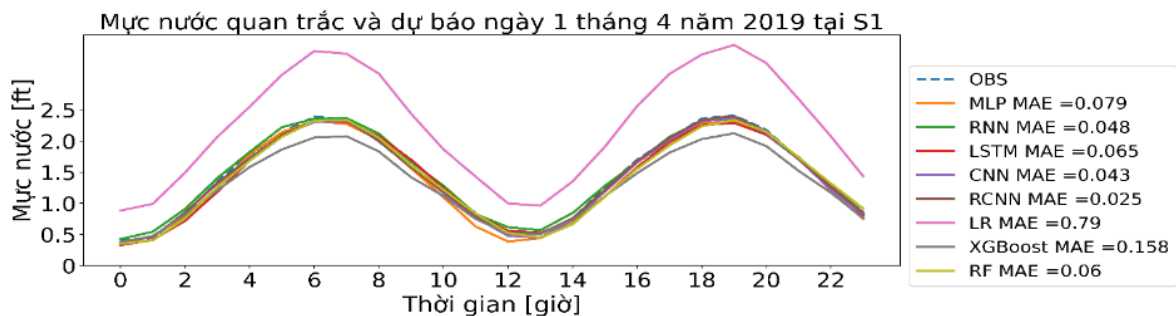
Bảng 3 Sai số bình phương trung bình (RMSE) của các mô hình DL của tất cả các vị trí (S1, 25A, S25B, S26) tại các bước thời gian dự đoán khác nhau (giờ)

	1 giờ	7 giờ	13 giờ	19 giờ	Tất cả
RCNN	0.07993	0.08492	0.08109	0.08708	0.08889
RNN	0.11918	0.139	0.12862	0.1462	0.0973
CNN	0.08784	0.08742	0.09097	0.0899	0.13259
LR	0.10537	0.11376	0.11551	0.11659	0.13672
LSTM	0.1889	0.19189	0.20395	0.20534	0.14902
MLP	0.1637	0.15599	0.15664	0.14545	0.15718
RF	0.14655	0.16314	0.14462	0.16614	0.1996
XGBOOST	0.30889	0.28521	0.31022	0.28527	0.2685

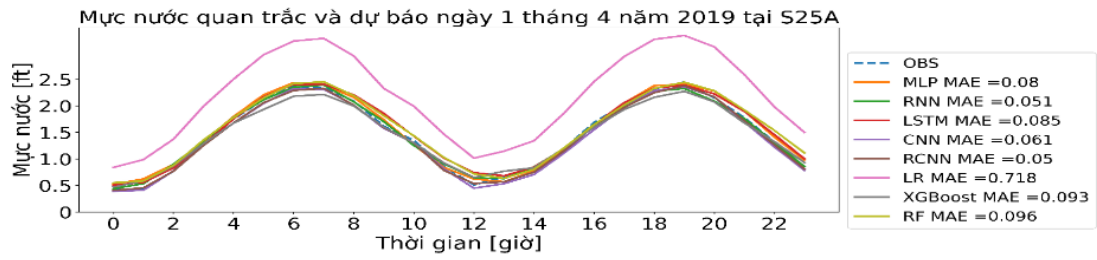
Ở trong bảng 2 và bảng 3, cột ‘Tất cả’ biểu thị giá trị trung bình (mean) của chỉ số (MAE/RMSE) tổng hợp qua các mốc thời gian dự báo 1h, 8h, 16h, 24h và qua các vị trí S1, S25A, S25B, S26. RCNN là mô hình hiệu quả nhất với MAE và RMSE thấp nhất, cho thấy độ chính xác và ổn định cao ở cả dự đoán ngắn và dài hạn. RNN phù hợp với dự báo ngắn hạn, CNN và MLP cho kết quả ổn định ở trung và dài hạn. LSTM không vượt trội như kỳ vọng và cho sai số cao hơn RNN. Các mô hình học máy có hiệu suất kém, không phù hợp với bài toán dự báo mực nước nhiều bước thời gian.

3.2. Trực quan hóa kết quả

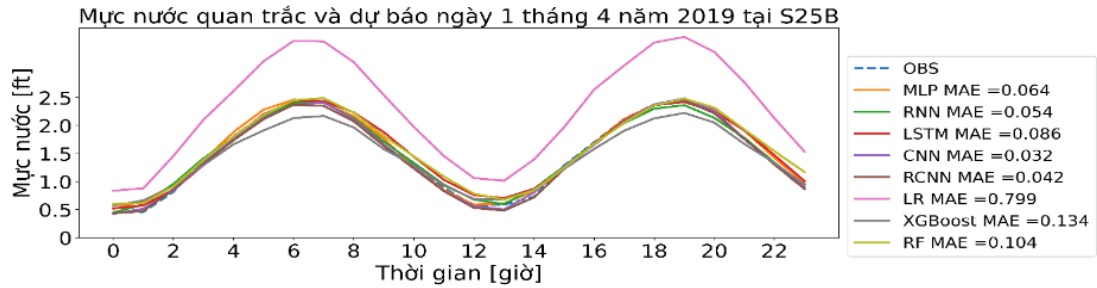
3.2.1. Giai đoạn bình thường



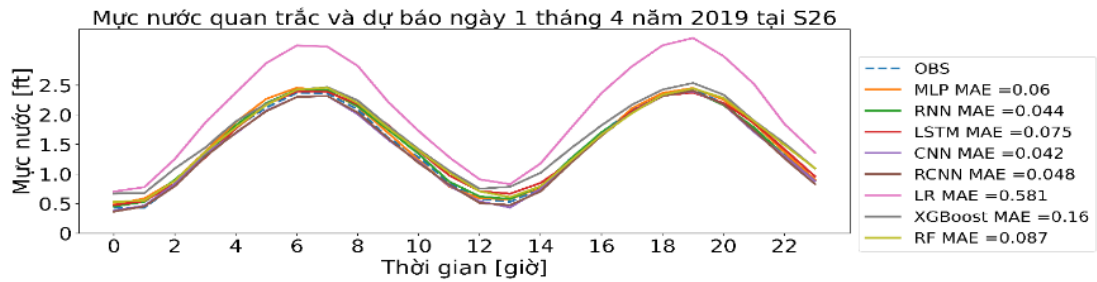
a) Mực nước tại trạm S1



b) Mức nước hạ lưu tại trạm S25A



c) Mức nước hạ lưu tại trạm S25B

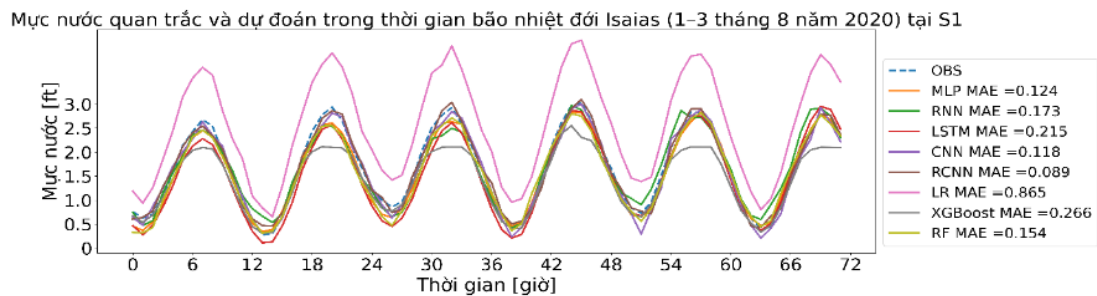


d) Mức nước hạ lưu tại trạm S26

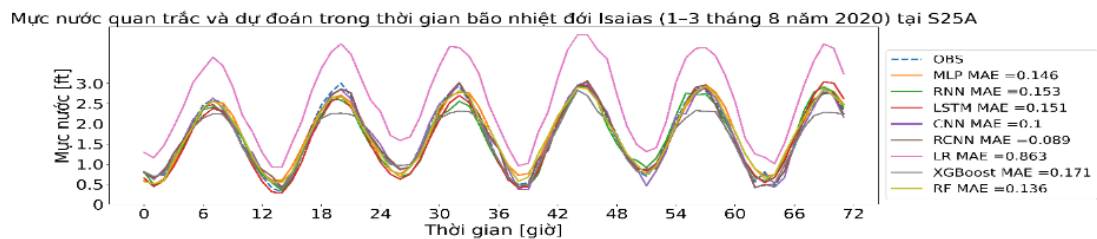
Hình 3. Mức nước quan trắc và dự đoán tại 4 vị trí vào ngày 1 tháng 4 năm 2019 gồm: (a), (b), (c), (d)

RCNN là mô hình dự báo hiệu quả nhất tại cả bốn trạm, với sai số MAE thấp và độ chính xác cao, đặc biệt nổi bật tại các trạm S1 và S25A. Nhờ kết hợp giữa CNN và RNN, RCNN khai thác tốt đặc trưng không gian–thời gian trong dữ liệu mực nước, giúp mô hình hóa biến động phức tạp hiệu quả.

3.2.2. Giai đoạn có bão

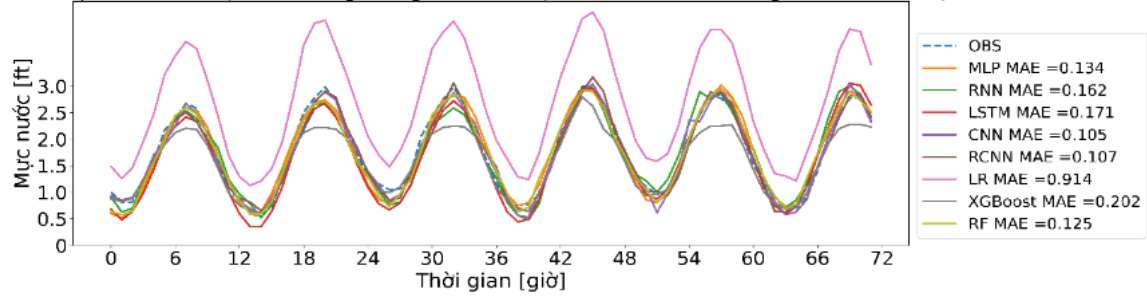


a) Mức nước tại trạm S1



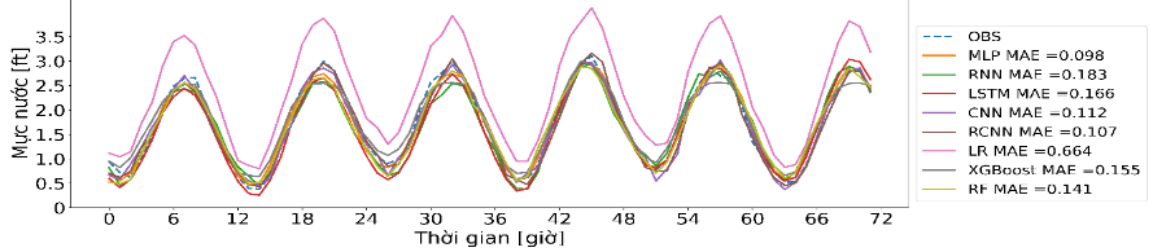
b) Mức nước hạ lưu tại trạm S25A

Mức nước quan trắc và dự đoán trong thời gian bão nhiệt đới Isaias (1-3 tháng 8 năm 2020) tại S25B



c) Mức nước hạ lưu tại trạm S25B

Mức nước quan trắc và dự đoán trong thời gian bão nhiệt đới Isaias (1-3 tháng 8 năm 2020) tại S26



d) Mức nước hạ lưu tại trạm S26

Hình 4. Mức nước quan trắc và dự đoán vào ngày 3-8 tháng 4 năm 2019 tại 4 vị trí: (a), (b), (c), (d)

RCNN cho thấy hiệu suất vượt trội khi đạt MAE thấp nhất tại S1 và S25A (0.089) và duy trì sai số nhỏ tại S25B và S26 (0.107).

4. KẾT LUẬN

Kết quả thực nghiệm cho thấy mô hình RCNN vượt trội so với các mô hình khác về độ chính xác và tính ổn định, đạt sai số tuyệt đối trung bình (MAE) thấp nhất (0.0629) và sai số bình phương trung bình (RMSE) thấp nhất (0.0888) trên tất cả các trạm và mốc thời gian dự báo. Hiệu quả của RCNN được khẳng định trong điều kiện phức tạp: tận dụng hiệu quả khả năng trích xuất đặc trưng không gian của CNN và đặc trưng thời gian của RNN, giúp cải thiện đáng kể hiệu suất dự báo, đặc biệt trong các điều kiện thời tiết phức tạp như bão nhiệt đới Isaias (tháng 8/2020). Các mô hình học sâu khác như RNN và CNN cũng cho kết quả tốt, nhưng không ổn định bằng RCNN. Dù vậy, RCNN vẫn tồn tại một số hạn chế như chi phí huấn luyện lớn, yêu cầu dữ liệu lịch sử dày và dễ suy giảm hiệu năng khi dữ liệu đầu vào thiếu hoặc nhiễu. Các nghiên cứu tiếp theo sẽ hướng tới tối ưu kiến trúc để giảm chi phí tính toán, tích hợp mô hình lai với kiến thức thủy văn – thủy lực để tăng khả năng giải thích, và mở rộng so sánh với các kiến trúc tiên tiến nhằm đánh giá toàn diện hơn.

5. TÀI LIỆU THAM KHẢO

- [1] H. L.Bui, 2005, "LŨ LỤT Ở NƯỚC TA VÀ ĐỊNH ĐỊNH HƯỚNG CÁC GIẢI PHÁP PHÒNG TRÁNH," *Tạp chí Khí tượng Thủy văn*.
- [2] Bekir Z. Demiray, 2024, "Transformer models for 120-hour streamflow prediction", Preprint.
- [3] e. a. M Herath, 2023, "Deep Machine Learning-Based Water Level Prediction Model for Colombo Flood Detention Area," *Applied Sciences*, vol. 13.4, p. 2194.
- [4] H. X. M Liang, 2015, "Recurrent convolutional neural network for object recognition.," *Proceedings of the IEEE conference on computer vision and pattern recognition*.

TÍNH LỖI, LỖM CỦA ĐỒ THỊ HÀM SỐ VÀ BẤT ĐẲNG THỨC

Nguyễn Thị Lý, Nguyễn Hữu Thọ
 Trường Đại học Thủy lợi, email: lycs2@tlu.edu.vn

1. GIỚI THIỆU CHUNG

Đây là một báo cáo tổng quan về một chủ đề mà chúng tôi đã từng sử dụng trong quá trình ôn luyện trong nhiều năm cho đội tuyển Olympic Toán học của Trường Đại học Thủy lợi tham dự các kỳ thi Olympic Toán học sinh viên Toàn quốc. Từ quá trình dạy học, chúng tôi đúc kết thành một báo cáo mang tính tổng quan để có thể áp dụng và phát triển tiếp trong những năm tiếp theo.

2. NỘI DUNG BÁO CÁO

A. Cơ sở lý thuyết

a. Định nghĩa ([3]). Cho hàm số $y = f(x)$ liên tục trên đoạn $[a; b]$ có đồ thị là (C) , xét hai điểm $(a; f(a))$ và $(b; f(b))$ nằm trên đồ thị (C) . Khi đó

i) Đồ thị (C) được gọi là lồi trên $[a; b]$ nếu tiếp tuyến tại mọi điểm thuộc cung AB luôn nằm phía trên đồ thị (C) .

ii) Đồ thị (C) được gọi là lõm trên $[a; b]$ nếu tiếp tuyến tại mọi điểm thuộc cung AB luôn nằm phía dưới đồ thị (C) .

b. Dấu hiệu đồ thị lồi, lõm

Định lý 2.1 ([2]). Cho hàm số $y = f(x)$ thuộc $C^2(a; b)$, khi đó

i) Nếu $f''(x) > 0$ với mọi $x \in (a; b)$ thì đồ thị của hàm số là lõm trên $(a; b)$.

ii) Nếu $f''(x) < 0$ với mọi $x \in (a; b)$ thì đồ thị của hàm số là lồi trên $(a; b)$.

c. Một số bất đẳng thức cơ bản

Định lý 2.2 ([1]). (Bất đẳng thức tiếp tuyến)

Cho hàm số $y = f(x)$ liên tục và có đạo hàm đến cấp hai trên đoạn $[a; b]$, khi đó

i) Nếu $f''(x) \geq 0$ với mọi $x \in [a; b]$ thì ta có $f(x) \geq f'(x_0)(x - x_0) + f(x_0), \forall x_0 \in [a; b]$.

ii) Nếu $f''(x) \leq 0$ với mọi $x \in [a; b]$ thì ta có $f(x) \leq f'(x_0)(x - x_0) + f(x_0), \forall x_0 \in [a; b]$.

Dấu "=" đạt được khi và chỉ khi $x = x_0$.

Định lý 2.3 ([1]). (Bất đẳng thức cát tuyến)

Cho hàm số $y = f(x)$ liên tục và có đạo hàm đến cấp hai trên đoạn $[a; b]$, khi đó

i) Nếu $f''(x) \geq 0$ với mọi $x \in [a; b]$ thì ta có $f(x) \geq \frac{f(a) - f(b)}{a - b}(x - a) + f(a), \forall x \in [a; b]$.

ii) Nếu $f''(x) \leq 0$ với mọi $x \in [a; b]$ thì ta có $f(x) \leq \frac{f(a) - f(b)}{a - b}(x - a) + f(a), \forall x \in [a; b]$.

Dấu "=" đạt được khi và chỉ khi $x = a$ hoặc $x = b$.

B. Một số ví dụ áp dụng

Ví dụ 1. Cho các số thực dương a, b, c thỏa mãn $a + b + c = 1$. Hãy chứng minh rằng

$$\frac{a}{\sqrt{a^2+1}} + \frac{b}{\sqrt{b^2+1}} + \frac{c}{\sqrt{c^2+1}} \leq \frac{3\sqrt{10}}{10}.$$

Giải. Từ giả thiết suy ra $a, b, c \in (0; 1)$, do tính đối xứng của a, b, c nên dễ thấy tại $a = b = c = \frac{1}{3}$

thì ta nhận được dấu " $=$ ". Xét hàm số $f(x) = \frac{x}{\sqrt{x^2+1}}$ với $x \in (0; 1)$, khi đó ta có

$$f'(x) = \frac{1}{\sqrt{(x^2+1)^3}} \text{ và } f''(x) = -\frac{3x}{\sqrt{(x^2+1)^5}}.$$

Dễ thấy $f''(x) = -\frac{3x}{\sqrt{(x^2+1)^5}} < 0 \quad \forall x \in (0; 1)$. Áp dụng bất đẳng thức tiếp tuyến lần lượt tại

$x = a, b, c$ và $x_0 = \frac{1}{3}$ ta nhận được

$$f(a) \leq f'\left(\frac{1}{3}\right)\left(a - \frac{1}{3}\right) + f\left(\frac{1}{3}\right),$$

$$f(b) \leq f'\left(\frac{1}{3}\right)\left(b - \frac{1}{3}\right) + f\left(\frac{1}{3}\right),$$

$$f(c) \leq f'\left(\frac{1}{3}\right)\left(c - \frac{1}{3}\right) + f\left(\frac{1}{3}\right).$$

Cộng hai vế của các bất đẳng thức trên, khi đó

$$f(a) + f(b) + f(c) \leq f'\left(\frac{1}{3}\right)(a + b + c - 1) + 3f\left(\frac{1}{3}\right) = \frac{3}{\sqrt{10}} = \frac{3\sqrt{10}}{10}.$$

Đẳng thức xảy ra khi và chỉ khi $a = b = c = \frac{1}{3}$.

Ví dụ 2. Cho các số thực không âm a, b, c thỏa mãn $a^2 + b^2 + c^2 = 3$. Hãy chứng minh rằng

$$\frac{1}{\sqrt{1+8a}} + \frac{1}{\sqrt{1+8b}} + \frac{1}{\sqrt{1+8c}} \geq 1.$$

Giải. Từ giả thiết suy ra $a, b, c \in [0; \sqrt{3}]$, do tính đối xứng nên dễ thấy tại $a = b = c = 1$ thì ta nhận được dấu " $=$ ". Xét hàm số $f(x) = \frac{1}{\sqrt{1+8x}}$ với $x \in [0; \sqrt{3}]$, khi đó ta có

$$f'(x) = -\frac{4}{\sqrt{(1+8x)^3}} \text{ và } f''(x) = \frac{48}{\sqrt{(1+8x)^5}}.$$

Dễ thấy $f''(x) > 0 \quad \forall x \in [0; \sqrt{3}]$. Áp dụng bất đẳng thức tiếp tuyến lần lượt tại $x = a, b, c$ và $x_0 = 1$ ta nhận được

$$f(a) \geq f'(1)(a - 1) + f(1),$$

$$f(b) \geq f'(1)(b - 1) + f(1),$$

$$f(c) \geq f'(1)(c - 1) + f(1).$$

Cộng hai vế của các bất đẳng thức trên, khi đó

$$f(a) + f(b) + f(c) \geq f'(1)(a+b+c-3) + 3f(1).$$

Mặt khác, theo bất đẳng thức Bunhiacopski ta có $(a+b+c)^2 \leq 3(a^2+b^2+c^2) = 9$, suy ra

$$a+b+c \leq 3 \text{ hay là } a+b+c-3 \leq 0, \text{ hơn nữa } f'(1) = -\frac{4}{27} < 0;$$

Như vậy $f(a) + f(b) + f(c) \geq f'(1)(a+b+c-3) + 3f(1) \geq 3f(1) = 1$. Đẳng thức xảy ra khi và chỉ khi $a=b=c=1$.

Chú ý. i) Phương pháp được sử dụng trong hai ví dụ trên có thể mở rộng cho việc chứng minh các bất đẳng thức dạng $f(a_1) + f(a_2) + \dots + f(a_n) \geq k$ hoặc $f(a_1) + f(a_2) + \dots + f(a_n) \leq k$ với $a_i (i=1,2,\dots,n)$ là các số thực cho trước.

ii) Nếu bất đẳng thức có dạng $f(a_1) \cdot f(a_2) \cdots f(a_n) \geq k$ với $f(a_i) > 0 (i=1,2,\dots,n); k > 0$ thì ta có thể lấy loga cơ số tự nhiên hai vế rồi sau đó sử dụng phương pháp như trên.

Ví dụ 3. Cho các số thực dương a, b, c thỏa mãn $a+b+c=3$. Tìm giá trị lớn nhất của biểu thức

$$P = \left(a + \sqrt{1+a^2}\right)^b \left(b + \sqrt{1+b^2}\right)^c \left(c + \sqrt{1+c^2}\right)^a.$$

Giải. Bằng cách lấy loga cơ số tự nhiên ta có

$$\ln(P) = b \ln\left(a + \sqrt{1+a^2}\right) + c \ln\left(b + \sqrt{1+b^2}\right) + a \ln\left(c + \sqrt{1+c^2}\right).$$

Từ giả thiết suy ra $a, b, c \in (0;1)$. Xét hàm số $f(x) = \ln\left(x + \sqrt{1+x^2}\right)$ với $x \in (0;1)$, khi đó ta có

$$f'(x) = \frac{1}{\sqrt{x^2+1}} \text{ và } f''(x) = \frac{-x}{\sqrt{(1+x^2)^3}}.$$

Để thấy $f''(x) < 0 \forall x \in (0;1)$. Áp dụng bất đẳng thức tiếp tuyến lần lượt tại $x=a$ và $x_0=1$

$$f(a) \leq f'(1)(a-1) + f(1) = f'(1)a + f(1) - f'(1).$$

Nhân hai vế của bất đẳng thức trên với b nhận được

$$bf(a) \leq f'(1)ab + [f(1) - f'(1)]b,$$

tương tự ta cũng có:

$$cf(b) \leq f'(1)cb + [f(1) - f'(1)]c,$$

$$af(c) \leq f'(1)ac + [f(1) - f'(1)]a.$$

Cộng hai vế của các bất đẳng thức trên, khi đó:

$$\ln(P) \leq f'(1)(ab+bc+ca - (a+b+c)) + f(1)(a+b+c).$$

Mặt khác do $a+b+c=3 \geq ab+bc+ca$ nên

$$\ln(P) \leq f'(1)(ab+bc+ca - (a+b+c)) + f(1)(a+b+c) \leq 3 \ln(1+\sqrt{2})$$

$$\ln(P) \leq \ln(1+\sqrt{2})^3.$$

Và từ đó $P \leq (1+\sqrt{2})^3$. Đẳng thức xảy ra khi và chỉ khi $a=b=c=1$.

Ví dụ 4. Cho các số thực dương a, b, c thỏa mãn $a + b + c = 1$. Chứng minh rằng

$$10(a^3 + b^3 + c^3) - 9(a^5 + b^5 + c^5) \geq 1.$$

Giải. Từ giả thiết suy ra $a, b, c \in (0; 1)$, do tính đối xứng nên dễ thấy tại $a = b = c = \frac{1}{3}$ thì ta nhận được dấu " $=$ ". Không mất tính tổng quát giả sử rằng $a \geq b \geq c$.

Xét hàm số $f(x) = 10x^3 - 9x^5$, $x \in (0; 1)$, khi đó ta có

$$f'(x) = 30x^2 - 45x^4 \text{ và } f''(x) = 60x - 180x^3.$$

Dễ thấy $f''(x) = 0 \Leftrightarrow x = x_0 = \sqrt{\frac{1}{3}}$. Từ đó

$$f''(x) > 0 \text{ với } x \in (0; x_0) \text{ và } f''(x) < 0 \text{ với } x \in (x_0; 1).$$

Nếu $0 < a < x_0$: Áp dụng bất đẳng thức tiếp tuyến lần lượt tại $x = a, b, c$ và $x_0 = \frac{1}{3}$ ta nhận được

$$f(a) \geq f'\left(\frac{1}{3}\right)\left(a - \frac{1}{3}\right) + f\left(\frac{1}{3}\right),$$

$$f(b) \geq f'\left(\frac{1}{3}\right)\left(b - \frac{1}{3}\right) + f\left(\frac{1}{3}\right),$$

$$f(c) \geq f'\left(\frac{1}{3}\right)\left(c - \frac{1}{3}\right) + f\left(\frac{1}{3}\right).$$

Cộng hai vế của các bất đẳng thức trên, khi đó:

$$f(a) + f(b) + f(c) \geq f'\left(\frac{1}{3}\right)(a + b + c - 1) + 3f\left(\frac{1}{3}\right) = 1.$$

Nếu $x_0 < a < 1$: Áp dụng bất đẳng thức cát tuyến ta có

$$f(a) \geq \frac{f(1) - f(x_0)}{1 - x_0}(a - 1) + f(1) > f(1) = 1.$$

Ta lại áp dụng bất đẳng thức tiếp tuyến lần lượt tại $x = b, c$ và $x_0 = 0$ ta nhận được

$$f(b) \geq f'(0)(b - 0) + f(0) = 0,$$

$$f(c) \geq f'(0)(c - 0) + f(0) = 0.$$

Và từ đó $f(a) + f(b) + f(c) > 1$, kết hợp lại ta nhận được điều cần chứng minh.

Ví dụ 5. Cho tam giác nhọn ABC . Tìm giá trị lớn nhất của biểu thức

$$F = \sin A \cdot \sin^2 B \cdot \sin^3 C.$$

Giải. Ta có $\ln(F) = \ln(\sin A) + 2\ln(\sin B) + 3\ln(\sin C)$. Xét hàm số

$$f(x) = \ln(\sin x), x \in \left(0; \frac{\pi}{2}\right)$$

khi đó ta có

$$f'(x) = \cot x \text{ và } f''(x) = -\frac{1}{\sin^2 x}.$$

Dễ thấy $f''(x) < 0 \forall x \in \left(0; \frac{\pi}{2}\right)$. Áp dụng bất đẳng thức tiếp tuyến với tam giác nhọn MNP ta nhận được

$$\begin{aligned}f(A) &\leq f'(M)(A-M) + f(M) = (A-M)\cot M + \ln(\sin M), \\f(B) &\geq f'(N)(B-N) + f(N) = (B-N)\cot N + \ln(\sin N) \\f(C) &\geq f'(P)(C-P) + f(P) = (C-P)\cot P + \ln(\sin P).\end{aligned}$$

Cộng hai vế của các bất đẳng thức trên, khi đó

$$\tan M \cdot f(A) + \tan N \cdot f(B) + \tan P \cdot f(C) \geq \tan M \cdot \ln(\sin M) + \tan N \cdot \ln(\sin N) + \tan P \cdot \ln(\sin P).$$

Ta chọn tam giác nhọn MNP sao cho: $\frac{\tan M}{1} = \frac{\tan N}{2} = \frac{\tan P}{3} = k, (k > 0)$, suy ra

$$\begin{cases} \tan M = k \\ \tan N = 2k \\ \tan P = 3k. \end{cases}$$

Mặt khác lại có $\tan M + \tan N + \tan P = \tan M \cdot \tan N \cdot \tan P$, vậy nên $6k = 6k^3$ hay là $k = 1$ khi đó

$$\begin{cases} \sin M = \frac{\tan M}{\sqrt{1 + \tan^2 M}} = \frac{1}{\sqrt{2}} \\ \sin N = \frac{2}{\sqrt{5}} \\ \sin P = \frac{3}{\sqrt{10}}. \end{cases}$$

Và do đó

$$f(A) + f(B) + f(C) \leq \ln \frac{1}{\sqrt{2}} + 2 \ln \frac{2}{\sqrt{5}} + 3 \ln \frac{3}{\sqrt{10}} = \ln \frac{27}{25\sqrt{10}},$$

tức là $F \leq \frac{27}{25\sqrt{10}}$, đẳng thức xảy ra khi và chỉ khi $A = M; B = N; C = P$. Vậy $\max F = \frac{27}{25\sqrt{10}}$.

C. Một số bài tập tự luyện

1. Cho các số thực dương a, b, c thỏa mãn $a + b + c = 1$. Chứng minh rằng

$$\frac{a}{1+bc} + \frac{b}{1+ca} + \frac{c}{1+ab} \geq \frac{9}{10}.$$

2. Cho các số thực dương a, b, c . Chứng minh rằng

$$\frac{(b+c-a)}{(b+c)^2 + a^2} + \frac{(c+a-b)}{(c+a)^2 + b^2} + \frac{(a+b-c)}{(a+b)^2 + c^2} \geq \frac{3}{5}.$$

3. Cho tam giác nhọn ABC . Tìm giá trị nhỏ nhất của biểu thức

$$F = \tan A + 2 \tan B + 3 \tan C.$$

4. Cho các số thực dương a, b, c thỏa mãn $a + b + c = 1$. Tìm giá trị nhỏ nhất của biểu thức

$$P = a^3 + \sqrt{1+b^2} + \sqrt[4]{1+c^4}.$$

3. KẾT LUẬN

Báo cáo này là một tổng quan đạt được sau một quá trình lâu dài giảng dạy ôn luyện cho đội tuyển Olympic Toán học của Trường tham dự các kỳ thi Olympic Toán học sinh viên Toàn quốc. Trong báo cáo, chúng tôi sử dụng công cụ tính lời/lỗ của hàm số trong việc chứng minh các bài toán về bất đẳng thức, phương pháp này có ưu điểm: dễ chọn được hàm đặc trưng phù hợp và sau đó các bước tiến hành thường tuân theo lộ trình rõ ràng, dễ hiểu.

4. TÀI LIỆU THAM KHẢO

- [1] R. Castillo, H. Rafeiro. 2016. Convex Functions and Inequalities. CMS Books in Mathematics. Springer, Cham.
- [2] Nguyễn Văn Mậu, Lê Ngọc Lăng, Nguyễn Minh Tuấn. 2006. Các đề thi Olympic Toán sinh viên Toàn quốc, NXB. Giáo dục, 226 trang.
- [3] R. Tyrrell Rockafellar. 1970. Convex Analysis. Princeton University Press.
- [4] Kỷ yếu của các kỳ thi Olympic Toán học sinh viên từ 2013 đến 2025.

LẬP KẾ HOẠCH ĐƯỜNG BAY BAO PHỦ KHÔNG ĐỒNG NHẤT

Trần Thị Cẩm Giang, Phan Anh

Trường Đại học Thủy lợi, email: giangttc@tlu.edu.vn

1. GIỚI THIỆU CHUNG

Trong những năm gần đây, UAV (phương tiện bay không người lái) được ứng dụng rộng rãi trong các lĩnh vực như nông nghiệp thông minh, giám sát, cứu hộ và khảo sát địa hình [1]. Xu hướng sử dụng nhiều UAV hoạt động đồng thời thay vì đơn lẻ ngày càng phổ biến, đặc biệt là hệ thống UAV không đồng nhất (Heterogeneous UAVs) [3]. Nhờ sự đa dạng về tốc độ, tầm bay, cảm biến và tải trọng, các UAV này giúp khai thác tối ưu khả năng từng chiếc, nâng cao hiệu quả toàn hệ thống.

Tuy nhiên, sự không đồng nhất cũng làm phức tạp bài toán lập kế hoạch đường bay, đặc biệt trong các nhiệm vụ yêu cầu bao phủ toàn bộ khu vực (Coverage Path Planning - CPP) mà không bỏ sót hoặc chồng lặp [2] [4]. Chủ đề “Lập kế hoạch đường bay bao phủ cho UAV không đồng nhất” giải quyết bài toán này bằng cách cải tiến thuật toán APPA (ACS-based Path Planning Algorithm), nhằm tối ưu phân công nhiệm vụ và lộ trình, giảm thời gian hoàn thành và nâng cao hiệu suất vận hành. Trong khi các phương pháp truyền thống như chia ô lưới còn nhiều hạn chế trong phối hợp UAV và xử lý ràng buộc môi trường, thì các hướng tiếp cận gần đây dựa trên học máy và heuristic dù tiềm năng nhưng vẫn gặp trở ngại về độ phức tạp và chi phí tính toán.

2. PHÁT BIỂU BÀI TOÁN

Tập hợp các UAV không đồng nhất: được ký hiệu là $U = \{U_1, U_2, \dots, U_n\}$, mỗi UAV U_i được đặc trưng bởi các thông số chính như sau: $U_i = \{V_i^{\max}, W_i\}$. Trong đó: $V_i^{\max}$ là tốc độ bay tối đa của UAV U_i , W_i là độ rộng quét của cảm biến trên UAV U_i . Bài báo không xét đến các yếu tố khác như dung lượng pin, khả năng liên lạc trên mỗi UAV.

Tập hợp vùng cần bao phủ: gồm m vùng độc lập $R = \{R_1, R_2, \dots, R_m\}$, mỗi vùng có diện tích A_j (m^2), vị trí tọa độ trung tâm (x_j, y_j) . Ma trận khoảng cách: Biểu diễn khoảng cách giữa các khu vực, được định nghĩa như sau: $D = \{D_{j,k}\}$. Trong đó: $D_{j,k}$ là khoảng cách bay từ khu vực R_j đến khu vực R_k . Nếu R_j và R_k trùng nhau, ta có $D_{j,k} = 0$. Ngược lại, $D_{j,k}$ được tính là khoảng cách Euclidean từ khu vực R_j đến khu vực R_k .

$$D_{j,k} = \sqrt{(x_j - x_k)^2 + (y_j - y_k)^2} \quad (1)$$

Ma trận vận tốc quét: biểu diễn tốc độ tối đa mà UAV U_i có thể quét khu vực R_j . Ta có: $V = \{V_{i,j}\}$, trong đó: $V_{i,j}$ là tốc độ tối đa của UAV U_i khi thực hiện quét khu vực R_j (m^2/s), $d_{i,j}$ là hệ số cản. Nếu UAV U_i không thực hiện nhiệm vụ quét tại khu vực R_j , ta có: $V_{i,j} = 0$. Ngược lại, $V_{i,j}$ được tính như sau:

$$V_{i,j} = d_{i,j} \cdot V_i^{\max} \quad (2)$$

Biểu diễn thời gian:

Thời gian quét tại khu vực R_j bởi UAV U_i , nếu $V_{i,j} > 0$:

$$TS_{i,j} = \frac{A_j}{V_{i,j} \cdot W_i} \quad (3)$$

Thời gian bay từ khu vực R_j đến khu vực R_k của UAV U_i :

$$TF_{i,j,k} = \frac{D_{j,k}}{V_i^{\max}} \quad (4)$$

Ràng buộc:

- (1) UAV không được bay về căn cứ cho đến khi hoàn thành tất cả khu vực được giao.
- (2) Tổng số UAV thực hiện nhiệm vụ không vượt quá tổng số UAV có.
- (3) Mỗi UAV phải được xuất phát từ căn cứ hoặc không tham gia nhiệm vụ.
- (4) Mỗi khu vực chỉ được quét một lần bởi một UAV.

Hàm mục tiêu: Tối thiểu hóa thời gian hoàn thành nhiệm vụ của toàn bộ hệ thống UAV.

$$\min \max_{1 \leq i \leq n} F(U_i) \quad (5)$$

trong đó: $F(U_i)$ là tổng thời gian thực hiện nhiệm vụ của UAV U_i .

$$F(U_i) = TS_{i,j} + TF_{i,j,k} \quad (6)$$

3. THUẬT TOÁN

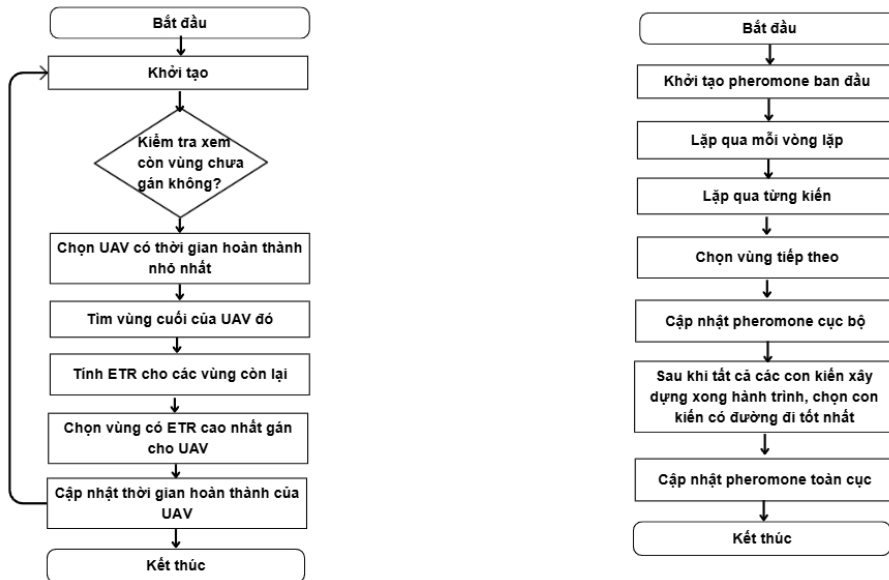
3.1. Tổng quan thuật toán

Để giải quyết bài toán lập kế hoạch đường bay bao phủ UAV không đồng nhất, đề xuất sử dụng thuật toán APPA. Thuật toán APPA được xây dựng dựa trên sự kết hợp giữa 2 giai đoạn:

Giai đoạn 1: Phân chia các vùng nhiệm vụ cho nhiều UAV không đồng nhất.

Giai đoạn 2: Lập kế hoạch đường đi cho từng UAV sao cho thời gian hoàn thành nhiệm vụ toàn cục là nhỏ nhất.

Độ phức tạp tính toán cao: bài toán bao gồm hai vấn đề tối ưu hóa lồng vào nhau đó là phân công nhiệm vụ và lập kế hoạch đường đi. Khác với [5], nghiên cứu này không chỉ sử dụng Ant Colony System (ACS) để tối ưu lộ trình mà còn kết hợp thêm một giai đoạn phân vùng nhiệm vụ dựa trên chỉ số *Effective Time Ratio* (ETR). Cách tiếp cận hai giai đoạn này cho phép cân bằng tải giữa các UAV không đồng nhất trước khi tiến hành tối ưu đường bay, từ đó giảm thiểu chênh lệch thời gian hoạt động và cải thiện đáng kể hiệu suất toàn cục. Đây là điểm khác biệt cốt lõi giúp APPA đạt được thời gian hoàn thành thấp và ổn định hơn trong nhiều kịch bản phức tạp.



Hình 1. Lưu đồ thuật toán APPA (phân vùng và tối ưu thứ tự)

3.2. Thuật toán APPA

3.1.1. Giai đoạn 1: Phân vùng

Giai đoạn này nhằm phân chia các vùng cho các UAV sao cho phân chia hợp lý để giảm chênh lệch thời gian hoạt động giữa các UAV. Các vùng sẽ được gán cho UAV nhân rồi dựa trên tỷ lệ thời gian hiệu quả (ETR).

Giả sử một UAV U_i đang ở vị trí R_j cần bay đến và quét khu vực R_k . Thời gian UAV U_i bay từ R_j đến R_k và quét khu vực R_k được ký hiệu là $TF_{i,j,k}$ và $TS_{i,k}$ tương ứng. Vì mục đích của UAV là quét khu vực R_k , nên thời gian cho việc quét khu vực R_k được gọi là thời gian hiệu quả. Tỷ lệ thời gian hiệu quả được định nghĩa là tỷ lệ giữa thời gian hiệu quả và tổng thời gian, theo công thức sau:

$$ETR_{i,j,k} = \frac{TS_{i,k}}{TF_{i,j,k} + TS_{i,k}} \quad (7)$$

Nếu giá trị ETR cao, điều đó có nghĩa là UAV dành nhiều thời gian cho nhiệm vụ quét hơn là di chuyển. Do đó, nếu vùng R_k có ETR cao sẽ được ưu tiên gán cho UAV U_i .

3.1.2. Giai đoạn 2: Lập kế hoạch đường bay bằng thuật toán ACS

Sau khi phân vùng, mỗi UAV được gán một tập hợp vùng nhiệm vụ. Giai đoạn tiếp theo là tối ưu thứ tự thăm các vùng này để rút ngắn thời gian hoàn thành. Thuật toán ACS được áp dụng riêng cho từng UAV trong quá trình lập lộ trình.

Khởi tạo

- Khởi tạo giá trị pheromone

Trên mỗi cạnh (i,j) , pheromone ban đầu được khởi tạo như sau:

$$\tau(i,j) = \tau_0 = \frac{1}{m \times L} \quad (8)$$

trong đó: m là số lượng vùng được gán cho UAV đang xét, L là độ dài hành trình được ước lượng bằng thuật toán Nearest Neighbor (NN).

- Khởi tạo độ hấp dẫn heuristic

Giá trị heuristic trên cung (i,j) được định nghĩa là:

$$\eta(i,j) = \frac{1}{D_{i,j}} \quad (9)$$

trong đó: $D_{i,j}$ là ma trận khoảng cách giữa các vùng. Giá trị $\eta(i,j)$ càng lớn thì cung đường (i,j) càng hấp dẫn, tức càng có khả năng được lựa chọn bởi kiến trong quá trình xây dựng lời giải.

Chọn khu vực tiếp theo

Tại mỗi bước đi, con kiến đang đứng sẽ chọn vùng tiếp theo để đi trong tập các vùng chưa đến theo một trong hai cách:

- Nếu $q < q_0$ chọn khai thác:

$$R_k = \arg \max_{R_j \in \pi_a} \tau(i,j)^\alpha \eta(i,j)^\beta \quad (10)$$

- Ngược lại, nếu $q > q_0$, chọn theo xác suất:

$$p(i,j) = \begin{cases} \frac{\tau(i,j)^\alpha \eta(i,j)^\beta}{\sum_{\forall k \in \pi_a} \tau(i,k)^\alpha \eta(i,k)^\beta}, & R_j \in \pi_a \\ 0, & \text{ngược lại} \end{cases} \quad (11)$$

trong đó: π_a là tập các vùng còn lại mà kiến chưa đi qua; α, β là hệ số điều chỉnh mức độ ảnh hưởng tương ứng của pheromone và heuristic ($\alpha, \beta > 0$); q_0 ($0 < q_0 < 1$) là tham số điều khiển chiến lược lựa chọn bước đi.

Cập nhật pheromone cục bộ

$$\tau(i,j) = (1-p)\tau(i,j) + p\tau_0 \quad (12)$$

trong đó: p là hệ số điều chỉnh tốc độ bay hơi pheromone cục bộ nằm trong khoảng $(0,1)$.

Cập nhật pheromone toàn cục

Sau khi tắt cả các cá thể kiến hoàn tất việc xây dựng lời giải trong một vòng lặp, thuật toán sẽ chọn ra lời giải tốt nhất toàn cục (p_{gb}). Chỉ những cung (i,j) thuộc trong lời giải tốt nhất này mới được tăng cường lượng pheromone.

$$\tau(i, j) = (1 - \varepsilon)\tau(i, j) + \varepsilon\Delta\tau(i, j) \quad (13)$$

Nếu (R_i, R_j) thuộc p_{gb} thì:

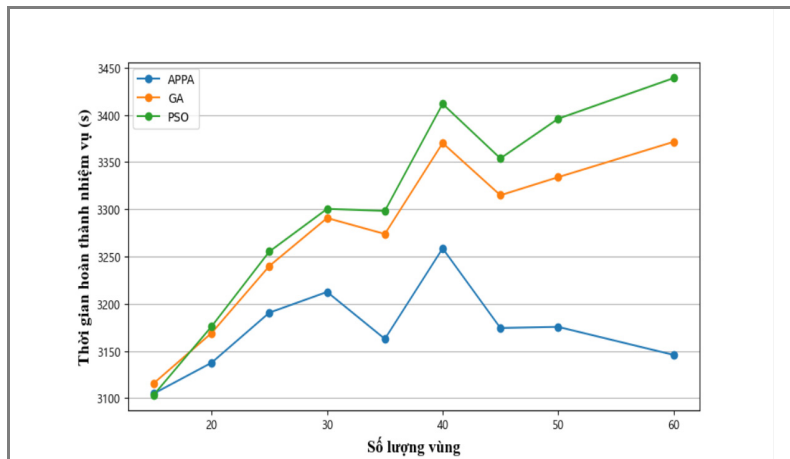
$$\Delta\tau(i, j) = \frac{1}{L_{gb}} \quad (14)$$

trong đó: L_{gb} là chi phí của lời giải tốt nhất toàn cục, ε là hệ số điều chỉnh tốc độ bay hơi pheromone toàn cục trong khoảng $(0,1)$.

4. KẾT QUẢ NGHIÊN CỨU

Các thí nghiệm được triển khai bằng Python 3.10.4 trên máy tính có bộ xử lý Intel Core i7, 16 GB RAM, hệ điều hành Windows 10. Các tham số chính của ACS được thiết lập ở mức chuẩn (số lượng kiến, số vòng lặp, hệ số bay hơi pheromone và tham số điều khiển q_0) để đảm bảo công bằng giữa các phương pháp so sánh. Với mỗi cấu hình, 50 bộ dữ liệu ngẫu nhiên được sinh ra và chạy lặp lại, kết quả báo cáo là giá trị trung bình thời gian hoàn thành nhiệm vụ.

4.1. Ảnh hưởng của số lượng vùng



Hình 2. Ảnh hưởng của số lượng vùng

Bảng 1. Thời gian hoàn thành nhiệm vụ khi hệ thống thay đổi số lượng vùng từ 15 đến 60

Số vùng	APPA	GA	PSO
15	3104.9	3115.8	3103.6
20	3137.5	3169.0	3176.0
25	3190.6	3240.0	3255.3
30	3212.5	3290.8	3300.5
35	3162.8	3273.9	3298.4
40	3258.6	3370.2	3411.6
45	3174.3	3314.9	3353.9
50	3175.5	3334.1	3395.9
60	3145.8	3371.5	3439.0

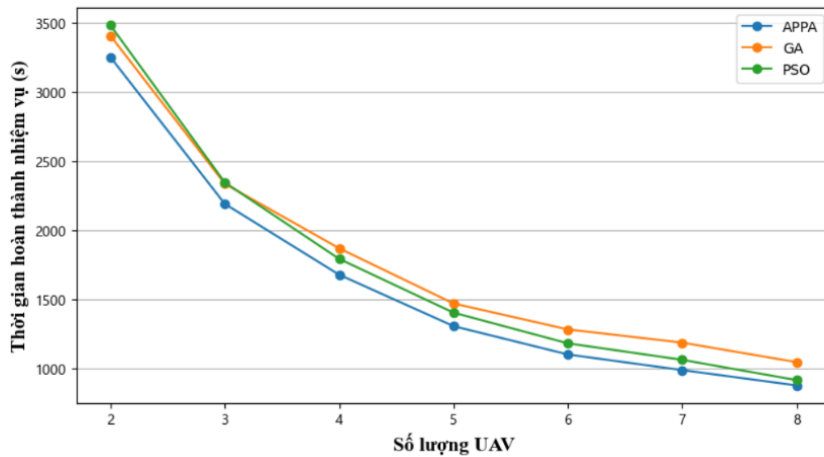
Số lượng vùng tăng kéo theo thời gian hoàn thành nhiệm vụ có xu hướng tăng, do khối lượng công việc của UAV tăng lên (nhiều vùng cần quét, nhiều đoạn đường cần bay hơn). Tuy nhiên, vì dữ liệu đầu vào là ngẫu nhiên, có thể xuất hiện các cấu hình thuận lợi giúp thuật toán chia vùng và

tối ưu hóa đường đi hiệu quả hơn, khiến thời gian không tăng đều mà có những điểm “gãy giảm”. Thuật toán APPA thể hiện khả năng phân chia nhiệm vụ và tối ưu hóa đường đi rất hiệu quả, đặc biệt rõ khi số lượng vùng lớn.

4.2. Ảnh hưởng của số lượng UAV

Bảng 2. Thời gian hoàn thành nhiệm vụ khi hệ thống thay đổi số lượng UAV từ 2 đến 8

Số UAV	APPA	GA	PSO
2	3249.4	3404.8	3479.3
3	2190.0	2334.9	2344.4
4	1678.3	1868.9	1792.5
5	1306.9	1470.6	1404.2
6	1102.5	1283.3	1182.8
7	989.7	1188.0	1063.8
8	878.8	1047.0	917.2



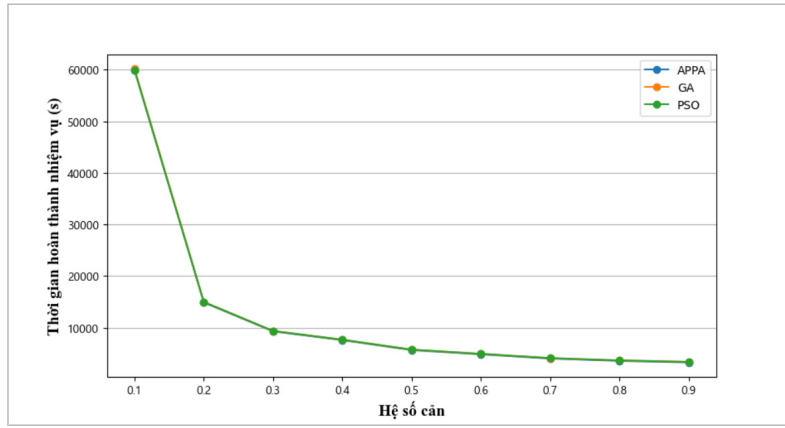
Hình 3. Biểu đồ ảnh hưởng của số lượng UAV

Khi số UAV tăng từ 2 đến 8, thời gian hoàn thành nhiệm vụ của các thuật toán giảm rõ rệt nhờ phân chia khối lượng công việc tốt hơn. Tuy nhiên, từ 6–8 UAV, mức giảm dần bão hòa do nhiệm vụ đã được chia nhỏ tối đa. Với ít UAV, mỗi UAV phải xử lý nhiều vùng nên mất thời gian hơn; khi tăng UAV, vùng được phân đều hơn, giảm tải và rút ngắn thời gian. Ở mọi mức UAV, APPA luôn đạt thời gian hoàn thành thấp nhất.

4.3. Ảnh hưởng hệ số cân

Bảng 3. Thời gian hoàn thành nhiệm vụ khi hệ thống thay đổi hệ số cân từ 0.1 đến 0.9

Hệ số cân	APPA	GA	PSO
0.1	59920.4	60148.5	59921.2
0.2	14891.9	14993.8	14971.5
0.3	9284.5	9369.4	9360.1
0.4	7584.6	7649.5	7674.4
0.5	5656.4	5732.7	5746.2
0.6	4832.9	4907.4	4925.9
0.7	4010.9	4083.9	4105.5
0.8	3562.4	3636.6	3660.3
0.9	3263.4	3364.9	3363.7



Hình 4. Biểu đồ ảnh hưởng của hệ số cân

Trong nghiên cứu này, hệ số cân (*di_j*) không mang ý nghĩa vật lý như lực cân, mà được hiểu là hệ số hiệu suất quét. Hệ số càng cao thì khu vực càng dễ quét, UAV di chuyển nhanh và tiêu tốn ít thời gian hơn. Ngược lại, hệ số càng thấp thì khu vực khó quét, UAV cần bay chậm và xử lý phức tạp hơn, làm tăng thời gian hoàn thành nhiệm vụ.

Kết quả cho thấy khi hệ số cân tăng, thời gian hoàn thành nhiệm vụ giảm đáng kể. Điều này phản ánh đúng bản chất bài toán: vùng càng thuận lợi thì UAV hoạt động càng hiệu quả. Tuy nhiên, khi hệ số cân tiệm cận giá trị tối đa (0.9), mức độ cải thiện thời gian giảm dần, do môi trường đã gần đạt điều kiện lý tưởng, và các thuật toán khó tiếp tục tối ưu sâu hơn.

5. KẾT LUẬN

Nghiên cứu này triển khai và đánh giá thuật toán APPA cho bài toán lập kế hoạch đường bay bao phủ với đội UAV không đồng nhất, gồm hai giai đoạn: phân vùng nhiệm vụ và tối ưu lộ trình bay. So với GA và PSO, APPA chiếm ưu thế về thời gian hoàn thành, đặc biệt trong kịch bản phức tạp với số vùng, số UAV và hệ số cân thay đổi. Ngoài ra APPA cũng đóng góp tính mới trong việc đề xuất cơ chế phân vùng dựa trên tỷ lệ thời gian hiệu quả (ETR) nhằm phân công nhiệm vụ hợp lý giữa các UAV không đồng nhất, kết hợp với tối ưu lộ trình bằng ACS cho từng UAV.

Thực nghiệm cho thấy APPA không chỉ rút ngắn thời gian mà còn duy trì độ ổn định cao nhờ phân chia vùng hợp lý, cân bằng tải giữa các UAV. Việc tối ưu đồng thời phân vùng và lộ trình giúp APPA đạt hiệu suất vượt trội, phù hợp với các ứng dụng thực. Tích hợp thêm yếu tố như pin, liên lạc và ràng buộc môi trường trong tương lai sẽ nâng cao tính ứng dụng thực tế.

6. TÀI LIỆU THAM KHẢO

- [1] N. K. Shubhani Aggarwal, 2019, "Path Planning Techniques for Unmanned Aerial Vehicles: A Review, Solutions, and Challenges", Computer Communications.
- [2] F. L. Y. Z. T. Y. Y. L. X. D. Jinchao Chen, 2022, "Coverage path planning of heterogeneous unmanned aerial vehicles based on ant colony system", Swarm and Evolutionary Computation.
- [3] Datsko, Denys. 2024. *Energy-Aware Multi-UAV Coverage Mission Planning with Optimal Speed of Flight*. IEEE Robotics and Automation Letters.
- [4] B. B. L. M. Gaurav Singhal, 2018, "Unmanned Aerial Vehicle classification, Applications and challenges: A Review", Preprints.
- [5] Şimşek 2009, "Free and forced vibration of a functionally graded beam subjected to a concentrated moving harmonic load", Composite Structures.

NHẬN DẠNG GIỌNG BA MIỀN VIỆT NAM DỰA TRÊN KẾT HỢP ĐẶC TRƯNG ASR VÀ SV

Lê Minh Đức¹, Đỗ Văn Hải²

¹Đại học Công nghệ - Đại học Quốc gia Hà Nội

²Trường Đại học Thủy lợi

1. GIỚI THIỆU

Trong những năm gần đây, nhận dạng vùng miền (accent recognition) trong giọng nói tiếng Việt ngày càng thu hút sự quan tâm nhờ tiềm năng ứng dụng rộng rãi trong các hệ thống tương tác người-máy, trợ lý ảo, cũng như nghiên cứu phân tích mạng xã hội. Việc phân biệt rõ Bắc, Trung, Nam không chỉ góp phần nâng cao chất lượng hiểu ngôn ngữ tự nhiên, mà còn giúp cá nhân hóa trải nghiệm người dùng (ví dụ: tùy chỉnh phản hồi âm thanh) và thăm dò hành vi, quan điểm theo địa bàn địa lý.

Trước đây, các phương pháp thống kê như GMM-HMM kết hợp các đặc trưng MFCC, formant và F0 từng được áp dụng để phân loại vùng miền. Những phương pháp này tuy đơn giản, ít tốn kém tài nguyên, nhưng độ chính xác chưa cao. Sự bùng nổ của học sâu và kỹ thuật transfer learning đã mở ra hướng đi mới: các mô hình self-supervised như Wav2Vec 2.0 [1], PhoWhisper [2] và WavLM [3] cho phép thu được embedding ngữ âm giàu biểu diễn, trong khi các kiến trúc speaker verification đem đến các đặc trưng giọng nói đặc thù cá nhân. Một số công trình cũng đã thử kết hợp hai nguồn thông tin này thông qua fine-tuning hoặc transfer learning chung, nhưng hầu hết mới dừng lại ở khâu ghép nối (concat) hoặc cộng đơn giản, chưa khai thác triệt để tính chất local-global của mỗi luồng đặc trưng.

Nghiên cứu này giới thiệu cơ chế Attentive Feature Fusion cho hai luồng embedding Automatic Speech Recognition (ASR) và speaker verification (SV). Sau khi chiếu xuống cùng không gian qua Multi Layer Perceptron (MLP), mỗi khung thời gian sẽ được cân chỉnh trọng số kết hợp thông qua module attention song song lặp (Parallel iAFF), cho phép điều chỉnh dần dần và giữ lại chi tiết quan trọng của cả hai embedding.

2. CÁC NGHIÊN CỨU LIÊN QUAN

Trong nghiên cứu [4] đã sử dụng các mô hình học tự giám sát như Wav2Vec 2.0 [1] và PhoWhisper [2] trên bộ dữ liệu ViMD, nhằm khai thác đặc trưng biểu diễn ngữ âm phong phú trong dữ liệu nói. Bên cạnh đó, [6] triển khai các kiến trúc dựa trên CNN ([7], [8]) trong thị giác máy tính, để xử lý mel-spectrogram cho bài toán này. Nghiên cứu [12] đã chỉ ra mối liên hệ giữa các đặc trưng Automatic Speech Recognition (ASR) và các đặc trưng từ bài toán speaker verification với đặc trưng vùng miền phương ngữ, từ đó đề xuất học chuyển giao từ các mô hình ASR và speaker verification để đạt hiệu suất cao hơn.

Bên cạnh hướng tiếp cận qua biểu diễn âm học và transfer learning, một hướng tiếp cận hiệu quả là kết hợp các luồng đặc trưng từ nhiều nguồn hữu ích cho bài toán. Các phương pháp kết hợp dựa trên chú ý (attentive fusion) đóng vai trò then chốt nhờ khả năng học trọng số động cho từng nguồn đầu vào, thay vì kết hợp đơn giản như cộng hoặc xếp chồng. SENet [11], vốn được thiết kế cho thị giác máy tính, trích xuất trọng số chú ý cho từng kênh sử dụng hai bước squeeze và excitation. Áp dụng cơ chế channel-wise attention trên toàn bộ đặc trưng đầu vào thông qua hai bước squeeze và excitation. Mặc dù đơn giản, SE chỉ khai thác ngữ cảnh toàn cục mà bỏ qua các biến đổi cục bộ theo từng frame, dẫn đến có thể bỏ sót các chi tiết nhỏ trong luồng tín hiệu âm thanh. MS-CAM [9] mở rộng SENet bằng cách bổ sung thêm nhánh attention local: song song với global average pooling, một cấu trúc bottleneck khác được áp dụng trực tiếp lên từng frame, giúp mô hình vừa nắm bắt được bối cảnh tổng thể vừa bảo toàn các đặc trưng cục bộ. Nhờ vậy, MS-CAM giảm thiểu tình

trạng "làm mượt" quá mức của SE, nhưng vẫn chỉ tập trung vào một lần sinh attention, chưa tối ưu trong việc điều chỉnh trọng số kết hợp theo từng cặp đặc trưng. Dựa trên MS-CAM, AFF [9] sử dụng attention để học trọng số kết hợp cho từng nguồn, từ đó điều khiển mức đóng góp của từng luồng đặc trưng. Bằng cách cho phép mô hình học tỉ lệ kết hợp phù hợp cho từng phân tử, AFF khắc phục nhược điểm của phép cộng (mượt chi tiết) và phép concat (tăng tham số). Tuy vậy, AFF chỉ thực hiện một bước attention duy nhất với input đã được ghép, dẫn đến hạn chế trong việc tận dụng sâu hơn thông tin ban đầu để sinh trọng số attention. Để khắc phục điểm này, iterative AFF [9] đề xuất lặp lại việc sinh trọng số attention nhiều lần, mỗi lần dùng đầu ra của bước trước làm đầu vào cho module MS-CAM. Cơ chế lặp giúp mô hình hiệu chỉnh dần trọng số kết hợp, giúp sinh ra trọng số attention tốt hơn, nhưng vẫn sử dụng chung một attention map, chưa tách riêng được ảnh hưởng của mỗi luồng. Một hướng tiếp cận khác là Parallel AFF [10], trong đó mô hình sinh song song hai bản đồ attention riêng biệt cho mỗi luồng đặc trưng, từ đó kết hợp mà không "pha trộn" lẫn nhau quá sớm. Cách làm này bảo tồn tối đa thông tin đặc trưng riêng, nhưng lại thiếu cơ chế lặp để tinh chỉnh trọng số kết hợp qua nhiều bước.

3. PHƯƠNG PHÁP KẾT HỢP ĐẶC TRƯNG

Các đặc trưng sẽ được kết hợp (feature fusion) giữa hai mô hình tiền huấn luyện: mô hình nhận dạng tiếng nói (ASR) *PhoWhisper-small* [2] và mô hình xác minh người nói (Speaker Verification) *WavLM-Base-Plus* for Speaker Verification [3]. Đặc trưng ASR sẽ được sử dụng như đặc trưng cục bộ $X_{raw} \in R^{T \times D_1}$ với độ phân giải $T = 1500$ frame cho 10s. Đặc trưng Speaker Verification sẽ được sử dụng như đặc trưng toàn cục $Y_{raw} \in R^{1 \times D_2}$. Trong đó D_1 và D_2 thể hiện số chiều ban đầu của hai đặc trưng. Hai đặc trưng này sẽ được chiếu xuống cùng chiều không gian $X \in R^{T \times D}$, $Y \in R^{1 \times D}$ sử dụng hai mô hình Multi-layer Perceptron (MLP) 2 lớp: $X = \text{MLP}(X_{raw})$; $Y = \text{MLP}(Y_{raw})$, giúp đảm bảo cùng số chiều và giữ sự tương thích ngữ nghĩa giữa hai đặc trưng.

3.1. Phương pháp fusion cơ bản

Phương pháp cộng (Add Fusion) và phương pháp xếp chồng (Concat Fusion) sẽ được sử dụng làm hai phương pháp baseline. Cụ thể, đặc trưng $Z \in R^{T \times D}$ được kết hợp theo phương pháp Add Fusion như sau: $Z = X \oplus Y$, và đặc trưng $Z \in R^{T \times 2D}$ sẽ được kết hợp theo phương pháp Concat Fusion: $Z = \text{Concat}(X, Y)$, trong đó $\oplus$ là phép cộng broadcast, và Concat là phép xếp chồng các đặc trưng. Nhược điểm của phép cộng là thường cần hai chiều giống nhau và thiếu linh hoạt, trong khi phép xếp chồng lại làm tăng số lượng tham số và khiến mô hình trở nên cồng kềnh. Vì những hạn chế này, ở phần tiếp theo các cơ chế fusion dựa trên attention sẽ được thử nghiệm.

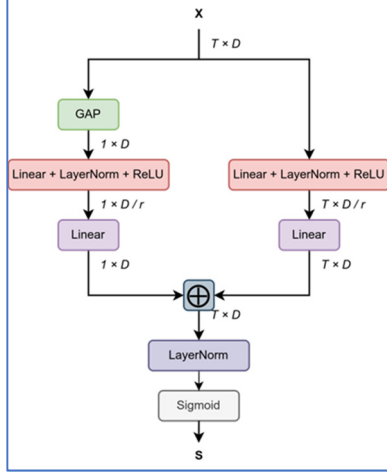
3.2. Multi-scale channel attention

Nghiên cứu [9] đề xuất module multi-scale channel attention (MS-CAM) như một cải tiến cho module Squeeze-and-Excitation (SE) [11] bằng cách cho phép mô hình học được attention từ ngữ cảnh hai góc nhìn local và global (multi-scale), được mô tả như trong Hình 1(a).

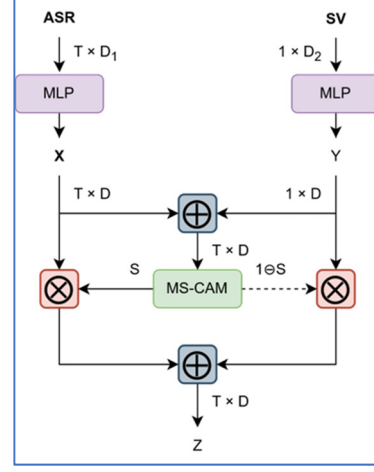
Cụ thể, với đặc trưng đầu vào $X \in R^{L \times D}$ gồm L frames là các vector D chiều, ngữ cảnh đặc trưng toàn cục $\mathbf{g}(X) \in R^D$ trong module MS-CAM được tính như sau: $\mathbf{g}(X) = W_2 \delta \left(LN \left(W_1 g(X) \right) \right)$, trong đó $g(X) = \frac{1}{L} \sum_{i=1}^L X[i, :]$ được tính bằng Global Average Pooling (GAP), δ là hàm kích hoạt Rectified Linear Unit (ReLU), và LN là lớp layer normalization (LN). Cấu trúc bottleneck gồm có hai lớp fully connected (FC), trong đó $W_1 \in R^{\frac{D}{r} \times D}$ là lớp giảm chiều, còn $W_2 \in R^{D \times \frac{D}{r}}$ là lớp tăng chiều, với r là tỷ lệ giảm chiều của đặc trưng. Tương tự, ngữ cảnh đặc trưng cục bộ $L(X) \in R^{L \times D}$ được tính thông qua một cấu trúc bottleneck như sau: $L(X) = FC_2 \left(LN \left(FC_1(X) \right) \right)$.

Kích thước ma trận trọng số của FC_1 và FC_2 lần lượt là $\frac{D}{r} \times D$ và $D \times \frac{D}{r}$. Lưu ý rằng $L(X)$ có cùng kích thước với đặc trưng đầu vào, giúp bảo toàn và làm nổi bật các chi tiết nhỏ trong đặc trưng

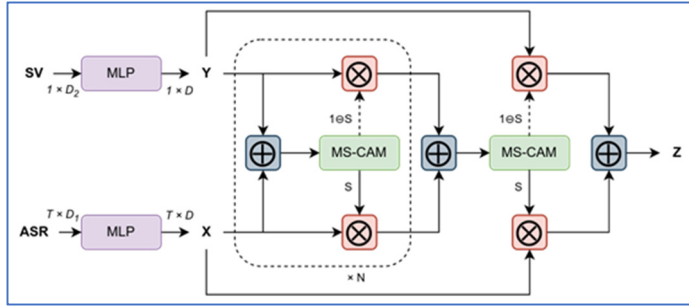
cục bộ. Với ngữ cảnh kênh toàn cục $g(X)$ và ngữ cảnh kênh cục bộ, đặc trưng đã được tinh chỉnh $X' \in R^{L \times D}$ thông qua MS-CAM được tính như sau: $X' = X \otimes M(X) = X \otimes \sigma \left(LN(L(X) + g(X)) \right)$, trong đó $M(X)$ là trọng số attention được sinh ra bởi MS-CAM, $\oplus$ biểu thị phép cộng broadcasting $\otimes$ biểu thị phép nhân element-wise, và σ là hàm Sigmoid.



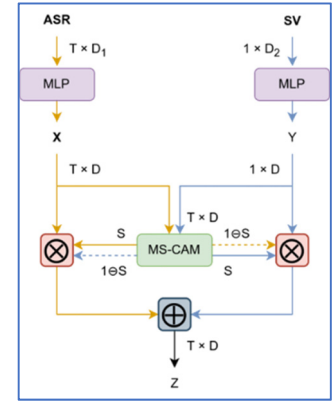
(a) MS-CAM



(b) AFF



(c) iAFF



(d) AFF Parallel.

Hình 1. Các phương pháp kết hợp đặc trưng

3.3. Attentional feature fusion

Để kết hợp hai đặc trưng $X \in R^{T \times D}$ và $Y \in R^{1 \times D}$, dựa trên module MS-CAM M , module *Attentional Feature Fusion* (AFF) tạo ra đặc trưng kết hợp $Z \in R^{T \times D}$ như sau:

$$Z = M(X \oplus Y) \otimes X \oplus (1 - M(X \oplus Y)) \otimes Y \quad (1)$$

Module AFF được minh họa trong Hình 1(b), trong đó đường nét đứt thể hiện trọng số attention $1 - M(X \oplus Y)$. Ở đây, trọng số kết hợp $M(X \oplus Y)$ và $1 - M(X \oplus Y)$ đều có giá trị trong khoảng 0 đến 1, giúp mô hình thực hiện lựa chọn đặc trưng bằng phép trung bình có trọng số giữa X và Y .

3.3.1. Iterative attentional feature fusion

Các phương pháp sử dụng đầy đủ ngữ cảnh thường gặp phải một vấn đề tất yếu, đó là cách tích hợp đặc trưng đầu vào ban đầu, cụ thể ở đây là đầu vào cho module MS-CAM M . Đầu vào cho module attention có thể ảnh hưởng đến trọng số kết hợp cuối cùng. Vì vậy [9] đề xuất sử dụng thêm một module attention khác để tích hợp đặc trưng đầu vào ban đầu cho MS-CAM, gọi là iterative Attentional Feature Fusion (iAFF), được minh họa trong Hình 1(c). Khi đó, phép tích hợp ban đầu $X \oplus Y$ được viết lại như sau:

$$X \uplus Y = M(X \oplus Y) \otimes X + (1 - M(X \oplus Y)) \otimes Y \quad (2)$$

3.3.2. Parallel attentional feature fusion

Nghiên cứu [10] đề xuất một phương pháp khác nhằm giải quyết vấn đề tích hợp đặc trưng ban đầu cho module attention, bằng cách áp dụng song song module MS-CAM cho các đặc trưng X và Y để sinh ra hai trọng số attention riêng biệt S^X và S^Y . Sau đó trọng số sẽ được sử dụng song song để kết hợp hai đặc trưng X và Y ban đầu, như được minh họa trong hình 4. Bằng cách tạo ra các trọng số attention S^X và S^Y dựa trên nội dung riêng của X và Y, mô hình có thể linh hoạt giữ lại tối đa các chi tiết quan trọng trong từng luồng và chỉ kết hợp chúng ở giai đoạn cuối. Quá trình tính toán được mô tả như sau:

$$S^X = \text{MS-CAM}(X) \quad (3)$$

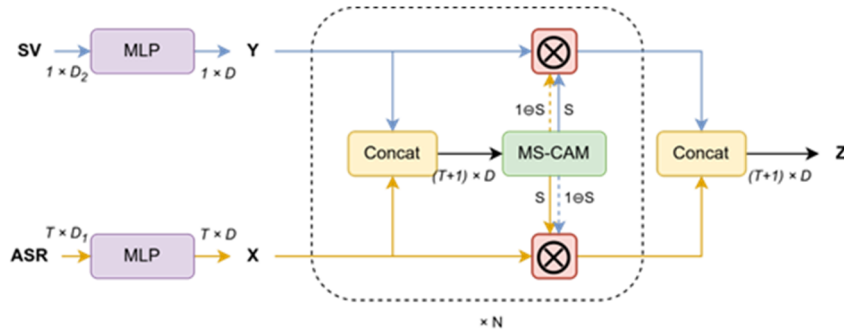
$$S^Y = \text{MS-CAM}(Y) \quad (4)$$

$$Z = X \otimes S^X \otimes (1 \ominus S^Y) + Y \otimes (1 \ominus S^X) \otimes S^Y \quad (5)$$

trong đó S^X và S^Y là các trọng số attention được sinh ra dựa trên nội dung của X và Y tương ứng.

4. PHƯƠNG PHÁP ĐỀ XUẤT

Mặc dù đã cải thiện khả năng biểu diễn, phương pháp iAFF vẫn sử dụng một đầu vào duy nhất cho module attention, trong khi phương pháp Parallel AFF không sử dụng lặp để cải thiện trọng số attention. Để kết hợp những ưu điểm của iAFF và Parallel AFF, phương pháp *Parallel iterative Attentional Feature Fusion* (Parallel iAFF) được đề xuất nhằm giải quyết triệt để hơn vấn đề tích hợp đặc trưng ban đầu cho module attention. Cụ thể, với đầu vào lần lặp thứ i là $X^{(i)}$ và $Y^{(i)}$, dữ liệu đầu vào sẽ được ghép theo chiều thời gian thành $C^{(i)} \in R^{(T+1) \times D}$. Nhờ giữ nguyên toàn bộ T+1 frames, MS-CAM sẽ học được đặc trưng cục bộ toàn cục trong cùng một tensor $C^{(i)}$. Module MS-CAM sẽ được áp dụng song song hai lần cho $C^{(i)}$ để đưa ra hai trọng số attention $L^{(i)}, G^{(i)}$, giúp mô hình lựa chọn đặc trưng từ từng luồng mà không gặp vấn đề tích hợp đặc trưng ban đầu. Module được minh họa trong Hình 2. Quá trình tính toán được mô tả như sau. Trước tiên đặc trưng sẽ được xếp chồng theo chiều thời gian như sau: $C^{(i)} = \text{Concat}(X^{(i)}, Y^{(i)})$.



Hình 2: Minh họa module Parallel iAFF

MS-CAM sẽ được tách thành hai nhánh: $\text{MS-CAM}_{\text{local}}$ học trọng số trên từng frame, còn $\text{MS-CAM}_{\text{global}}$ học trọng số toàn cục. $L^{(i)} = \text{MS-CAM}_{\text{local}}(C^{(i)})$ và $G^{(i)} = \text{MS-CAM}_{\text{global}}(C^{(i)})$

Sau đó, mỗi bộ trọng số attention được tách ra thành các thành phần dành cho luồng local và global: $L_X^{(i)} = L_{1:T,:}^{(i)}, L_Y^{(i)} = L_{T+1,:}^{(i)}$ và $G_X^{(i)} = G_{1:T,:}^{(i)}, G_Y^{(i)} = G_{T+1,:}^{(i)}$.

Tiếp theo, đặc trưng local và global tại bước lặp i+1 được cập nhật tuần tự qua hai phép gộp có trọng số:

$$X^{(i+1)} = X^{(i)} \otimes L_X^{(i)} \otimes (1 - L_Y^{(i)}) + Y^{(i)} \otimes (1 - L_X^{(i)}) \otimes L_Y^{(i)} \quad (6)$$

$$Y^{(i+1)} = \text{GAP} \left(G_X^{(i)} \otimes X^{(i)} \otimes (1 - G_Y^{(i)}) + (1 - G_X^{(i)}) \otimes Y^{(i)} \otimes G_Y^{(i)} \right) \quad (7)$$

Trong đó GAP nhằm mục đích chuyển đặc trưng local-global thành một đặc trưng global duy nhất $Y^{(i+1)}$. Cuối cùng, đặc trưng được kết hợp được tính như sau: $Z^{(i+1)} = \text{Concat}(X^{(i+1)}, Y^{(i+1)})$

Qua đó, Parallel iAFF vừa giữ được chi tiết local lẫn global, vừa tận dụng được đặc trưng đầu vào qua các bước lặp bổ trợ, giúp cải thiện tính biểu diễn so với các phương pháp trước.

5. THỬ NGHIỆM

5.1. Thiết lập thử nghiệm

- Dữ liệu: Dữ liệu được sử dụng là bộ dữ liệu accent Zalo với 2879 file wav được gán nhãn giới tính và 3 miền Bắc, Trung, Nam. Các file audio dài khoảng từ 1s đến 32s với tần số lấy mẫu 16kHz. Đầu vào được sử dụng là các chunk 10s. Những file dài dưới 10s sẽ được áp dụng cyclic padding. Dữ liệu được chia ngẫu nhiên thành ba tập training (65%), validation (15%), testing (20%).

- Metrics: Đánh giá bằng Accuracy và F1.

- Các phương pháp so sánh:

+ Baseline Finetune W2V2-base: Fine-tune Wav2Vec 2.0 base cho bài toán phân vùng miền.

+ Linear Probing WavLM (SV): Chỉ huấn luyện lớp phân loại trên embedding của WavLM-Base-Plus SV.

+ Finetune WavLM (SV): Fine-tune toàn bộ WavLM-Base-Plus SV cho bài toán vùng miền.

+ Linear Probing PhoWhisper (ASR): Chỉ huấn luyện lớp phân loại trên embedding ASR từ PhoWhisper-small.

+ Add Fusion/Concat Fusion: Kết hợp embedding SV + ASR bằng cộng broadcast và ghép chuỗi.

+ Parallel iAFF K=2, K=3: Phương pháp Parallel iterative AFF với số vòng lặp K.

5.2. Kết quả thử nghiệm

Bảng 1. Kết quả thử nghiệm của phương pháp đề xuất và các phương pháp cơ sở khác

Phương pháp	Accuracy	F1
Baseline Finetune Wav2Vec 2.0	0.7639	0.7586
Linear Probing WavLM (SV)	0.6788	0.6651
Finetune WavLM (SV)	0.8524	0.8443
Linear Probing PhoWhisper (ASR)	0.8854	0.8813
Add Fusion	0.9080	0.9054
Concat Fusion	0.8767	0.8710
Proposed Parallel iAFF K=2	0.9115	0.9078
Proposed Parallel iAFF K=3	0.9132	0.9085

Có một số kết luận được rút ra từ bảng trên:

- **Embedding đơn luồng:** Linear probing trên PhoWhisper (ASR) đạt Acc = 88.54%, F1 = 88.13%, cho thấy embedding ngữ âm giàu thông tin semantic có khả năng nhận dạng vùng miền tốt hơn embedding SV (Acc = 67.88%, F1 = 66.51%). Fine-tune WavLM cho SV cải thiện đáng kể (Acc = 85.24%, F1 = 84.43%), nhưng vẫn thấp hơn ASR.

- **Fusion đơn giản:** Add Fusion đạt Acc = 90.80%, F1 = 90.54%, cho thấy việc kết hợp hai luồng embedding bổ sung lẫn nhau. Ngược lại, Concat Fusion chỉ đạt Acc = 87.67%, F1 = 87.10%, chứng tỏ phép ghép chuỗi dễ gây over-parameter và không hiệu quả bằng Add Fusion.

- **Attention-based Fusion:** Phương pháp Parallel iAFF với K = 2 vượt lên Acc = 91.15%, F1 = 90.78%, và tiếp tục cải thiện nhẹ với K = 3 (Acc = 91.32%, F1 = 90.85%). Kết quả này khẳng định

cơ chế song song giúp mô hình vừa giữ được chi tiết local–global, vừa tinh chỉnh trọng số kết hợp hiệu quả qua các vòng, mang lại độ chính xác cao nhất.

6. KẾT LUẬN

Trong nghiên cứu này, báo cáo tập trung giải quyết bài toán nhận dạng vùng miền tiếng nói bằng cách kết hợp hai loại đặc trưng: đặc trưng cục bộ từ mô hình ASR (PhoWhisper-small) và đặc trưng toàn cục từ mô hình Speaker Verification (WavLM-Base-Plus). Các đóng góp chính của nghiên cứu này bao gồm: (1) thiết kế module Parallel iterative Attentional Feature Fusion, cho phép tích hợp hai luồng embedding ASR và SV thông qua cơ chế chú ý lặp song song; (2) thực hiện đánh giá toàn diện trên Zalo Accent Corpus, so sánh hiệu năng với nhiều phương pháp baseline và các chiến lược fusion đơn giản. Kết quả thực nghiệm cho thấy, phương pháp Parallel iAFF đề xuất với số vòng lặp $K=3$ đã đạt được cải thiện 2.78% accuracy và 2.72% F1; đồng thời vượt trội hơn hẳn so với các phương pháp fusion đơn giản (Add Fusion, Concat Fusion) và các phương pháp dựa trên attention AFF, iAFF, Parallel AFF.

7. TÀI LIỆU THAM KHẢO

- [1] Alexei Baevski, Henry Zhou, Abdelrahman Mohamed, & Michael Auli. (2020). wav2vec 2.0: A Framework for Self-Supervised Learning of Speech Representations.
- [2] Thanh-Thien Le, Linh The Nguyen, & Dat Quoc Nguyen. (2024). PhoWhisper: Automatic Speech Recognition for Vietnamese.
- [3] Chen, S., Wang, C., Chen, Z., Wu, Y., Liu, S., Chen, Z., Li, J., Kanda, N., Yoshioka, T., Xiao, X., Wu, J., Zhou, L., Ren, S., Qian, Y., Qian, Y., Wu, J., Zeng, M., Yu, X., & Wei, F. (2022). WavLM: Large-Scale Self-Supervised Pre-Training for Full Stack Speech Processing. *IEEE Journal of Selected Topics in Signal Processing*, 16(6), 1505–1518.
- [4] Nguyen Van Dinh, Thanh Chi Dang, Luan Thanh Nguyen, & Kiet Van Nguyen. (2024). Multi-Dialect Vietnamese: Task, Dataset, Baseline Models and Challenges.
- [5] Shergill, V. (2021). Accent and Gender Recognition from English Language Speech and Audio Using Signal Processing and Deep Learning. In *Hybrid Intelligent Systems* (pp. 62–72). Springer International Publishing.
- [6] Nguyen, T., Nguyen, U., Pham, T., Nguyen, T., & Nguyen, B. (2024). Gender and Dialect Classification for the Vietnamese Language. In *Proceedings of the 38th Pacific Asia Conference on Language, Information and Computation* (pp. 389–397). Tokyo University of Foreign Studies.
- [7] Kaiming He, Xiangyu Zhang, Shaoqing Ren, & Jian Sun. (2015). Deep Residual Learning for Image Recognition.
- [8] Gao Huang, Zhuang Liu, Laurens van der Maaten, & Kilian Q. Weinberger. (2018). Densely Connected Convolutional Networks.
- [9] Yimian Dai, Fabian Gieseke, Stefan Oehmcke, Yiquan Wu, & Kobus Barnard. (2020). Attentional Feature Fusion.
- [10] Bei Liu, Zhengyang Chen, & Yanmin Qian (2022). Attentive Feature Fusion for Robust Speaker Verification. In *Interspeech 2022* (pp. 286–290).
- [11] Jie Hu, Li Shen, Samuel Albanie, Gang Sun, & Enhua Wu. (2019). Squeeze-and-Excitation Networks.
- [12] Tạ Bảo, T., Le, N., & Đỗ, H. (2024). Transfer learning methods for low-resource speech accent recognition: A case study on Vietnamese language. *Engineering Applications of Artificial Intelligence*, 132, 107895.
- [13] Hung, P., Van Loan, T., & Quang, N. (2016). Statistical Analysis of Vietnamese Dialect Corpus and Dialect Identification Experiments.

VIETNAMESE TEXT NORMALIZATION FOR TEXT-TO-SPEECH

Nguyen Tung Luong¹, Do Van Hai²

¹*Viettel AI, Viettel Group*

²*Thuyloi University*

1. INTRODUCTION

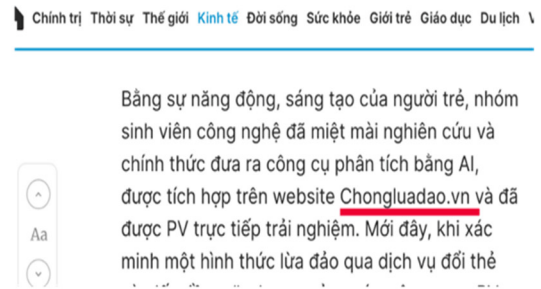
Text-to-Speech (TTS) technology plays an essential role in many modern applications, including virtual assistants, automated customer service, audiobooks, and assistive tools for people with visual impairments [1]. The goal of a TTS system is to convert written text into clear, natural-sounding speech that closely mimics human prosody and pronunciation [1].

To achieve high-quality synthesis, the input text must be clean, grammatically correct, and fully normalized [2, 3]. However, in real-world scenarios, especially in Vietnamese, input text often suffers from multiple issues: lack of diacritics, missing spaces between words, or inclusion of special spans such as email addresses, URLs, or search keywords that do not follow standard language rules [4].

For example, a user might input a query like *xinchaoquyvi* instead of the properly spaced and diacritic-rich version *Xin chào quý vị*. Such errors can significantly degrade the quality of generated speech, leading to mispronunciations or unnatural intonation [8-10]. These problems often occur in contexts such as email addresses and usernames like *hotrokhachhang@gmail.com*, domain names and brand names like *dienmayxanh*, or OCR outputs from scanned documents like *Xinchao*. Humans resolve these forms rapidly via contextual reasoning, prior lexical and idiomatic knowledge, and cognitive flexibility, highlighting the difficulty of replicating such inference in computational models [5, 12-14]. Consequently, TTS/NLP systems should aim to learn robust representations of Vietnamese structure, exploit wider sentential context beyond local cues, and jointly perform word segmentation with diacritic restoration to accurately reconstruct noisy text [11, 13, 14].



Source: thanhnien.vn



Source: Vietnam Net

Figure 1. Examples of typical non-diacritic, unsegmented Vietnamese inputs in real scenarios

In this report, we propose a unified text-normalization framework for Vietnamese that jointly performs word segmentation and diacritics restoration, converting raw noisy strings into TTS-ready text. The method casts spacing and accent recovery as a single sequence-editing task with a joint decoder, allowing cross-task signals (e.g., predicted syllable boundaries) to disambiguate homographs. Training uses span-level supervision: only the corrupted n-gram is masked while the remaining context is copied, mirroring real post-OCR cleanup and reducing unnecessary edits. To capture long-range dependencies, each example couples the full sentence with the localized corrupted span, enabling the model to learn agreement cues (e.g., tense, named entities) that n-gram baselines miss.

This integrated design avoids the cascading error loop of pipeline systems and delivers cleaner inputs for modern Vietnamese TTS. To the best of our knowledge, this is the first approach to formulate Vietnamese spacing and diacritic recovery as a single, end-to-end task.

2. METHODOLOGY

This section details our method for Vietnamese text normalization, comprising context-rich training data and two model-design choices.

2.1. Synthetic Data Generation Pipeline

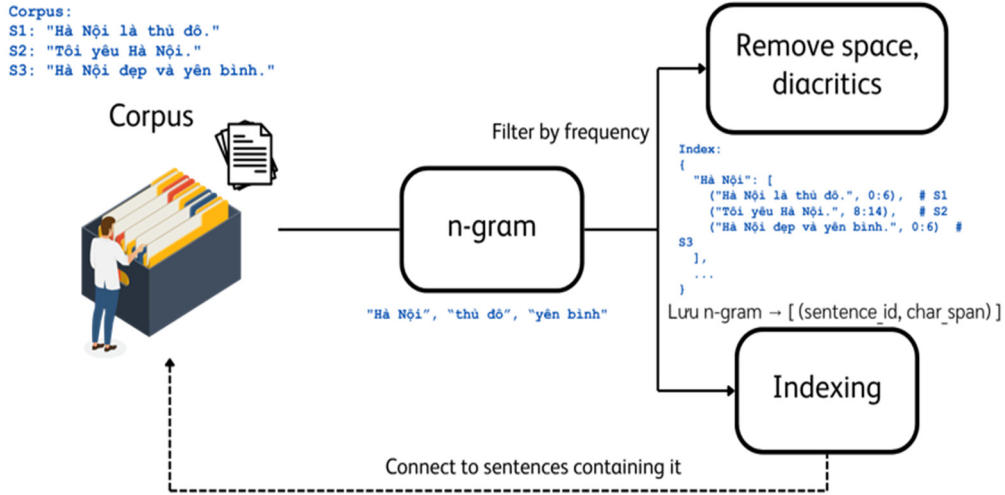


Figure 2. End-to-end synthetic data pipeline. Starting from a clean corpus we extract frequent n-grams, remove spaces and diacritics to generate realistic noise spans, and finally produce aligned (input, target, span) triples. Both “no_context” and “full_context” variants are emitted so the model can learn phrase-level and sentence-level cues simultaneously.

To compensate for the scarcity of naturally corrupted Vietnamese text, we construct a large-scale synthetic dataset from clean corpora using a controlled noise injection framework. As shown in Figure 2, starting with 100,000 sentences from a Vietnamese news corpus [6], we extract frequent 2–6-grams, filtering by length-aware frequency thresholds to preserve meaningful patterns while discarding noise. These n-grams are indexed with precise character spans, then normalized by lowercasing, removing diacritics, and stripping whitespace to simulate common real-world corruption. For each n-gram, a limited number of contexts are sampled to balance frequency and diversity. We generate two training formats: isolated phrase restoration and in-context sentence correction with span annotations. Finally, the dataset is split into stratified train and validation sets, preserving n-gram diversity and preventing data leakage.

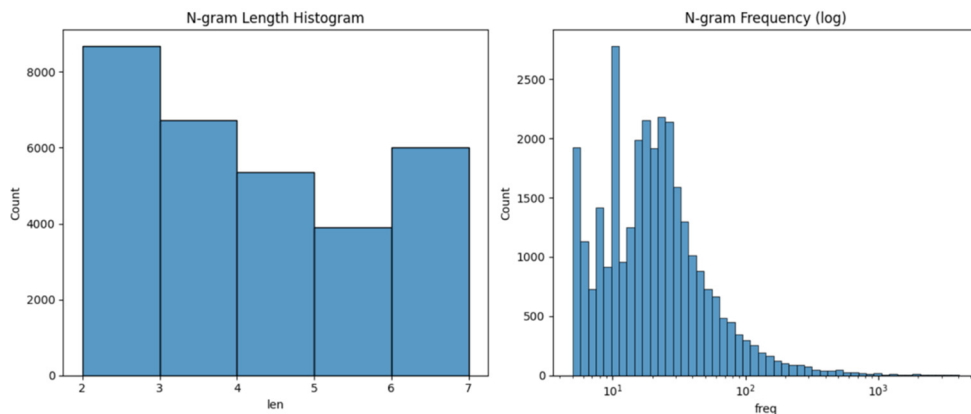


Figure 3. Left: Histogram of n-gram lengths (from 2 to 6). Right: Frequency distribution of n-grams on a log scale.

The final synthetic corpus includes 18,905 distinct n -grams and 112,734 aligned sentence contexts. This coverage ensures a balanced representation of common phrases and rare expressions. Figure 3 visualises the distribution statistics: the left panel shows the length histogram, while the right panel presents the frequency distribution on a logarithmic scale.

In Figure 3, the length histogram (left) confirms that shorter n -grams (2–3 words) dominate: frequent collocations tend to be short, whereas longer expressions are rarer but often more meaningful for context. The tail of longer n -grams (6–7) ensures the model is exposed to multi-word units that help resolve ambiguity in real sentences. The frequency plot (right) shows a clear heavy-tail distribution: most n -grams occur fewer than 100 times, but a small subset appears thousands of times. This motivates the frequency thresholding and quota sampling strategy, which reduces oversampling of extremely common short collocations and preserves rare but linguistically valid phrases [15].

Together, these evidence indicates that our synthetic pipeline yields a corpus faithful to natural language frequency laws while still covering the diversity needed to train a robust restoration model.

2.2. Restoration pipeline

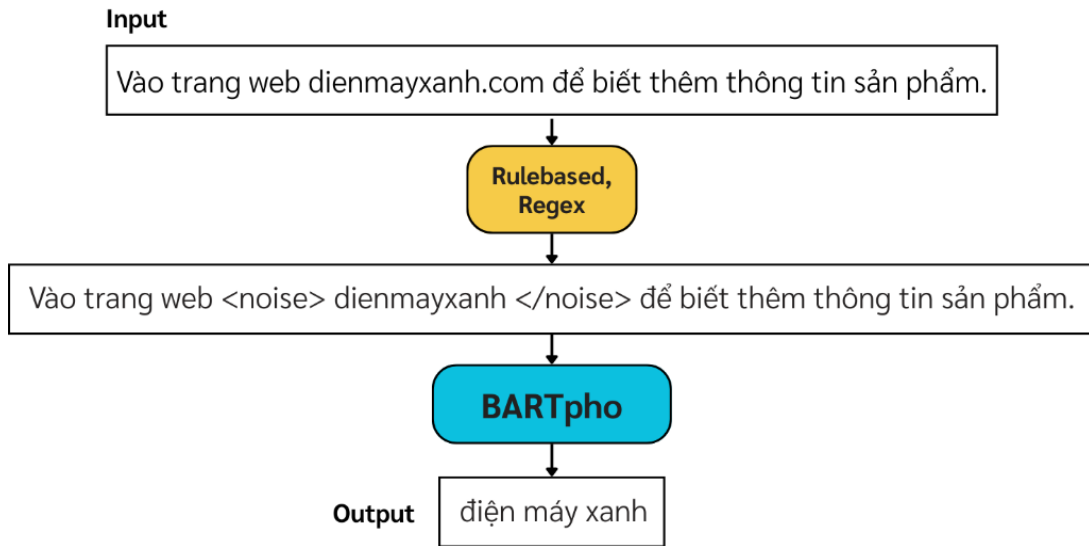


Figure 4. Proposed end-to-end restoration pipeline. A rule-based detector locates the noisy phrase (here “dienmayxanh”), wraps it with `<noise>` tags, and re-inserts it into the original sentence. BARTpho then generates the clean span “điện máy xanh” conditioned on full sentence context

Our restoration module builds on BARTpho’s syllable-level sequence-to-sequence architecture [13]. Unlike vanilla mBART, BARTpho is pre-trained exclusively on Vietnamese, which makes its representations more robust to local spelling variations and word order. Previous diacritic restorers often operate at the phrase level or treat each corrupted token in isolation. This neglects crucial syntactic and semantic cues: Vietnamese tone marks can flip word meaning entirely, so the surrounding words often disambiguate homographs. For example, the string `dienmayxanh` without context could mean anything, but when embedded in “Vào trang web `dienmayxanh` để biết thêm thông tin” the broader sentence clarifies it must be the well-known brand “Điện Máy Xanh”.

By explicitly feeding the *whole sentence* and highlighting the noisy span using `<noise>` markers, the model can:

- Leverage long-range grammatical signals (subject–verb agreement, collocation).
- Better handle named entities and domain-specific terms absent in dictionaries.
- Avoid over-correcting or hallucinating words outside the corrupted region.

This synergy between local and global context turns a generic pre-trained encoder–decoder into a targeted text normaliser that mimics how human readers fix garbled input: they do not guess a word in isolation but interpret it coherently with the entire sentence.

Figure 4 illustrates the complete workflow: Integrated detection–restoration pipeline. A simple rulebased or regex detector scans each sentence for suspicious merged or unaccented tokens. Detected spans are wrapped in <noise> tags and the modified sentence is then passed to the fine-tuned BARTpho. The decoder restores the correct diacritics and inserts missing word boundaries, producing clean, TTS-ready output in a single forward pass.

This architecture is designed to enhance disambiguation through contextual information, preserve the integrity of already correct text, and maintain modularity, allowing for seamless integration of improved detectors without needing to retrain the generator. Empirical results 4 show that context-aware decoding reduces unintended substitutions and improves both word error rate (WER) and sentence fluency.

3. EXPERIMENTS

3.1. Experimental Setup

To evaluate our method, we compare three models:

- Bi-LSTM (baseline): A strong sequence labeller trained on the same corpus without explicit span masking or context injection.
- BARTpho (no context): A standard BARTpho fine-tuned to restore corrupted spans but without feeding the surrounding sentence.
- BARTpho with context: Our proposed model, fine-tuned on full sentences with explicit <noise> tags to highlight the span.

All models are trained on the same synthetic corpus (2.1), with identical tokenisation and maximum sequence lengths. Inference uses on a single NVIDIA RTX 4070.

3.2. Evaluation Metrics

We report four complementary metrics:

- WER (Word Error Rate): Proportion of words incorrectly restored.
- CER (Character Error Rate): Finer-grained version at the character level.
- BLEU: Measures n-gram overlap between predicted and reference spans.
- ROUGE-L: Measures longest common subsequence overlap.

Lower WER/CER and higher BLEU/ROUGE indicate better performance.

3.3. Experimental results

Table 1. Comparison between the Bi-LSTM baseline and BARTpho without context.
Lower WER/CER and higher BLEU/ROUGE-L indicate better performance

Model	WER (↓)	CER (↓)	BLEU (↑)	ROUGE-L (↑)
Bi-LSTM (baseline)	0.43	0.31	39.20	0.69
BARTpho no-context	0.15	0.08	60.52	0.92

Table 2. Comparison of BARTpho with and without added context.
Lower WER/CER and higher BLEU/ROUGE-L indicate better restoration

Model	WER (↓)	CER (↓)	BLEU (↑)	ROUGE-L (↑)
BARTpho	0.15	0.08	60.52	0.93
BARTpho + context	0.09	0.06	91.25	0.95

Table 3. Average inference time per sentence for BARTpho variants.
Higher context usage increases runtime due to longer input sequences

Model	Average Inference Time (s) (↓)
BARTpho	0.01
BARTpho + context	0.21

Table 1 and 2 summarises the restoration accuracy of all models, while Table 3 shows the average inference time.

The Bi-LSTM baseline achieves 43% WER and 31% CER, highlighting the challenge of restoring both diacritics and spacing with purely sequential labelling. Switching to BARTpho without context reduces WER to 15% and CER to 8%, confirming the advantage of span-level generation. When full sentence context is added, WER drops further to 9% and CER to 6%, while BLEU jumps from 60.52 to 91.25 and ROUGE-L reaches 0.95, demonstrating that global syntax cues help disambiguate difficult cases.

However, this improvement comes at a cost: inference time rises from 0.01s (no context) to 0.21s (context) because transformer self-attention scales quadratically with tokens, memory pressure rises, and autoregressive decoding becomes progressively slower as more tokens are generated.

For batch offline TTS pipelines or high-quality voice assistants where restoration accuracy is critical, we recommend deploying BARTpho with context despite its higher runtime. For real-time applications with tight latency constraints (such as on-device TTS), the context-free variant still offers robust restoration at minimal cost.

Discussion: To compress a BART-style encoder–decoder for joint restoration without sacrificing accuracy, we recommend knowledge distillation [16] that aligns intermediate encoder layers (patient KD) [18] and distills token/sequence distributions on the decoder [17], optionally augmented with on-policy [19] distillation to curb exposure bias by training on the student’s rollouts. For aggressive targets, combine capacity-gap strategies with post-training quantization and structured pruning to meet strict latency budgets. To address detection errors, replace the pure regex stage with a hybrid detector: retain high-precision patterns for canonical formats (emails/URLs/brands), add a lightweight neural span-tagger to boost recall on open-class noise [20], leverage upstream confidence signals and lexicon mismatches, and use active learning plus a small contextual verifier to suppress false positives. This preserves precision while improving coverage, robustness, and runtime efficiency.

4. CONCLUSION

We introduced a unified, context-aware approach for joint Vietnamese diacritic restoration and word segmentation, addressing key shortcomings of traditional and task-separated methods. Our solution combines realistic synthetic data, span-level fine-tuning of BARTpho, and a lightweight rule-based noise detector. Experiments show substantial improvements in word error rate, hallucination reduction, and disambiguation of homographs and named entities. Nonetheless, limitations remain in rule-based span detection, data coverage, and inference efficiency. Future work should explore learned noise detectors, expanded corruption patterns, and model compression for real-time applications. Overall, this framework offers a strong foundation for robust Vietnamese text normalization in TTS and beyond. To the best of our knowledge, this is the first end-to-end framework that jointly tackles Vietnamese spacing and diacritic recovery, providing a novel foundation for robust text normalization in TTS and related applications.

In our future work, we will replace the rule-only detector with a hybrid ML approach using active learning, compress the BART-style model via patient/on-policy distillation with pruning and quantization for low latency, optimize inference (caching, batching, beam tuning), and run end-to-end TTS evaluations to guide deployment.

6. REFERENCES

- [1] Vu, Nguyen, Le, Pham, and Nguyen. 2025. Zero-Shot Text-to-Speech for Vietnamese: PhoAudiobook Dataset and Benchmark. *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (ACL 2025 Short Papers)*.
- [2] Nguyen, Le, and Nguyen. 2024. ViLexNorm: A Lexical Normalization Corpus for Vietnamese Social Media Text. *arXiv preprint arXiv:2401.16403*.
- [3] Nguyen, Nguyen, and Nguyen. 2024. A Weakly Supervised Data Labeling Framework for Machine Lexical Normalization in Vietnamese Social Media. *arXiv preprint arXiv:2409.20467*.
- [4] Le, Lam, and Nguyen. 2025. A Survey of Vietnamese Document Analysis and Recognition: Challenges and Future Directions. *arXiv preprint arXiv:2506.05061*.
- [5] Nguyen, Vu, and Hoang. 2025. Vietnamese Words Are Not Constructed from Syllables. *Proceedings of the 39th AAAI Conference on Artificial Intelligence (AAAI 2025)*.
- [6] Binh. 2018. Vietnamese news corpus.
- [7] Clark, Bastings, et al. 2022. CANINE: Pre-training an efficient tokenization-free encoder for language representation. *TACL*.
- [8] Dang and Nguyen. 2020. TDP: A hybrid diacritic restoration with transformer decoder. *PACLIC 34*.
- [9] Han and Nguyen. 2019. Deep neural sequence-to-sequence models for Vietnamese diacritic restoration. *ICT Conference*.
- [10] Hung. 2019. Vietnamese diacritics restoration using deep learning approach. *International Journal of Advanced Computer Science*.
- [11] Nguyen, Nguyen, Phan, Nguyen, and Ha. 2006. Vietnamese word segmentation with CRFs and SVMs: An investigation. *PACLIC 20*.
- [12] Nguyen and Nguyen. 2020. PhoBERT: Pre-trained language models for Vietnamese. *Findings of EMNLP*.
- [13] Tran, Le, and Nguyen. 2021. BARTpho: Pre-trained sequence-to-sequence models for Vietnamese. *Interspeech*.
- [14] Xue, Constant, Roberts, et al. 2022. ByT5: Towards a token-free future with pretrained byte-to-byte models. *TACL*.
- [15] Zipf. 1949. *Human Behavior and the Principle of Least Effort*. Addison–Wesley.
- [16] Hinton, G., Vinyals, O., Dean, J. 2015. Distilling the Knowledge in a Neural Network. *NIPS Deep Learning Workshop*.
- [17] Kim, Y., Rush, A. M. 2016. Sequence-Level Knowledge Distillation. *Proceedings of EMNLP*.
- [18] Sun, S., Cheng, Y., Gan, Z., Liu, J. 2019. Patient Knowledge Distillation for BERT Model Compression. *Proceedings of EMNLP-IJCNLP*.
- [19] Kumar, A., et al. 2024. On-Policy Distillation of Language Models. *Proceedings of ICLR*. discussion 2.
- [20] Do, Dinh-Truong; Nguyen, Ha Thanh; Bui, Thang Ngoc; Vo, Dinh Hieu. 2021. *VSEC: Transformer-based Model for Vietnamese Spelling Correction*.

EFFICIENT VIETNAMESE TEXT-TO-SPEECH WITH ENGLISH BORROWINGS VIA SYNTHESIZED HYBRID SPEECH DATA

Trần Hồng Nhật^{1,2}, Đỗ Văn Hải³

¹Hanoi University of Science and Technology, ²Viettel AI,

³Thuyloi University

ABSTRACT

Text-to-Speech (TTS) technology plays a crucial role in various applications such as virtual assistants, customer service automation, and assistive devices. In the Vietnamese language context, incorporating English borrowings into TTS models is essential, especially in sectors like technology and business. However, existing Vietnamese speech datasets, particularly those from single speakers, often lack sufficient English content, limiting a model's ability to generate natural and accurate bilingual speech. This paper addresses the issue by using a zero-shot multilingual TTS model, XTTS, to generate hybrid Vietnamese-English speech data, which was then used to fine-tune the FastSpeech 2 model. The fine-tuned model demonstrated improved accuracy in pronouncing English borrowings within Vietnamese speech while significantly reducing the computational load, achieving a sixteen-time increase in efficiency compared to XTTS. This approach offers a lightweight and effective solution for developing Vietnamese TTS systems capable of handling English borrowings in real-world applications.

1. INTRODUCTION

In recent years, Text-to-Speech (TTS) technology has become a critical component in various applications, including virtual assistants, automated customer service, and assistive devices for individuals with visual impairments. High-quality TTS systems significantly enhance user experience by providing natural and intelligible speech synthesis, thereby improving communication and accessibility. In the Vietnamese language context, the integration of English borrowings is increasingly common, particularly in sectors such as technology and business. However, existing Vietnamese speech datasets, particularly those collected from a single speaker, often fail to capture the linguistic nuances required for effective bilingual speech synthesis. These datasets typically contain limited English content, restricting the model's ability to learn and replicate English pronunciations accurately. Moreover, the speaker's limited English proficiency and inconsistent accent can lead to unnatural or inaccurate pronunciation, creating challenges for models aiming to deliver high-quality bilingual TTS.

A potential solution to this problem is to leverage multilingual zero-shot TTS models (ZS-TTS), such as XTTS [1] or VALL-E-X [6], which are capable of synthesizing high-quality speech across multiple languages with minimal additional training. Despite their effectiveness, these models are considerably larger and more computationally intensive compared to traditional TTS models like FastSpeech 2 [3] or StyleTTS [4]. This makes them less suitable for deployment in scenarios that require efficient and real-time inference or adaptation to new speakers.

To address these challenges, this paper proposes a novel approach: using fine-tuned ZS-TTS models to generate hybrid Vietnamese-English speech data, which is then used to train a more efficient traditional TTS model, such as FastSpeech 2. This method not only mitigates the limitations of existing datasets but also enables the development of lightweight TTS models capable of accurately pronouncing English borrowings while maintaining high-quality Vietnamese speech synthesis.

The contributions of this paper are summarized as follows: (1) A pipeline was developed to generate hybrid Vietnamese-English speech data using a large language model in conjunction with the XTTS multilingual TTS system. (2) This synthesized data was used to train a FastSpeech 2 model, enabling it to accurately pronounce common English borrowings within Vietnamese speech. (3) The resulting model is sixteen times more efficient and thirty times smaller than the XTTS model, making it highly suitable for practical deployment.

2. RELATED WORK

Recent advancements in Text-to-Speech (TTS) technology have focused on improving speech quality and introducing multilingual capabilities. Models like StyleTTS and FastSpeech 2 [3] have been adopted for their efficiency, but they struggle with multilingual and code-switching tasks. To address this, zero-shot multilingual TTS models like XTTS [1] and VALL-E-X [6] have emerged, capable of generating speech in multiple languages with minimal fine-tuning. However, these models are large and resource-intensive, making them impractical for real-time use.

In Vietnamese TTS, little work has been done to address the challenge of synthesizing speech with English borrowings, which are common in Vietnamese, especially in domains like business and technology. Existing Vietnamese datasets lack significant English content, limiting the ability of monolingual models to handle such borrowings. Some efforts have explored synthetic data generation using large language models to supplement underrepresented linguistic features, but this remains underutilized in the Vietnamese-English context.

This paper builds on these approaches by using a finetuned XTTS model to generate hybrid Vietnamese-English speech, which is then used to train a more efficient FastSpeech 2 model, addressing both dataset and computational challenges.

3. METHODOLOGY

3.1. Prerequisites

In this subsection, we provide an overview of the XTTS - Zero-shot Multilingual TTS model and the FastSpeech 2 TTS model, highlighting their key features and relevance to this paper.

3.1.1 XTTS model

XTTS [1] is a multilingual, zero-shot, multi-speaker TTS model designed to adapt efficiently to new languages and speakers. Pretrained on 27,000 hours of multilingual data, including 14,000 hours of English, XTTS excels in generating high-quality English speech. As illustrated in Fig. 1, the model architecture consists of two primary components: a GPT-based text encoder and a Vector Quantized-Variational Autoencoder (VQ-VAE) [5].

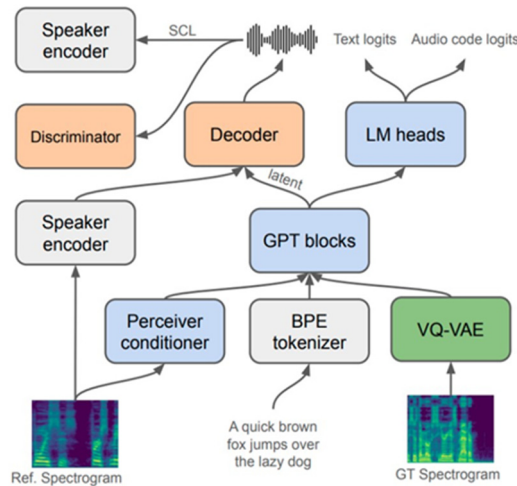


Figure 1. Overall architecture of XTTS [1]

During training, the model processes text, reference audio, and ground truth audio, which is converted into mel-spectrograms. The VQ-VAE encodes the mel-spectrogram into tokens, while the GPT-2 encoder predicts these VQ-VAE codes from the input text tokens. A conditioning encoder extracts speaker-specific features from the reference audio to capture voice characteristics, and a HiFi-GAN [2] vocoder converts the predicted tokens back into waveforms. A speaker embedding model computes speaker consistency loss (SCL) to ensure similarity between the generated speech and the reference speaker.

During inference, XTTS generates speech using the input text and reference speaker audio. Although not originally trained on Vietnamese, XTTS can be fine-tuned on Vietnamese datasets to handle both Vietnamese speech and English borrowings. Despite its ability to produce high-quality outputs, XTTS's large size and slow inference limit its practicality for real-time applications. Additionally, the absence of a duration predictor occasionally leads to abnormal speech durations and gibberish at the end of outputs.

3.1.2 FastSpeech 2 model

FastSpeech 2 [3] is a lightweight, non-autoregressive text-to-speech (TTS) model designed for fast and efficient speech synthesis. Unlike autoregressive models such as CosyVoice [7] that generate speech sequentially, FastSpeech 2 synthesizes speech frames in parallel, drastically improving inference speed while maintaining high speech quality. Fig. 2 illustrates the overall architecture of FastSpeech 2.

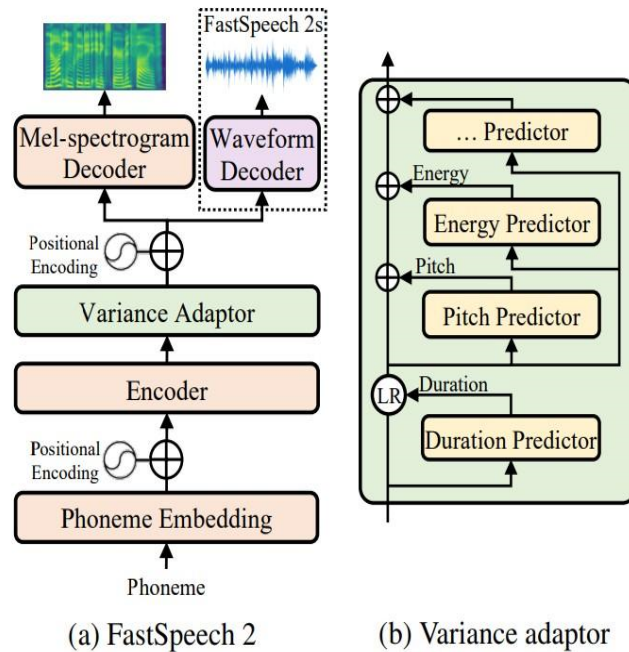


Figure 2. Overall architecture of FastSpeech 2

The model comprises several key components: a phoneme encoder, a variance adaptor, and a melspectrogram decoder. The phoneme encoder processes the input text, converting it into phoneme-level representations. The variance adaptor then introduces critical variation factors such as duration, pitch, and energy, allowing for more expressive and natural-sounding speech. Finally, the melspectrogram decoder generates a mel-spectrogram, which is then converted into speech using a vocoder like HiFiGAN [2].

FastSpeech 2's non-autoregressive nature allows for much faster inference compared to autoregressive models, making it ideal for real-time applications. Additionally, its smaller size ensures greater deployment efficiency, especially in monolingual TTS tasks like Vietnamese. When fine-tuned with hybrid Vietnamese-English data, FastSpeech 2 can generate high-quality Vietnamese speech while accurately pronouncing English borrowings.

3.2. Hybrid data generation

To generate hybrid Vietnamese-English speech data, we began by selecting a diverse corpus of Vietnamese text. From this corpus, we filtered the most common English borrowings. Using these borrowings, we then constructed a Vietnamese text dataset where each sentence contained at least one out-of-vocabulary (OOV) English borrowing. This dataset was generated by using an LLM via a custom prompt. Next, we trained the XTTS model on a single-speaker Vietnamese dataset to ensure it could accurately synthesize Vietnamese speech in the speaker’s voice. Once trained, using the original Vietnamese speech data as reference, we employed XTTS to synthesize speech from the constructed sentences, producing the final hybrid Vietnamese-English training data. For more details on the datasets used, please refer to the Dataset section.

Additionally, since FastSpeech 2 learns to speak based on phonemes, the FastSpeech model trained on pure Vietnamese and common English words cannot adapt well to unseen phoneme patterns in Vietnamese that are prevalent in English such as double or triple phonemes. To tackle this issue, we also generate hybrid data based on words containing these sequences to allow the model to be able to speak unseen English words with these sequences.

3.3. FastSpeech 2 Training

In this subsection, we describe the training process for the FastSpeech 2 models: FS-baseline and FS-hybrid. The FS-baseline model was trained on a dataset of Vietnamese speech from a single female speaker. This dataset provided a robust foundation for learning Vietnamese phonetic and prosodic features.

For the FS-hybrid model, the dataset was extended by incorporating synthetic hybrid Vietnamese-English speech. This hybrid dataset was generated using the XTTS model, which was fine-tuned on the same single-speaker Vietnamese dataset. The goal of this additional data was to improve the model’s ability to handle English borrowings in Vietnamese speech, enhancing pronunciation accuracy and fluency when integrating foreign words into natural sentences.

Both models were trained using the same FastSpeech 2 architecture with identical hyperparameters and a standard training-validation split. The FS-hybrid model, however, leveraged the hybrid dataset to develop the capability to synthesize accurate Vietnamese-English speech, while FS baseline focused solely on Vietnamese content. For results of these models’ comparison, please refer to the experiment details and results section.

4. EXPERIMENTS

4.1. Dataset

4.1.1 Text corpus

The text corpus for this paper was built from two main sources: transcriptions of YouTube speech content and text crawled from various online newspapers. The data was manually labeled by annotators to ensure quality. From this corpus, we filtered out the 200 most common out-of-vocabulary (OOV) English borrowings, which became the foundation for generating hybrid Vietnamese-English speech data.

Half of the 5,000 synthetic sentences were generated using these filtered OOVs by prompting a large language model (Gemini-pro) to naturally incorporate the English borrowings into Vietnamese sentences. The remaining half of the sentences were designed to include complex phoneme sequences uncommon in Vietnamese (e.g., double or triple consonants). This approach created a comprehensive dataset with both typical borrowings and challenging phoneme patterns, resulting in approximately five hours of hybrid speech data.

The testing dataset followed a similar method. It was split into three parts: 500 sentences featuring common English borrowings, 500 sentences with more difficult phoneme sequences and 500 pure Vietnamese sentences. These three subsets were used to create the test sets, referred to as `test_common`, `test_difficult`, and `test_viet`.

4.1.2 Speech datasets

The speech data used in this paper consists of 10 hours of high-quality, studio-recorded Vietnamese audio from a single female speaker. This dataset was employed to train the baseline FastSpeech 2 model and fine-tune the XTTS model. Once fine-tuned, the XTTS model was used to generate 5 hours of hybrid Vietnamese-English speech based on the hybrid text mentioned earlier. The resulting 15 hours of combined speech data (10 hours of Vietnamese and 5 hours of hybrid) were then used to train the hybrid FastSpeech 2 model. A 9:1 training-validation split was applied consistently throughout the training process to ensure model stability and robust performance.

4.2. Speech quality analysis

In this subsection, we compare the performance of the XTTS model, the FastSpeech 2 baseline (FS-baseline), and the FastSpeech 2 model trained on hybrid data (FS-hybrid) using the Word Error Rate (WER) as the primary metric. WER objectively measures how accurately the models generate speech, particularly focusing on their ability to handle English borrowings in Vietnamese. Following XTTS, we also employed the open-source UTMOS model [9] to predict the Naturalness Mean Opinion Score (nMOS) on a 1-5 scale.

Table 1. WER and UTMOS of XTTS and FastSpeech 2 models on the test sets

Model	test_common (WER %)	test_diff (WER %)	test_viet (WER%)	UTMOS (overall ↑)
XTTS	2.54	2.13	1.38	3.43
FS-baseline	4.94	9.86	0.39	2.95
FS-hybrid (ours)	4.46	4.86	0.44	2.92

As shown in Table 1, FS-hybrid improves WER over both XTTS and the FS-baseline, with the biggest gains on **test_difficult** (harder English borrowings). On **test_common**, where borrowings are simpler, the gap to the FS-baseline is smaller but still in FS-hybrid’s favor. While XTTS records the lowest overall WER, it has practical drawbacks, including a large model size, odd duration behavior, and occasional content inconsistencies.

FS-hybrid strengthens handling of English borrowings while keeping the efficiency benefits discussed earlier. It also maintains robustness on Vietnamese speech, matching or exceeding XTTS on **test_viet** in WER. By contrast, UTMOS favors XTTS in predicted perceptual quality. Overall, FS-hybrid offers a good balance of accuracy and efficiency for code-switching use cases.

4.3. Performance Analysis

In this section, we compare the FS-hybrid model with the XTTS model in terms of model size and inference speed, which are key factors for practical deployment. As presented in Table 2, the FS-hybrid model is nearly 30 times smaller than the XTTS model, allowing for more efficient use of computational resources. This smaller size contributes to a significantly faster training process, as FS-hybrid completes more epochs per hour due to its streamlined architecture. Moreover, FS-hybrid outperforms XTTS in inference speed, offering a 16x improvement, making it far more suitable for real-time speech synthesis applications.

Table 2. Performance comparison of XTTS and FS-hybrid

Model	Inference time (s)	#Params
XTTS	1.44	480M
FS-hybrid (ours)	9×10^{-2}	16M

5. CONCLUSION

This report presents an efficient method for building a Vietnamese TTS model that handles English borrowings. We use XTTS to generate hybrid Vietnamese–English data and fine-tune a FastSpeech 2 model to synthesize both languages seamlessly. This approach addresses the lack of English in existing Vietnamese datasets and yields a model roughly ten times smaller and faster than XTTS, making it suitable for real-time use. Overall, it offers a practical, scalable path to improve Vietnamese TTS with integrated English content while improving both performance and deployment efficiency.

6. REFERENCES

- [1] Edresson Casanova, et al. XTTS: a massively multilingual zero-shot text-to-speech model. arXiv preprint arXiv:2406.04904, 2024.
- [2] Jungil Kong, et al. Hifigan: Generative adversarial networks for efficient and high fidelity speech synthesis. *Advances in neural information processing systems*, 33:17022–17033, 2020.
- [3] Yi Ren, et al. FastSpeech 2: Fast and high-quality end-to-end text to speech. arXiv preprint arXiv:2006.04558, 2020. 1, 2.
- [4] Li, Yinghao Aaron, et al. "Styletts 2: Towards human-level text-to-speech through style diffusion and adversarial training with large speech language models." *Advances in Neural Information Processing Systems* 36 (2023): 19594-19621.
- [5] Aaron Van Den Oord, et al. Neural discrete representation learning. *Advances in neural information processing systems*, 30, 2017.
- [6] Ziqiang Zhang, et al. Speak foreign languages with your own voice: Cross-lingual neural codec language modeling. arXiv preprint arXiv:2303.03926, 2023.
- [7] Du, Zhihao, et al. "Cosyvoice: A scalable multilingual zero-shot text-to-speech synthesizer based on supervised semantic tokens." arXiv preprint arXiv:2407.05407 (2024).
- [8] Saeki, T., Xin, D., Nakata, W., Koriyama, T., Takamichi, S., Saruwatari, H. (2022) UTMOS: UTokyo-SaruLab System for VoiceMOS Challenge 2022. *Proc. Interspeech 2022*, 4521-4525, doi: 10.21437/Interspeech.2022-439.

PHƯƠNG PHÁP ADVANTAGE ACTOR-CRITIC CHO ĐIỀU HƯỚNG UAV LIÊN TỤC TRONG MÔI TRƯỜNG KHE HẸP

Tạ Chí Hiếu, Nguyễn Anh Huy, Phan Thị Phương Anh

Trường Đại học Thủy lợi, email: hieutc@thu.edu.vn

1. GIỚI THIỆU CHUNG

UAV (Unmanned Aerial Vehicle – phương tiện bay không người lái) ngày càng được ứng dụng rộng rãi trong các nhiệm vụ như giám sát công trình, cứu hộ và giao hàng thông minh [1]. Một trong những thách thức lớn là làm sao để UAV có thể điều hướng chính xác và an toàn trong môi trường chật hẹp với các khe hẹp liên tiếp – như hành lang kỹ thuật, khu công nghiệp hoặc đô thị đông đúc.

Trong các hướng tiếp cận bằng học tăng cường sâu (Deep Reinforcement Learning – DRL), thuật toán Proximal Policy Optimization (PPO) thường được sử dụng để huấn luyện UAV trên các hành trình dài được chia thành từng đoạn [1]. Tuy nhiên, trong các môi trường có cấu trúc lặp lại và không bị phân đoạn như chuỗi khe hẹp liên tục, việc tái sử dụng chính sách theo đoạn không còn hiệu quả.

Đề tài này lựa chọn thuật toán **Advantage Actor-Critic (A2C)** – một phương pháp học tăng cường theo chính sách – nhằm đánh giá khả năng học chính sách điều khiển **ổn định và toàn cục** cho UAV khi di chuyển trong môi trường có khe hẹp tuần hoàn [2]. Việc lựa chọn A2C xuất phát từ ưu điểm về tính ổn định trong cập nhật và khả năng thích nghi tốt với môi trường có cấu trúc lặp lại. Qua đó, nghiên cứu hướng đến đánh giá tiềm năng của A2C trong các ứng dụng UAV đòi hỏi tính bền vững và liên tục trong điều hướng.

2. PHƯƠNG PHÁP NGHIÊN CỨU

2.1. Mô hình vật lý UAV và môi trường mô phỏng

UAV được mô hình hóa trong mặt phẳng 2D với các biến trạng thái gồm vị trí (x, z) , góc nghiêng (θ) , vận tốc theo phương ngang và phương thẳng đứng $(\dot{x}, \dot{z})$ và tốc độ góc $(\dot{\theta})$. Độ cao được điều khiển bằng PID (Proportional Integral Derivative – bộ điều khiển tỷ lệ tích phân vi phân), còn hướng bay điều chỉnh thông qua mô-men quay τ [3].

Mô hình động học được mô tả bằng các phương trình sau:

$$\ddot{x} = \frac{T \sin \theta}{m} - \frac{k_d}{m} \dot{x} \quad (1)$$

$$\ddot{z} = \frac{T \cos \theta}{m} - g - \frac{k_d}{m} \dot{z} \quad (2)$$

$$\ddot{\theta} = \frac{\tau}{I} \quad (3)$$

trong đó: T là lực đẩy tác động theo trục thân UAV, τ là moment quay, m là khối lượng UAV, I là mô men quán tính, g là gia tốc trọng trường và k_d là hệ số cản không khí.

Để duy trì độ cao ổn định, bộ điều khiển PID tính toán lực đẩy như sau:

$$T = k_p e_z + k_i \int e_z dt + k_d \frac{de_z}{dt} \quad (4)$$

với các hệ số của bộ điều khiển PID: $k_p = 10.0$, $k_i = 0.03$, $k_d = 4.0$, $e_z = z_{\text{target}} - z$ là sai số độ cao tại thời điểm z và lực đẩy (thrust) được giới hạn trong khoảng $[20, 35]$ N. Trong đó:

- Quá trình huấn luyện, nhận thấy lực tối thiểu 20N đủ thắng trọng lực và lực tối đa 30N tránh quá tải.
- Thành phần tỷ lệ $k_p e_z$: Tạo lực đẩy tỷ lệ thuận với độ lệch so với mục tiêu
- Thành phần tích phân $k_i \int e_z dt$: Loại bỏ sai số dư do nhiễu bên ngoài với hệ số nhỏ $k_i = 0.03$
- Thành phần vi phân $k_d \frac{de_z}{dt}$: Dự đoán xu hướng thay đổi độ cao và tạo lực cản để giảm dao động, giúp ổn định nhanh hơn.
- Hàm e_z đại diện cho sai số về độ cao của UAV so với độ cao mục tiêu trong môi trường 2D. Cụ thể, e_z là một hàm đo độ lệch giữa độ cao thực tế của UAV và độ cao mong muốn, dùng để tính toán phần thưởng r_h nhằm khuyến khích UAV duy trì độ cao ổn định trong quá trình bay. Hàm này có thể được định nghĩa là hàm khoảng cách tuyệt đối hoặc bình phương sai số về độ cao, giúp đưa ra tín hiệu điều chỉnh lực đẩy để UAV không bị lệch quá cao hoặc quá thấp so với mục tiêu. Điều này rất quan trọng trong môi trường 2D bởi vì ngoài việc di chuyển ngang, UAV cũng cần kiểm soát chặt chẽ chiều cao để tránh va chạm hoặc rơi, đặc biệt khi bay qua các khe hẹp giữa chướng ngại vật.

Môi trường mô phỏng có dạng “FlappyDroneEnv”, thiết lập $x_{\text{target}} = 100\text{m}$, với các chướng ngại vật hình chữ nhật bố trí cách nhau 10m. Không gian quan sát là một vector 10 chiều:

$$\left[\frac{x}{x_{\text{target}}}, \frac{z-5.0}{5.0}, \frac{\theta}{\pi}, \frac{\dot{x}}{5.0}, \frac{\dot{z}}{2.0}, \frac{\dot{\theta}}{2.0}, \frac{dx}{x_{\text{target}}}, \frac{\text{sign}(dx)\dot{x}}{5.0}, \frac{d_{\text{obs}}}{x_{\text{target}}}, \frac{z-z_{\text{gap}}}{5.0} \right] \quad (5)$$

trong đó $dx = x_{\text{target}} - x$, d_{obs} là khoảng cách đến chướng ngại vật tiếp theo, và $z_{\text{gap}} \in [4.0, 6.0]\text{m}$ là tâm khe hở. Không gian hành động là một mô-men xoắn duy nhất $\tau \in [-0.6, 0.6]$, được chuẩn hóa về khoảng $[-1.2, 1.2]$ N·m.

2.2. Hàm phần thưởng

Hàm phần thưởng là tổng có trọng số của các thành phần thúc đẩy tiến độ di chuyển, duy trì độ cao ổn định và tránh va chạm với chướng ngại vật:

$$r_t = r_h + r_p + r_v + r_{\text{gap}} + r_{\theta} + r_{\text{term}} \quad (6)$$

Trong đó:

- r_h : giữ UAV ở gần độ cao mục tiêu (5m);
- r_p : thúc đẩy tiến độ hướng tới mục tiêu;
- r_v : duy trì vận tốc ổn định;
- r_{gap} : cản giữa UAV vào vị trí khe hẹp;
- r_{θ} : kiểm soát độ nghiêng;

$$- r_{\text{term}} = \begin{cases} r_{\text{target}} = 3000 & \text{nếu khoảng cách } |dx| < 0.5\text{m} \\ r_{\text{coll}} = -1000 & \text{nếu UAV va chạm với chướng ngại vật.} \\ r_{\text{altitude}} = -1000 & \text{nếu độ cao } z < 1\text{m hoặc } z > 14\text{m} \\ r_{\text{angle}} = -500 & \text{nếu góc nghiêng } |\theta| > \pi / 2.12 \text{ (} 85^\circ \text{)} \end{cases}$$

Phần thưởng r_{term} là phần thưởng hoặc phạt được tính vào cuối mỗi tập (episode) bay của UAV, nhằm củng cố kết quả của toàn bộ hành trình bay. Cụ thể:

- Nếu UAV bay thành công, tức là tiếp cận mục tiêu với sai số nhỏ theo trục hoành (trong môi trường 2D), thì r_{term} sẽ có giá trị dương lớn để khuyến khích hành vi đó.
- Ngược lại, nếu UAV va chạm vào chướng ngại vật, vi phạm giới hạn an toàn về độ cao hoặc góc nghiêng, hoặc không đạt mục tiêu trong thời gian quy định, thì r_{term} sẽ là một phần phạt nặng, nhằm ngăn chặn các hành động nguy hiểm hoặc không hiệu quả.

Phần thưởng này đóng vai trò như một phần đánh giá tổng thể cho cả quãng đường bay, không chỉ dựa trên từng bước di chuyển mà còn nhằm đảm bảo UAV hoàn thành nhiệm vụ một cách an toàn và hiệu quả.

Khác với bài PPO chia phần thưởng theo đoạn, hàm phần thưởng A2C được áp dụng toàn cục trên một đoạn 100m duy nhất, nhấn mạnh **sự ổn định liên tục** thay vì chuyển đổi chính sách giữa các đoạn.

2.3. Thuật toán A2C và quy trình huấn luyện

Thuật toán Advantage Actor-Critic (A2C) được triển khai dưới dạng cập nhật đồng bộ theo chính sách (on-policy), sử dụng kết hợp hai mạng: Actor (chính sách) và Critic (ước lượng giá trị), với cấu trúc mạng đa lớp (MLP), cập nhật đồng thời cả module hành động và module đánh giá đều được cập nhật cùng lúc. Không dùng replay buffer như SAC, A2C yêu cầu cập nhật ngay sau mỗi batch, dẫn tới tốc độ hội tụ chậm hơn nhưng ổn định hơn trong môi trường có cấu trúc định kỳ [2]. Quá trình huấn luyện diễn ra trong 1 triệu bước thời gian, có checkpoint lưu mô hình và đánh giá dựa trên 100 tập kiểm thử.

Thuật toán Advantage Actor–Critic (A2C) cho điều hướng UAV:

- 1: Khởi tạo môi trường UAV (FlappyDroneEnv với độ cao mục tiêu 5.0m).
- 2: Khởi tạo mạng Actor-Critic:
 - Mạng Actor dự đoán chính sách $\pi_\theta(a|s)$.
 - Mạng Critic ước lượng giá trị trạng thái $V_w(s)$.
 - Các mạng có kiến trúc: 2 lớp ẩn MLP, mỗi lớp 256 nơ-ron, hàm kích hoạt ReLU.
- 3: Lặp cho đến đủ số bước huấn luyện (ví dụ 1 triệu):
 - Thu thập một batch các trải nghiệm bằng cách tương tác với môi trường theo chính sách hiện tại.
 - Tính toán lợi ích (advantage) $A_t = R_t - V_w(s_t)$, trong đó R_t là phần thưởng thực nhận sau bước t [4].
 - Cập nhật mạng Critic bằng cách giảm thiểu mất mát: $L(w) = (R_t - V_w(s_t))^2$ [5].
 - Cập nhật mạng Actor bằng cách tối đa hóa hàm mục tiêu: $J(\theta) = \mathbb{E}_t[\log \pi_\theta(a_t|s_t)A_t]$ [5].
 - Cập nhật tham số mạng bằng gradient ascent/descent theo hàm mất mát/lợi ích tương ứng.

4: Định kỳ đánh giá mô hình trên tập thử nghiệm, lưu mô hình tốt nhất.

5: Kết thúc, xuất mô hình đã huấn luyện.

Thuật toán A2C cho phép UAV học chính sách điều khiển liên tục, thích nghi tốt với môi trường có sự thay đổi và tính phức tạp của chướng ngại vật khe hẹp.

3. KẾT QUẢ NGHIÊN CỨU

Sau khi huấn luyện, mô hình đạt hiệu suất cao nhất (dựa trên tổng phần thưởng trong giai đoạn đánh giá) được kiểm thử trên 100 tập độc lập với các bố trí chướng ngại vật được sinh ngẫu nhiên. Một tập được coi là thành công nếu UAV tiếp cận mục tiêu ở khoảng cách 100m. Các trường hợp thất bại bao gồm: va chạm với chướng ngại vật, vi phạm độ cao cho phép, góc nghiêng vượt quá giới hạn an toàn hoặc kết thúc sớm do vượt quá số bước tối đa.

3.1. Kết quả định lượng

Để đánh giá hiệu quả đào tạo UAV bằng thuật toán Advantage Actor-Critic (A2C), chúng tôi thực hiện huấn luyện và kiểm thử trên 100 tập bay với độ dài quãng đường 100 m trong môi trường 2D có nhiều khe hẹp và chướng ngại vật. Để đảm bảo tính khách quan, kết quả A2C được đối sánh đồng thời với các thuật toán phổ biến gồm PPO và SAC, trong đó các tham số môi trường và quá trình kiểm thử đều giữ nguyên. Kết quả SAC sử dụng dữ liệu thực nghiệm từ báo cáo nội bộ nhóm tác giả (chưa xuất bản), đảm bảo đồng nhất điều kiện so sánh.

Các chỉ số đánh giá gồm: tỉ lệ thành công (success rate), số bước trung bình, tỉ lệ va chạm, tỉ lệ vi phạm an toàn, cũng như giá trị phần thưởng trung bình. Bảng sau trình bày kết quả so sánh định lượng chi tiết giữa các thuật toán.

Bảng 1. Kết quả đánh giá các thuật toán A2C, PPO, SAC trên môi trường 2D trong bài toán điều hướng UAV có khe hẹp

Thuật toán	Tỉ lệ thành công [%]	Số bước trung bình	Tỉ lệ va chạm [%]	Tỉ lệ vi phạm an toàn [%]	Trung bình phần thưởng
A2C	93.0	229.5 ± 36.08	6.0	1.0	1850
PPO	90.5	240.1 ± 38.7	7.5	2.0	1760
SAC	94.0*	129.0 ± 28.2*	4.0*	2.0*	1980*

(*Kết quả SAC từ báo cáo nội bộ nhóm tác giả, chưa xuất bản chính thức nhưng kiểm thử cùng điều kiện thiết lập.)

Phân tích chi tiết cho thấy SAC có ưu thế nổi bật với tỉ lệ thành công cao nhất, số bước trung bình thấp và phần thưởng trung bình lớn hơn so với các phương pháp còn lại. Đường cong học của SAC hội tụ nhanh, ổn định và phân phối phần thưởng đều, trong khi A2C cũng đạt hiệu quả ổn định, hội tụ tốt. PPO mặc dù có khả năng thích nghi nhất định nhưng tỉ lệ va chạm và số bước trung bình đều cao hơn, cho thấy quá trình học chính sách không tối ưu bằng SAC/A2C trong môi trường này.

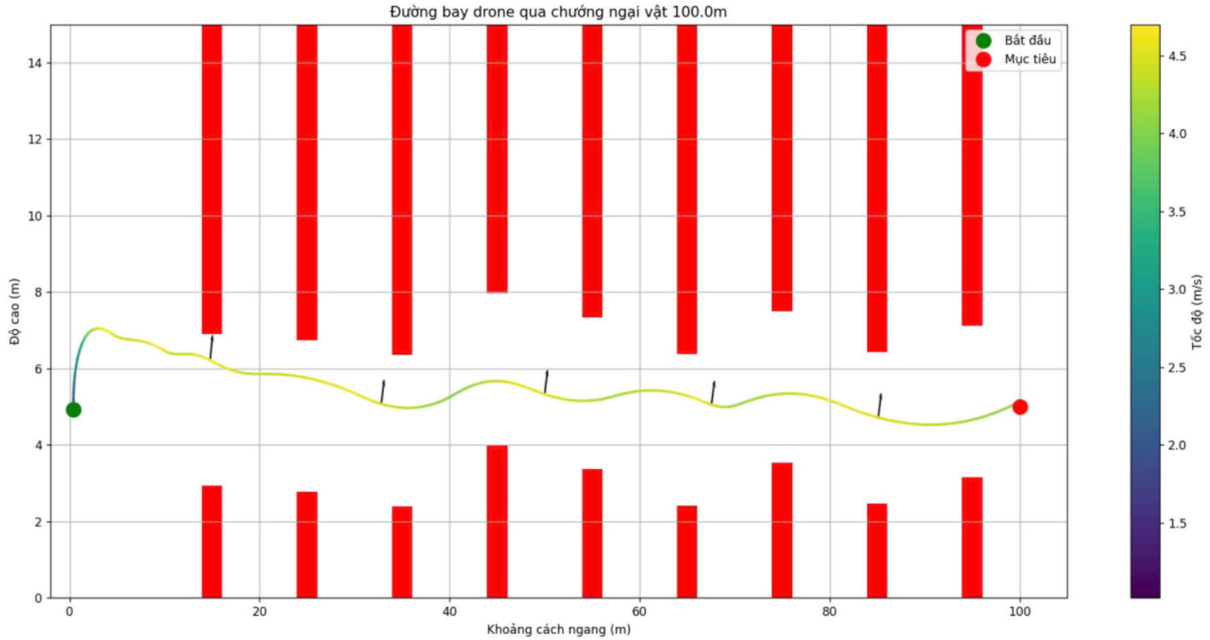
Kết quả này làm cơ sở thuyết phục cho tính hiệu quả của các thuật toán off-policy như SAC trong điều hướng UAV phức tạp, đồng thời xác nhận vai trò của A2C trong các môi trường đòi hỏi sự ổn định về chính sách và tốc độ học tập.

3.2. Phân tích quỹ đạo

Để đánh giá chiến lược tránh chướng ngại vật dựa trên mô hình A2C, tiến hành nhiều mô phỏng trong một môi trường 2D phức tạp, dày đặc các rào chắn thẳng đứng. Drone được huấn luyện nhằm tiếp cận mục tiêu đặt tại khoảng cách 100 m theo phương ngang (trục x), qua đó kiểm tra khả năng thích ứng và điều hướng an toàn trong các bố cục đầy thử thách.

Môi trường bay và quỹ đạo của drone được trực quan hóa trên mặt phẳng 2D, trong đó:

- Điểm xanh lá biểu thị vị trí xuất phát của drone.
- Điểm màu đỏ là vị trí mục tiêu
- Đường đi có màu sắc thay đổi theo vận tốc (sử dụng bảng màu viridis) mô tả quỹ đạo di chuyển của drone.
- Các mũi tên minh họa hướng quay của drone tại các thời điểm khác nhau.
- Hình chữ nhật màu đỏ là các chướng ngại vật



Hình 1. Kết quả điều hướng của máy bay không người lái ở khoảng cách 100 mét

Phân tích quỹ đạo bay từ mô hình A2C cho thấy UAV học được chính sách liên tục và linh hoạt để vượt qua các khe hẹp. Hành động điều chỉnh độ cao, hướng bay, góc nghiêng diễn ra mượt mà và nhất quán, giúp UAV hạn chế va chạm, duy trì nhiệm vụ điều hướng an toàn.

So với các thuật toán phản xạ cục bộ như PPO, A2C thể hiện ưu điểm nổi bật là học được một chính sách tổng quát cho toàn hành trình bay, phù hợp với môi trường có tính tuần hoàn hoặc lặp lại cấu trúc không gian.

4. KẾT LUẬN

Nghiên cứu triển khai và đánh giá hiệu quả thuật toán A2C trong điều khiển UAV bay vượt qua môi trường chướng ngại vật khe hẹp cho thấy A2C có khả năng học chính sách hiệu quả, duy trì tỷ lệ thành công cao và độ ổn định tốt. So với PPO, A2C có ưu điểm về kiểm soát hành vi và khả năng thích nghi trong môi trường tuần hoàn, dù chưa đạt hiệu suất tối ưu như SAC khi sử dụng buffer và entropy regularization.

Kết quả cho thấy A2C có khả năng học được một chính sách điều khiển tương đối hiệu quả, với tỷ lệ thành công 93% trên đoạn 100m. So với PPO và SAC:

- A2C ít yêu cầu bộ nhớ hơn (do không dùng replay buffer),
- Dễ triển khai trong môi trường tuần hoàn do cập nhật trực tiếp theo phân phối trạng thái hiện tại,
- Tuy nhiên, hội tụ chậm hơn và phụ thuộc nhiều vào việc tinh chỉnh phần thưởng và kiến trúc mạng.

Một hạn chế được ghi nhận là A2C có xu hướng bị quá khớp với môi trường huấn luyện nếu cấu trúc khe hẹp thay đổi mạnh, do đặc điểm on-policy không hỗ trợ học từ kinh nghiệm cũ như SAC. Ngoài ra, chính sách học được dễ gặp vấn đề khi gặp khe có khoảng cách/chiều cao thay đổi đột ngột.

Từ kết quả này, có thể đề xuất các hướng mở rộng:

- Tích hợp perception (thị giác máy tính) để UAV học chính sách từ ảnh hoặc lidar,
- So sánh với A2C async (A3C), PPO-LSTM để tăng khả năng ghi nhớ ngữ cảnh,
- Kết hợp A2C với kế hoạch đường đi offline để giảm không gian hành động.

Nghiên cứu đóng góp vào bức tranh toàn cảnh về học tăng cường UAV bằng cách chỉ ra rằng A2C là một lựa chọn cân bằng giữa ổn định, hiệu quả và đơn giản triển khai, đặc biệt thích hợp với các bài toán điều hướng có cấu trúc không gian tuần hoàn hoặc giới hạn địa hình rõ rệt.

5. TÀI LIỆU THAM KHẢO

- [1] T. Ch. Hiếu và P. V. Cường, "Segmented PPO-Based Transfer Learning for 2D UAV Navigation," NSA 2025 Proceeding.
- [2] S.-E. Shen và Y.-C. Huang, "Application of Reinforcement Learning in Controlling Quadrotor UAV Flight Actions," Drones, 2024.
- [3] K. Bouzgou, Y. Bestaoui, L. Benchikh, B. Ibari, và Z. Ahmed-Foitih, "Dynamic Modeling, Simulation and PID Controller of Unmanned Aerial Vehicle UAV," INTECH, Luton, United Kingdom, 2017.
- [4] A. Van de Kleut, "Actor-Critic Methods, Advantage Actor-Critic (A2C) and Generalized Advantage Estimation (GAE)," 2023. [Online]. Available: <https://avandekleut.github.io/a2c/>. [Accessed: Sep. 20, 2025].
- [5] GeeksforGeeks, "Actor-Critic Algorithm in Reinforcement Learning," updated Jul. 23, 2025. [Online]. Available: <https://www.geeksforgeeks.org/machine-learning/actor-critic-algorithm-in-reinforcement-learning/>. [Accessed: Sep. 20, 2025].

ỨNG DỤNG SOFT ACTOR-CRITIC CHO ĐIỀU HƯỚNG UAV 2D TRONG MÔI TRƯỜNG CHƯỚNG NGẠI VẬT SINH NGẪU NHIÊN

Tạ Chí Hiếu, Phan Thị Phương Anh, Nguyễn Anh Huy

Trường Đại học Thủy lợi, email: hieutc@thu.edu.vn

1. GIỚI THIỆU CHUNG

Trong những năm gần đây, sự phát triển nhanh chóng của UAV (Unmanned Aerial Vehicle) đã mở ra nhiều ứng dụng trong các lĩnh vực như giám sát nông nghiệp, giao hàng, khảo sát hạ tầng, và cứu hộ khẩn cấp. Một yêu cầu quan trọng trong các ứng dụng này là khả năng **điều hướng tự động an toàn** trong môi trường phức tạp có nhiều vật cản, đặc biệt trong các khu vực đô thị hoặc không gian hẹp như hành lang, rừng rậm, hoặc nhà máy [1].

Các thuật toán học tăng cường như DQN, PPO, A2C, SAC đã được nghiên cứu để giúp UAV học chính sách điều khiển tối ưu từ dữ liệu môi trường. Trong số đó, **PPO** thường được sử dụng cho các hành trình dài, chia thành từng đoạn huấn luyện [2]. Tuy nhiên, cách tiếp cận chia đoạn này có nhược điểm là phải **reset chính sách học lại** tại mỗi đoạn, làm giảm tính kế thừa và tính liên tục trong hành vi điều khiển [3].

Bài nghiên cứu này áp dụng thuật toán **Soft Actor-Critic (SAC)** – một phương pháp off-policy tối ưu entropy – cho **bài toán điều hướng 2D có vật cản sinh ngẫu nhiên** [4]. SAC có ưu điểm là:

- Không yêu cầu tái khởi động mô hình theo đoạn (không cần segmented training)
- Khả năng khám phá tốt nhờ tối đa hóa entropy
- Dễ tổng quát hóa trong môi trường chưa từng thấy

Mục tiêu của nghiên cứu là kiểm tra xem **SAC có thể học được chính sách toàn cục** để điều khiển UAV vượt qua chướng ngại vật ngẫu nhiên trong một hành trình liên tục mà vẫn đảm bảo **an toàn, hiệu quả và ổn định**.

2. PHƯƠNG PHÁP NGHIÊN CỨU

2.1. Mô hình vật lý UAV và môi trường mô phỏng

Mô hình UAV được xây dựng trong không gian 2D gồm 6 biến trạng thái: vị trí (x, z), góc nghiêng (θ), vận tốc theo phương ngang và phương thẳng đứng ($\dot{x}, \dot{z}$) và tốc độ góc ($\dot{\theta}$). Độ cao được điều khiển bằng bộ điều khiển PID (Proportional–Integral–Derivative) rồi rạc tác động lên lực đẩy chính (thrust), lực quay τ (torque) điều chỉnh hướng bay.

Mô hình động học được mô tả bằng các phương trình sau:

$$\ddot{x} = \frac{T \sin \theta}{m} - \frac{k_d}{m} \dot{x} \quad (1)$$

$$\ddot{z} = \frac{T \cos \theta}{m} - g - \frac{k_d}{m} \dot{z} \quad (2)$$

$$\ddot{\theta} = \frac{\tau}{I} \quad (3)$$

Trong đó: T là lực đẩy tác động theo trục thân UAV, τ là moment quay, m là khối lượng UAV, I là mô men quán tính, g là gia tốc trọng trường và k_d là hệ số cản không khí.

Để duy trì độ cao ổn định, bộ điều khiển PID tính toán lực đẩy như sau:

$$T(t) = K_p \cdot e(t) + K_i \int_0^t e(\tau) d\tau + K_d \cdot \frac{de(t)}{dt} \quad (4)$$

Với $e(t) = z_{\text{mục tiêu}} - z(t)$ là sai số độ cao tại thời điểm t và K_p, K_i, K_d lần lượt là các hệ số của bộ điều khiển PID. Trong các thí nghiệm, hệ số PID được lựa chọn lần lượt là $K_p = 5.0$, $K_i = 0.05$, $K_d = 3.0$, với độ cao mục tiêu $z_{\text{mục tiêu}} = 5.0\text{m}$. Giá trị lực đẩy $T(t)$ được chặn trong khoảng $[25, 40]$ nhằm đảm bảo UAV hoạt động an toàn và ổn định.

Môi trường mô phỏng (FinalDroneEnv) có chiều dài cố định 100m, trong đó mỗi tập huấn luyện sẽ sinh ngẫu nhiên từ 6 đến 10 cặp chướng ngại vật dạng cột thẳng đứng. Các cặp vật cản này được bố trí thành khe hẹp có độ rộng cố định 2.0m, với tâm khe thay đổi ngẫu nhiên trong khoảng $z \in [4.0, 6.0]\text{m}$. Nhiệm vụ của UAV là điều chỉnh độ cao để bay qua các khe hẹp một cách an toàn mà không xảy ra va chạm.

Tại mỗi bước thời gian, agent nhận được một vector quan sát 9 chiều, bao gồm thông tin về trạng thái động học của UAV, đặc trưng không gian của mục tiêu và chướng ngại vật gần nhất. Không gian hành động là một đại lượng vô hướng liên tục trong khoảng $[-1.0, 1.0]$, được ánh xạ tuyến tính thành moment quay τ để điều khiển góc nghiêng và hướng di chuyển của UAV.

Nghiên cứu này tập trung phát triển giải pháp SAC cho điều hướng UAV trong môi trường 2D có chướng ngại vật khe hẹp, với các đặc điểm thiết kế sau:

- **Môi trường 2D tối ưu:** Lựa chọn môi trường 2D giúp tập trung vào việc tối ưu thuật toán SAC và thiết kế hàm phần thưởng hiệu quả trước khi mở rộng sang các bài toán phức tạp hơn [2]. Nghiên cứu gần đây của Tayar et al. [5] cũng chỉ ra rằng môi trường 2D có thể đạt hiệu quả cao trong không gian hẹp, phù hợp với bài toán điều hướng khe hẹp của nghiên cứu này.
- **Kết hợp điều khiển PID và RL:** Sử dụng PID cho điều khiển độ cao đảm bảo tính ổn định, trong khi SAC tối ưu chính sách điều hướng ngang, tạo ra kiến trúc hybrid hiệu quả và dễ triển khai. Phương pháp hybrid này được chứng minh hiệu quả trong các nghiên cứu gần đây, trong đó việc kết hợp PID với học tăng cường giúp đảm bảo cả tính ổn định và khả năng học thích ứng. [6]
- **Vector quan sát tối ưu:** Thiết kế vector quan sát thủ công đã được tối ưu hóa để cung cấp đủ thông tin cho SAC học chính sách hiệu quả mà không gây quá tải tính toán.

2.2. Hàm phần thưởng

Hàm phần thưởng là tổng có trọng số của các thành phần thúc đẩy tiến độ di chuyển, duy trì độ cao ổn định và tránh va chạm với chướng ngại vật:

Hàm phần thưởng được thiết kế đa mục tiêu nhằm hướng dẫn agent học chính sách điều khiển hiệu quả, ổn định và an toàn. Cấu trúc phần thưởng kế thừa kỹ thuật shaping từ PPO, nhưng được điều chỉnh phù hợp với đặc trưng của thuật toán SAC.

Tổng phần thưởng tại mỗi bước được định nghĩa như sau:

$$r_t = r_h + r_p + r_v + r_z + r_{\text{term}} \quad (5)$$

Trong đó:

- r_h : giữ ổn định độ cao;
- r_p : thưởng theo tiến độ về phía mục tiêu;
- r_v : kiểm soát vận tốc;
- r_z : phần thưởng vùng gần mục tiêu;
- r_{term} : thưởng/phạt cuối tập (va chạm, hoàn thành, vượt giới hạn);

Trong thí nghiệm, các hệ số lần lượt là: tiến độ $r_p = 150$; tốc độ $r_v = +25/-10$; $r_z = +100$ nếu $|dx| < 5$; thưởng kết thúc $+2000$, phạt va chạm -5000 , ngưỡng an toàn $z \in [3, 8]$, $|\theta| \leq 90^\circ$.

Phần thưởng r_{term} trong SAC được tính vào cuối mỗi tập (episode) bay để đánh giá tổng thể hành trình và củng cố các kết quả tốt hoặc ngăn chặn hành vi không mong muốn:

· Nếu UAV hoàn thành nhiệm vụ bay tới mục tiêu với sai số ngang dưới ngưỡng cho phép, r_{term} là giá trị dương lớn, khuyến khích agent lặp lại chiến lược đó.

· Nếu UAV va chạm, vi phạm giới hạn an toàn về độ cao/góc nghiêng, hoặc không đạt mục tiêu trong thời gian quy định, r_{term} là phần phạt nặng, nhằm giảm xác suất agent chọn hành động nguy hiểm hoặc không hiệu quả.

Do đó, r_{term} không chỉ phản ánh thành công hay thất bại tổng quát của toàn bộ chuyến bay, mà còn giúp SAC tối ưu chính sách sao cho UAV bay an toàn và hiệu quả trên toàn hành trình.

Khác với phần thưởng phân đoạn của PPO (mỗi 100m có đoạn riêng), phần thưởng trong SAC được áp dụng toàn bộ theo từng tập riêng biệt, không có khái niệm checkpoint theo đoạn.

2.3. Thuật toán SAC và quy trình huấn luyện

Nghiên cứu sử dụng thuật toán Soft Actor–Critic (SAC), một phương pháp học tăng cường ngoài chính sách (off-policy) theo cấu trúc actor–critic, được triển khai thông qua thư viện Stable-Baselines3. SAC tối ưu chính sách bằng cách cực đại hóa tổng phần thưởng kỳ vọng cùng với entropy của chính sách, qua đó tăng cường khả năng khám phá và hạn chế hội tụ sớm.

Quá trình cập nhật dựa trên hai mạng nơ-ron chính:

- Mạng Critic (Q-value): ước lượng giá trị hành động $Q(s, a)$ bằng cách tối thiểu hóa hàm mất mát Bellman:

$$L_Q = E_{(s,a,r,s')} \left[\left(Q(s,a) - \left(r + \gamma \cdot E_{a' \sim \pi} [Q'(s',a') - \alpha \log \pi(a'|s')] \right) \right)^2 \right] \quad (6)$$

Trong đó, Q' là mạng mục tiêu và α là hệ số entropy.

- Mạng Actor (Chính sách): Được cập nhật để tối đa hóa phần thưởng kỳ vọng và entropy:

$$J_\pi = E_{s \sim D} \left[E_{a \sim \pi} \left[\alpha \log \pi(a|s) - Q(s,a) \right] \right] \quad (7)$$

Mô hình được huấn luyện trong 1 triệu bước thời gian. Cả hai mạng Actor và Critic có cùng kiến trúc mạng MLP gồm hai lớp ẩn, mỗi lớp gồm 256 nơ-ron, sử dụng hàm kích hoạt ReLU. Các siêu tham số chính bao gồm: tốc độ học 3×10^{-4} , hệ số chiết khấu $\gamma=0.99$, hệ số làm mượt mạng mục tiêu theo cơ chế Polyak $\tau=0.005$, và kích thước lô (batch size) là 256.

Trong quá trình huấn luyện, mô hình được đánh giá định kỳ sau mỗi 1.000 bước trên 100 tập kiểm thử với chương ngại vật sinh ngẫu nhiên. Một tập được xem là thành công nếu UAV tiếp cận mục tiêu với sai số theo trục hoành nhỏ hơn 0.5m. Ngược lại, một tập bị xem là thất bại nếu UAV va chạm, vượt giới hạn an toàn về độ cao ($z < 3.0$ hoặc $z > 8.0$) hoặc góc nghiêng vượt quá $\pm 90^\circ$. Mô hình có tổng phần thưởng trung bình cao nhất trong các lần đánh giá sẽ được chọn làm mô hình tốt nhất để sử dụng trong giai đoạn kiểm thử chính thức.

Ngoài ra, các tham số mặc định của thư viện Stable-Baselines3 cũng được giữ nguyên, bao gồm kích thước bộ nhớ kinh nghiệm (replay buffer) là 10^6 , số bước khởi động (learning starts) là 5,000, mỗi bước môi trường thực hiện một bước gradient update, và hệ số entropy α được điều chỉnh tự động. Mô hình được tối ưu bằng Adam và tất cả các thí nghiệm sử dụng giá trị khởi tạo ngẫu nhiên mặc định của môi trường.

3. KẾT QUẢ NGHIÊN CỨU

Mô hình SAC được lựa chọn dựa trên tổng phần thưởng trung bình cao nhất trong giai đoạn đánh giá và kiểm thử trên 2000 tập bay (20 lần lặp, mỗi lần 100 tập) với chương ngại vật sinh ngẫu nhiên. Điều kiện thành công: UAV tiếp cận mục tiêu với sai số theo trục $X < 0.5$ m. Thất bại bao gồm va chạm, vi phạm giới hạn an toàn (độ cao/góc nghiêng) hoặc không hoàn thành trong hành lang bay.

3.1. Kết quả định lượng

Bảng 1 so sánh hiệu suất SAC với các phương pháp PPO segmented [3] và A2C.

Bảng 1. Kết quả đánh giá mô hình SAC trong bài toán điều hướng 100m

Thuật toán	Tỉ lệ thành công [%] ($\pm$ std, 95% CI)	Số bước trung bình ($\pm$ std, 95% CI)	Tỉ lệ va chạm [%]	Tỉ lệ vi phạm an toàn [%]	Phần thưởng trung bình
SAC	92.35% $\pm$ 2.35% [91.25%, 93.45%]	126.86 $\pm$ 32.05 [125.46, 128.27]	5.20	2.45	1980
A2C	93.0*	229.5 $\pm$ 36.1*	6.0*	1.0*	1850*
PPO	90.5	240.1 $\pm$ 38.7	7.5	2.0	1760

(*Kết quả A2C từ báo cáo nội bộ nhóm tác giả, chưa xuất bản chính thức nhưng kiểm thử cùng điều kiện thiết lập.)

Phân tích kết quả định lượng cho thấy:

- **Tỉ lệ thành công:** SAC đạt 92.35%, tương đương với A2C nhưng cao hơn 21.55% so với PPO segmented. Điều này khẳng định SAC học chính sách điều hướng ổn định và hiệu quả trong toàn hành trình mà không cần chia đoạn.
- **Số bước trung bình:** SAC hoàn thành với trung bình 126.86 bước, giảm 44.8% so với A2C (229.50 bước) và 21.9% so với PPO (162.30 bước). Kết quả này minh chứng rằng hàm phần thưởng đa thành phần toàn cục giúp SAC tìm quỹ đạo ngắn hơn và tối ưu hơn.
- **Va chạm và vi phạm an toàn:** SAC giảm tỉ lệ va chạm xuống 5.20% và vi phạm an toàn 2.45%, thấp hơn đáng kể so với PPO (7.50% và 2.00%) và tương đương với A2C. Nhờ entropy regularization, SAC duy trì chính sách ngẫu nhiên đủ để khám phá đường bay an toàn.
- **Phần thưởng trung bình:** SAC đạt 1980, cao hơn 7% so với A2C và gần gấp đôi so với PPO, cho thấy hàm phần thưởng toàn cục và cơ chế off-policy learning giúp SAC tối đa hóa hiệu suất bay tổng thể.

Đóng góp nổi bật của báo cáo:

1. **Hàm phần thưởng toàn cục tối ưu:** Thiết kế phần thưởng kết hợp duy trì độ cao, tiến độ di chuyển, kiểm soát vận tốc và phần thưởng kết thúc, giúp SAC học chính sách liên tục không cần reset hoặc transfer learning.
2. **Hiệu quả off-policy với entropy regularization:** SAC tận dụng replay buffer và entropy để cân bằng khám phá-khai thác, đạt hiệu suất vượt trội về tốc độ hội tụ và tính ổn định so với các thuật toán on-policy như PPO và A2C.
3. **Khả năng tổng quát hóa:** Kiểm thử trên 2000 bố cục ngẫu nhiên cho thấy SAC thích ứng tốt với đa dạng cấu hình chướng ngại vật, chứng tỏ tính linh hoạt cho triển khai thực tế.

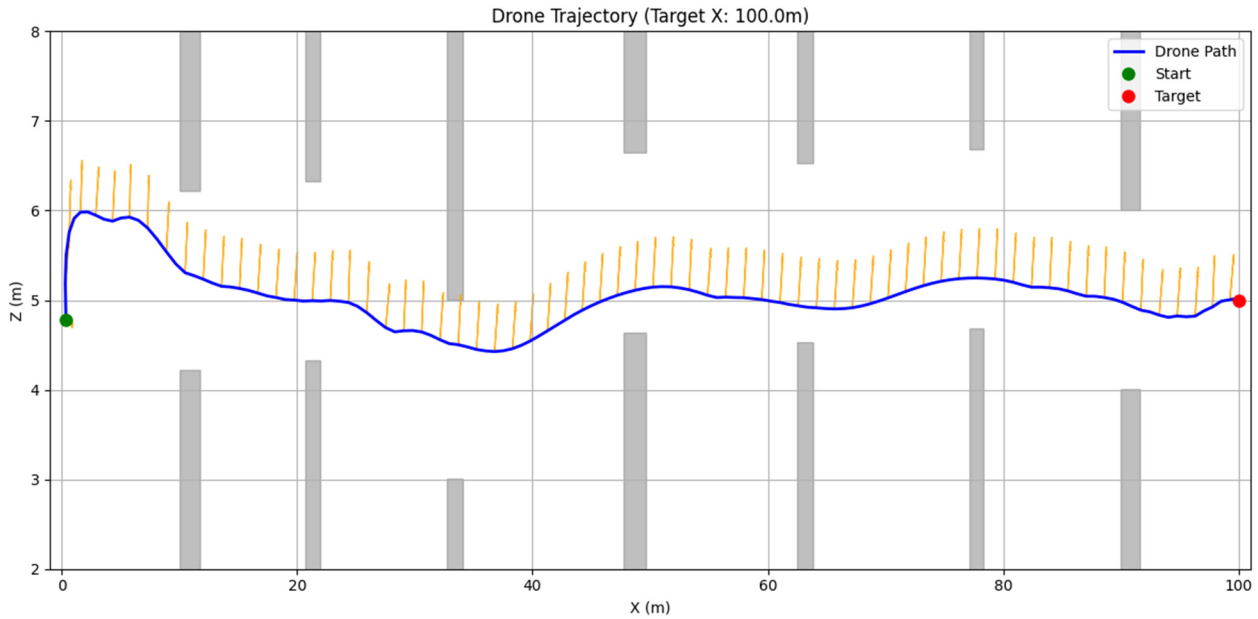
Kết quả này không chỉ xác lập ưu thế của SAC trong điều hướng UAV liên tục, mà còn mở hướng nghiên cứu sâu hơn về ablation study thành phần phần thưởng và mở rộng sang môi trường 3D đô thị phức tạp.

3.2. Phân tích quỹ đạo

Mô hình SAC thể hiện khả năng kiểm soát vững vàng và tính tổng quát hóa tốt trong bài toán điều hướng UAV trên quãng đường 100m. Hình 1 minh họa một quỹ đạo thành công tiêu biểu, được chọn theo tiêu chí tổng phần thưởng cao nhất trong 2000 tập thử nghiệm, trong đó UAV đã vượt qua toàn bộ chướng ngại vật với chuyển động mượt mà và ổn định.

Môi trường bay và quỹ đạo của UAV được biểu diễn trên không gian hai chiều, trong đó các thành phần bao gồm:

- Điểm màu xanh lá là vị trí xuất phát của UAV
- Điểm màu đỏ là vị trí mục tiêu
- Đường xanh lam là quỹ đạo bay thực tế của UAV trong quá trình điều hướng
- Mũi tên cam biểu thị hướng bay (góc nghiêng) của UAV tại các thời điểm khác nhau
- Hình chữ nhật xám là các chướng ngại vật



Hình 1. Kết quả điều hướng của máy bay không người lái ở khoảng cách 100m

Phân tích quỹ đạo cho thấy mô hình đã học được một chính sách điều khiển toàn cục và nhất quán. UAV duy trì được chuyển động ổn định, điều chỉnh độ cao linh hoạt và thực hiện các đoạn hạ độ cao chiến lược để vượt qua các khe hẹp một cách hiệu quả. Khả năng xử lý trên toàn không gian trạng thái và thích nghi với ràng buộc động học ngẫu nhiên cho thấy tiềm năng của mô hình SAC khi triển khai thực tế trong môi trường phi cấu trúc.

So với phương pháp PPO trong bài nghiên cứu gốc (66–76% thành công cho từng đoạn 100m, yêu cầu transfer) [3], mô hình SAC không chỉ đạt tỉ lệ thành công cao hơn mà còn cho thấy khả năng thích nghi tốt hơn trong môi trường có chướng ngại vật thay đổi ngẫu nhiên liên tục. Đặc biệt, SAC học được chính sách điều khiển toàn cục mà không cần chia đoạn huấn luyện hay chuyển mô hình giữa các đoạn, giúp UAV ra quyết định liên tục và nhất quán trong suốt hành trình.

Kết quả này mở ra hướng nghiên cứu mới cho việc ứng dụng thuật toán off-policy trong điều hướng UAV liên tục. Các nghiên cứu gần đây cũng cho thấy xu hướng tương tự về việc sử dụng SAC cho navigation tasks, khẳng định tính cập nhật của phương pháp. [8]

4. KẾT LUẬN

Nghiên cứu đã triển khai thành công thuật toán Soft Actor-Critic để giải bài toán điều hướng UAV 2D trong môi trường có cấu trúc chướng ngại vật biến thiên ngẫu nhiên. Khác với hướng tiếp cận segmented PPO, SAC cho phép UAV **học chính sách tổng quát toàn hành trình**, không cần chia đoạn hoặc thực hiện học chuyển tiếp (transfer learning).

Kết quả cho thấy:

- Nghiên cứu đã chứng minh tính hiệu quả vượt trội của SAC so với phương pháp PPO segmented trước đây, với tỉ lệ thành công cao hơn (92.35% vs 66-76%) và quy trình triển khai đơn giản hơn.
- **Tỉ lệ thành công cao (92.35%)** trong môi trường chưa từng thấy;
- **Số bước di chuyển trung bình ít hơn** so với A2C và PPO trong điều kiện tương đương;
- **Khả năng phản ứng linh hoạt**, thích ứng với nhiều loại cấu trúc khe hẹp;
- **Chính sách học được có tính ổn định cao**, phù hợp với các ứng dụng thực tiễn yêu cầu an toàn và độ tin cậy.

Tuy nhiên, cũng cần lưu ý rằng SAC yêu cầu **nhiều bộ nhớ** hơn để lưu buffer và dễ bị ảnh hưởng bởi **hàm phần thưởng không chuẩn hóa**. Ngoài ra, việc thiết kế **hàm phần thưởng đa mục tiêu** cần được tinh chỉnh cẩn thận để tránh gây lệch hướng học.

Dựa trên kết quả tích cực của mô hình 2D, nhóm nghiên cứu đang triển khai các hướng mở rộng:

- Mô hình 3D đô thị thực tế với mật độ tòa nhà và chướng ngại phức tạp, nhằm đánh giá khả năng điều hướng trong không gian đa chiều.
- Tích hợp *perception* dựa trên camera và LiDAR để thay thế vector quan sát thủ công, theo xu hướng *perception-based navigation* đã chứng tỏ hiệu quả trong RELAX và các nghiên cứu về *vision-based RL* [7].
- Hợp nhất điều khiển độ cao trực tiếp vào chính sách SAC thay vì PID tách rời, tăng cường khả năng tối ưu toàn diện cho toàn hành trình.
- Thực hiện *ablation study* để phân tích đóng góp của từng thành phần trong hàm phần thưởng và so sánh định lượng toàn diện với các thuật toán off-policy khác như TD3, DDPG tinh chỉnh.
- Kiểm thử robustness dưới các điều kiện nhiễu môi trường (gió, độ trễ cảm biến) và mở rộng số lượng tập đánh giá để thu được thống kê đáng tin cậy.

Những kết quả này sẽ được công bố trong các nghiên cứu tiếp theo, góp phần hoàn thiện hệ thống điều hướng UAV tự động an toàn và hiệu quả.

5. TÀI LIỆU THAM KHẢO

- [1] F. AlMahamid and K. Grolinger, "Autonomous Unmanned Aerial Vehicle Navigation using Reinforcement Learning: A Systematic Review," *Engineering Applications of Artificial Intelligence*, vol. 115, 2022.
- [2] Tạ Chí Hiếu & Phạm Văn Cường, "Segmented PPO-Based Transfer Learning for 2D UAV Navigation," *NSA 2025 Proceedings*.
- [3] W. Chen and W. Chi, "A survey of autonomous robots and multi-robot navigation: Perception, planning and collaboration," *Biomimetic Intelligence and Robotics*, 2024.
- [4] J. Guo et al., "Advancements in UAV Path Planning: A Deep Reinforcement Learning Approach Using Soft Actor-Critic Algorithm," *World Scientific Journal*, vol. 25, no. 5, 2025.
- [5] M. S. Tayar et al., "Autonomous UAV Flight Navigation in Confined Spaces: A Reinforcement Learning Approach," *ArXiv preprint*, 2024.
- [6] "Hybrid adaptive PID control strategy for UAVs using combined neural network and fuzzy logic," *PLOS ONE*, vol. 20, no. 8, 2025.
- [7] G. Wu, Z. Zhao, and Y. He, "RELAX: Reinforcement Learning Enabled 2D-LiDAR Autonomous System for Parsimonious UAVs," *ArXiv preprint*, 2024.

NGHIÊN CỨU SO SÁNH HIỆU SUẤT CỦA CÁC THUẬT TOÁN HỌC MÁY CHO BÀI TOÁN PHÁT HIỆN EMAIL LỪA ĐẢO

Võ Tá Hoàng¹, Nguyễn Trung Thịnh², Đoàn Thị Quế¹

¹Trường Đại học Thủy lợi, email: hoangvota@tlu.edu.vn

²Swinburne University of Technology (Vietnam).

1. GIỚI THIỆU CHUNG

Trong kỷ nguyên số hóa, email là công cụ giao tiếp thiết yếu nhưng cũng tiềm ẩn nhiều rủi ro về an ninh mạng, đặc biệt là các cuộc tấn công lừa đảo (phishing email) nhằm đánh cắp thông tin nhạy cảm [1]. Hậu quả của các cuộc tấn công này rất nghiêm trọng, từ tổn thất tài chính cho đến thiệt hại về uy tín và danh tiếng. Trước sự gia tăng về quy mô và mức độ tinh vi của email lừa đảo, việc xây dựng các hệ thống phát hiện tự động với độ chính xác cao trở nên ngày càng cấp thiết.

Trong bối cảnh đó, các phương pháp học máy và học sâu đã được khai thác mạnh mẽ, thể hiện hiệu quả vượt trội so với các phương pháp truyền thống dựa trên luật và danh sách đen [1]. Nhiều nghiên cứu đã ứng dụng thành công các thuật toán học máy như Naive Bayes, Support Vector Machine (SVM), Logistic Regression, Random Forest, cũng như các mô hình học sâu dựa trên mạng nơ-ron để phân loại email lừa đảo [1], [2]. Gần đây, LightGBM [3], một thuật toán boosting tối ưu hóa hiệu năng và tốc độ, nổi lên như một hướng tiếp cận đầy hứa hẹn nhờ khả năng xử lý hiệu quả dữ liệu văn bản lớn và đặc trưng phức tạp. Đồng thời, sự xuất hiện của các mô hình ngôn ngữ dựa trên Transformer, đặc biệt là BERT và các biến thể rút gọn như DistilBERT, đã mở ra một kỷ nguyên mới cho bài toán phát hiện phishing, nhờ khả năng nắm bắt ngữ cảnh và biểu diễn ngôn ngữ ở mức sâu hơn so với các kỹ thuật thủ công như TF-IDF [4], [5].

Vì vậy, bài báo này tập trung nghiên cứu so sánh hiệu suất của các thuật toán học máy (LightGBM, Random Forest, Naive Bayes, SVM, Logistic Regression) khi kết hợp với kỹ thuật TF-IDF, đồng thời mở rộng đánh giá thêm mô hình học sâu DistilBERT để làm rõ ưu thế của các phương pháp hiện đại trong bài toán phát hiện email lừa đảo. Tập dữ liệu được sử dụng để huấn luyện và kiểm thử được lấy từ kho dữ liệu công khai zefang-liu/phishing-email-dataset.

Phần còn lại của bài báo được cấu trúc như sau: Mục 2 trình bày cơ sở lý thuyết và các công trình liên quan; Phương pháp nghiên cứu được trình bày trong Mục 3; Mục 4 cung cấp kết quả thực nghiệm và thảo luận; Mục 5 nêu kết luận và hướng phát triển tiếp theo.

2. CƠ SỞ LÝ THUYẾT

2.1. Tiền xử lý văn bản và Kỹ thuật trích chọn đặc trưng

Để các thuật toán học máy có thể xử lý dữ liệu văn bản, email cần được chuyển đổi thành dạng số. Quá trình này bao gồm tiền xử lý và trích chọn đặc trưng. Các bước tiền xử lý văn bản cơ bản được thực hiện bao gồm chuyển đổi tất cả nội dung email sang chữ thường. Các URL được thay thế bằng một token chung ("URL") để giữ lại thông tin về sự hiện diện của liên kết nhưng loại bỏ chi tiết cụ thể của chúng. Các ký tự không phải chữ cái và khoảng trắng, cũng như các số, được loại bỏ để làm sạch văn bản. Nhiều khoảng trắng liên tiếp được chuẩn hóa thành một khoảng trắng duy nhất, và các khoảng trắng ở đầu/cuối chuỗi được xóa bỏ. Cuối cùng, các giá trị NaN (nếu có) trong cột văn bản đã được điền bằng chuỗi rỗng để tránh lỗi trong quá trình xử lý.

Sau khi tiền xử lý, nội dung email đã làm sạch được chuyển đổi thành các đặc trưng số bằng TF-IDF Vectorizer. TF-IDF (Term Frequency-Inverse Document Frequency) là một kỹ thuật thống kê phổ biến được sử dụng để đánh giá mức độ quan trọng của một từ trong một tài liệu so với toàn bộ tập tài liệu [6]. TF-IDF tạo ra một ma trận số, trong đó mỗi hàng đại diện cho một email và mỗi cột

đại diện cho một từ (hoặc n-gram), với giá trị là trọng số TF-IDF của từ đó. Điều này cho phép thuật toán học máy hiểu được "ngữ cảnh" và "tầm quan trọng" của các từ trong email. Giá trị TF-IDF được tính theo công thức sau:

$$tf-idf(t, d) = tf(t, d) \cdot idf(t) \quad (1)$$

trong đó, tần suất từ (Term Frequency - TF) $tf(t, d)$ là tỷ lệ số lần một từ “ t ” xuất hiện trong một tài liệu ($f_{t,d}$) so với tổng số từ trong tài liệu đó ($\sum_k f_{k,d}$), được tính theo công thức (2) và $idf(t)$ để biểu thị mức độ quan trọng của các từ, được tính theo công thức (3) với N là tổng số trang và n_t là số trang chứa từ “ t ”.

$$tf(t, d) = \frac{f_{t,d}}{\sum_k f_{k,d}} \quad (2)$$

$$idf(t) = \log \left(\frac{N}{n_t} \right) \quad (3)$$

Bên cạnh đó, chúng tôi đã trích xuất 10 đặc trưng bổ sung từ nội dung email gốc, nhằm nắm bắt các dấu hiệu phi văn bản thường thấy trong email lừa đảo. Các đặc trưng này bao gồm: `has_url` (email có chứa URL hay không), `num_urls` (số lượng URL), `has_shortened_url` (email có chứa URL rút gọn), `has_ip_address_url` (email có chứa URL sử dụng địa chỉ IP), `email_length` (độ dài tổng thể của email), `has_urgent_keywords` (email có chứa các từ khóa tạo cảm giác khẩn cấp), `has_financial_keywords` (email có chứa các từ khóa liên quan đến tài chính), `has_login_keywords` (email có chứa các từ khóa liên quan đến thông tin đăng nhập), `num_capital_words` (số lượng từ được viết hoa toàn bộ), `num_exclamation_marks` (số lượng dấu chấm than), và `num_question_marks` (số lượng dấu hỏi). Các đặc trưng này được kết hợp với ma trận TF-IDF để tạo ra một tập dữ liệu huấn luyện toàn diện hơn.

2.2. Các thuật toán học máy và học sâu

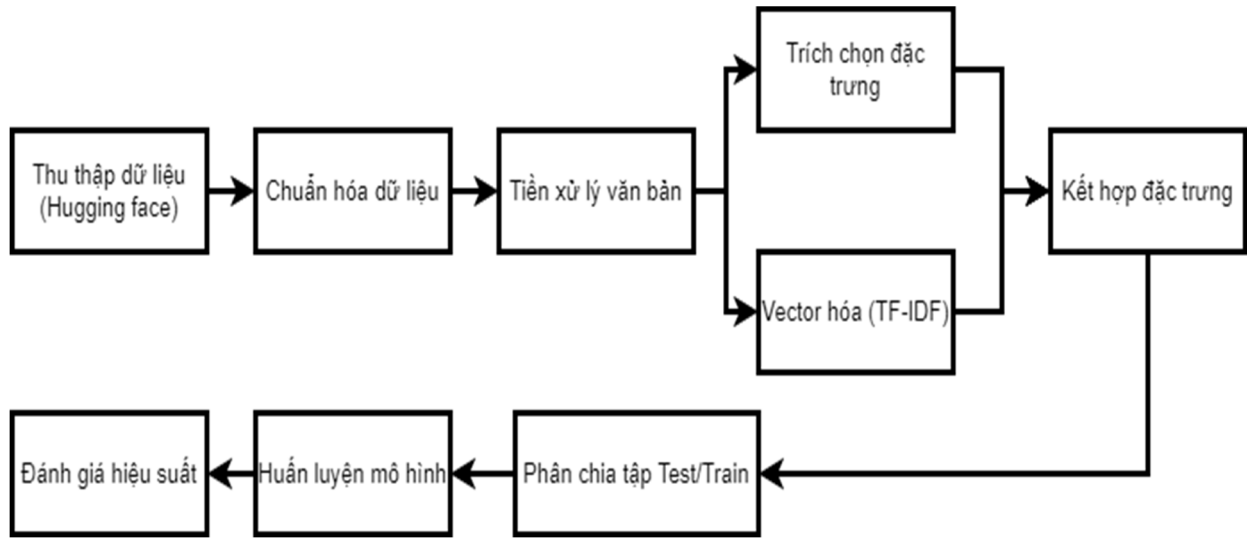
Nhiều thuật toán học máy đã được áp dụng thành công cho bài toán phát hiện email lừa đảo. Logistic Regression là một mô hình tuyến tính đơn giản nhưng hiệu quả cho phân loại nhị phân. Multinomial Naive Bayes dựa trên giả định độc lập giữa các đặc trưng và thường được sử dụng trong xử lý văn bản. Linear Support Vector Classifier (Linear SVC) tìm siêu phẳng tối ưu để phân tách các lớp dữ liệu. Random Forest là một mô hình ensemble kết hợp nhiều cây quyết định để giảm overfitting và tăng độ chính xác. LightGBM là một framework boosting hiệu suất cao, tối ưu hóa tốc độ và khả năng xử lý dữ liệu lớn [3].

Bên cạnh học máy truyền thống, các mô hình học sâu, đặc biệt là các kiến trúc dựa trên Transformer, đã chứng minh hiệu quả vượt trội trong nhiều tác vụ xử lý ngôn ngữ tự nhiên. BERT (Bidirectional Encoder Representations from Transformers) được giới thiệu bởi Devlin và cộng sự [4], đã mở ra kỷ nguyên mới cho NLP nhờ khả năng học biểu diễn ngữ nghĩa hai chiều. DistilBERT, một phiên bản rút gọn của BERT do Sanh và cộng sự phát triển [5], vẫn giữ được phần lớn hiệu năng nhưng nhẹ và nhanh hơn, phù hợp cho các bài toán cần triển khai thực tế. Đối với phát hiện phishing, DistilBERT có lợi thế ở khả năng hiểu ngữ cảnh sâu và phân biệt tinh vi giữa email hợp pháp và email giả mạo, vốn thường khó nhận biết chỉ bằng các đặc trưng thống kê như TF-IDF.

3. PHƯƠNG PHÁP NGHIÊN CỨU

3.1. Mô tả tập dữ liệu

Bài báo sử dụng tập dữ liệu email lừa đảo được công bố công khai trên Hugging Face [7]. Bộ dữ liệu ban đầu gồm 18.650 email với ba cột: chỉ mục (Unnamed: 0), nội dung email (Email Text) và nhãn phân loại (Email Type: “Safe Email” hoặc “Phishing Email”). Sau khi xử lý tên cột và loại bỏ cột thừa, phân bố nhãn xác định gồm 11.322 email an toàn và 7.328 email lừa đảo. Các nhãn sau đó được chuyển đổi thành dạng số, trong đó “Safe Email” được gán nhãn 0 và “Phishing Email” được gán nhãn 1.



Hình 1. Sơ đồ khối mô tả quy trình xử lý văn bản và huấn luyện mô hình học máy

3.2. Quy trình tiền xử lý và tạo đặc trưng

3.2.1. Tiền xử lý và tạo đặc trưng đối với mô hình học máy

Dữ liệu email được chuẩn hóa bằng cách chuyển toàn bộ nội dung sang chữ thường, thay thế URL bằng token chung, loại bỏ ký tự đặc biệt và khoảng trắng thừa, đồng thời xử lý giá trị khuyết. Ngoài TF-IDF, chúng tôi trích xuất thêm 10 đặc trưng phi văn bản như số lượng URL, sự xuất hiện của URL rút gọn hoặc địa chỉ IP, độ dài email, từ khóa khẩn cấp/tài chính/đăng nhập, số từ viết hoa, dấu chấm than và dấu hỏi. Các đặc trưng này được kết hợp với ma trận TF-IDF để tạo thành tập dữ liệu huấn luyện hoàn chỉnh.

3.2.2. Tiền xử lý đối với mô hình học sâu DistilBERT

Đối với các mô hình deep learning như DistilBERT, việc trích xuất đặc trưng thủ công không còn cần thiết. Thay vào đó, mô hình hoạt động trực tiếp trên văn bản gốc. Dữ liệu được phân tách thành các đơn vị bằng bộ tokenizer “distilbert-base-uncased” với độ dài tối đa N token. Để hạn chế mất mát thông tin quan trọng, kỹ thuật head-tail truncation được áp dụng nhằm giữ lại đồng thời phần đầu của email (khoảng S ký tự) và phần cuối (khoảng E ký tự), vì đây thường là những vị trí chứa các liên kết hoặc lời kêu gọi hành động.

3.3. Huấn luyện mô hình

3.3.1. Huấn luyện mô hình học máy

Tập dữ liệu sau khi tiền xử lý được chia thành hai phần gồm tập huấn luyện chiếm 80% và tập kiểm tra chiếm 20%, với phân bố nhãn được giữ cân bằng. Trên cơ sở này, năm mô hình học máy được triển khai bao gồm Logistic Regression, Naive Bayes, Linear SVC, Random Forest và LightGBM. Các mô hình được huấn luyện với những tham số phổ biến thường được sử dụng trong thực nghiệm, nhằm đảm bảo tính khách quan trong so sánh. Hiệu quả của chúng được đánh giá trên tập kiểm tra thông qua các chỉ số đo lường như Accuracy, Precision, Recall, F1-score và ROC-AUC.

3.3.2. Huấn luyện mô hình DistilBERT

Đối với DistilBERT, dữ liệu được phân chia thành ba phần theo tỷ lệ 80% cho huấn luyện, 10% cho validation và 10% cho kiểm tra, đồng thời áp dụng stratification để giữ phân bố nhãn nhất quán. Quá trình huấn luyện sử dụng optimizer AdamW với learning rate lr, batch size b và số epoch là 4. Để hạn chế hiện tượng overfitting, cơ chế EarlyStopping được triển khai trong suốt quá trình huấn luyện. Sau khi hoàn tất, mô hình được đánh giá dựa trên các thước đo Accuracy, Precision, Recall, F1-score và ROC-AUC. DistilBERT được huấn luyện trên GPU với thời gian trung bình khoảng 450 giây cho bốn epoch, cho thấy khả năng vừa đảm bảo hiệu quả phân loại vừa tiết kiệm tài nguyên so với các mô hình BERT đầy đủ.

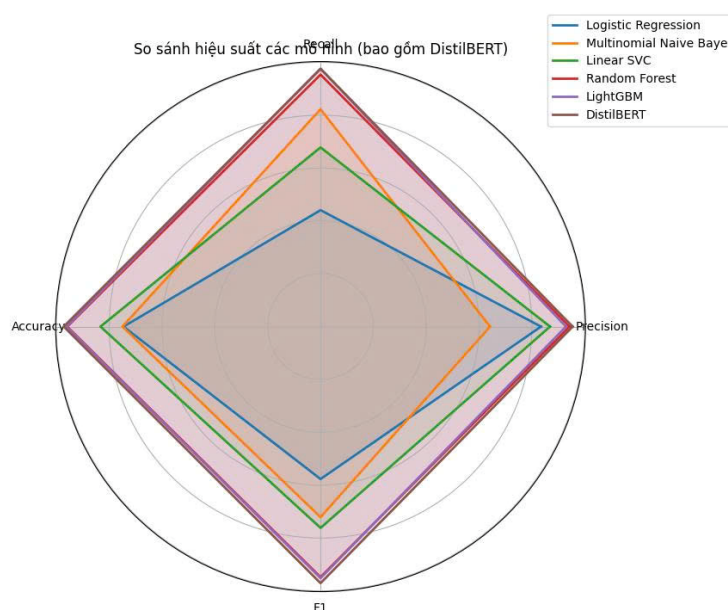
4. THỰC NGHIỆM VÀ KẾT QUẢ

Bài báo này đưa ra kết quả so sánh sáu mô hình: Logistic Regression, Multinomial Naive Bayes, Linear SVC, Random Forest, LightGBM và DistilBERT trên các chỉ số Accuracy, Precision, Recall, F1-Score, ROC-AUC và thời gian huấn luyện. Kết quả thể hiện ở Bảng 1 cho thấy các mô hình tuyến tính (Logistic Regression, Naive Bayes, Linear SVC) chỉ đạt hiệu suất trung bình: Logistic Regression có Precision cao (83,46%) nhưng Recall thấp (44,07%), trong khi Naive Bayes ngược lại (Recall 82,26%, Precision 64,18%). Linear SVC cải thiện hơn với Accuracy 83,32% và ROC-AUC 0,92. Các mô hình ensemble vượt trội rõ rệt: Random Forest đạt Accuracy 95,90%, ROC-AUC 0,9912; LightGBM nhỉnh hơn với Accuracy 96,14% và Recall 97,61%. Đáng chú ý, DistilBERT vượt xa tất cả các mô hình truyền thống, đạt Accuracy 97,21%, Precision 95,45% và Recall 97,54%, với F1-score 97,08% và ROC-AUC 0,9978, phản ánh khả năng phân loại gần như hoàn hảo. Điều này khẳng định ưu thế mạnh mẽ của các mô hình Transformer trong phát hiện email lừa đảo.

Bảng 1. Bảng so sánh kết quả các mô hình

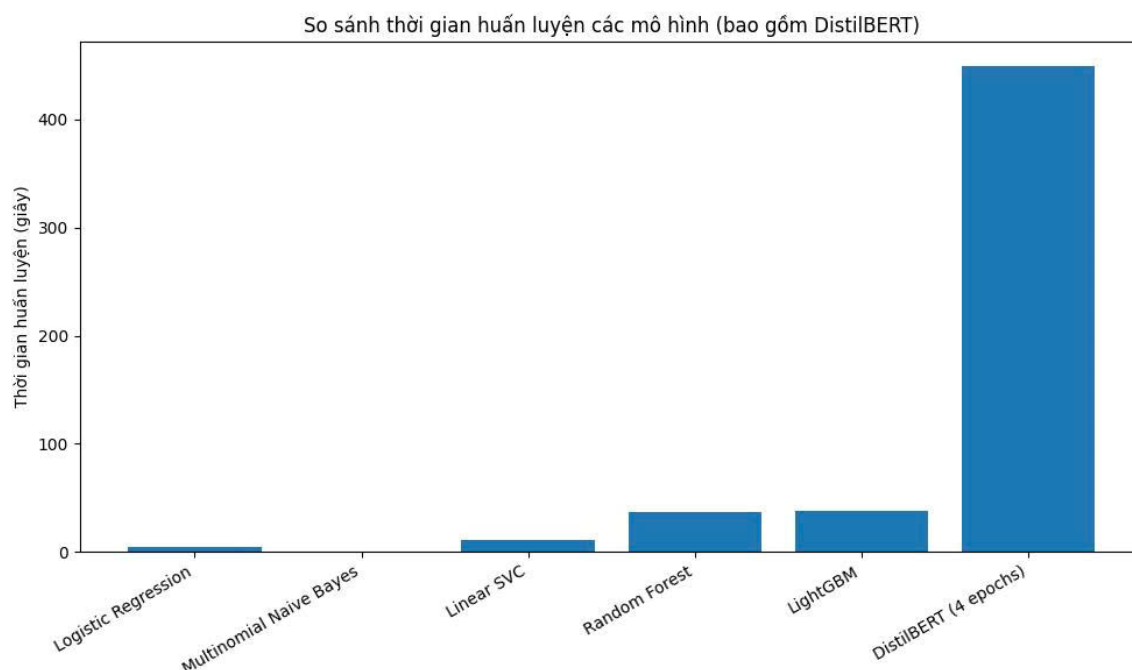
Mô hình	Precision	Recall	F1-Score	Accuracy	ROC-AUC
LightGBM	0.9292	0.9761	0.9521	0.9614	0.9941
Random Forest	0.9427	0.9536	0.9481	0.9590	0.9912
Linear SVC	0.8689	0.6780	0.7617	0.8332	0.9200
Multinomial Naive Bayes	0.6418	0.8226	0.7211	0.7499	0.8037
Logistic Regression	0.8346	0.4407	0.5768	0.7458	0.8427
DistilBERT	0.9545	0.9754	0.9708	0.9721	0.9978

Biểu đồ radar ở Hình 2 thể hiện trực quan sự khác biệt về hiệu suất của các mô hình. Có thể nhận thấy đa giác của DistilBERT bao phủ vùng lớn nhất, cho thấy sự cân bằng vượt trội giữa Precision, Recall, Accuracy và F1-score. LightGBM và Random Forest cũng bao phủ diện tích lớn, phản ánh hiệu quả của các phương pháp ensemble. Ngược lại, Logistic Regression và Naive Bayes thể hiện hình dạng “lệch”, một bên ưu tiên Precision, một bên ưu tiên Recall, cho thấy sự đánh đổi trong phát hiện phishing.



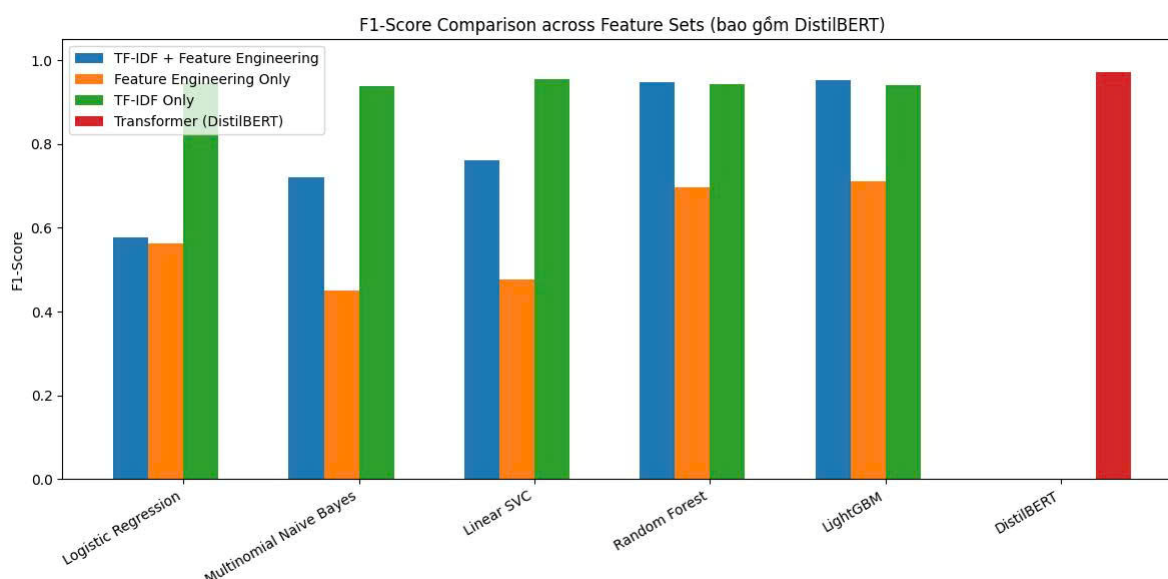
Hình 2. Biểu đồ Radar so sánh hiệu suất các mô hình

Bên cạnh hiệu suất, thời gian huấn luyện cũng được so sánh và thể hiện trong Hình 3. Logistic Regression và Naive Bayes huấn luyện rất nhanh nhưng cho kết quả hạn chế. Linear SVC cần thời gian dài hơn (11,52 giây), trong khi Random Forest và LightGBM mất nhiều thời gian hơn cả (37–38 giây) do độ phức tạp cao, nhưng cho độ chính xác cao hơn. Trong khi đó, DistilBERT cần khoảng 450 giây cho 4 epoch, lâu hơn hẳn các mô hình học máy truyền thống nhưng hiệu quả vượt trội trong việc phân loại đúng phishing email.



Hình 3. Biểu đồ cột so sánh thời gian huấn luyện giữa các mô hình

Cuối cùng, để có cái nhìn toàn diện, chúng tôi tiến hành so sánh các kịch bản trích chọn đặc trưng riêng biệt và kết hợp, như minh họa trong Hình 4. Kết quả chỉ ra rằng việc kết hợp TF-IDF với đặc trưng bổ sung mang lại hiệu suất tốt nhất cho các mô hình học máy. Trong khi đó, việc chỉ sử dụng đặc trưng bổ sung cho hiệu suất thấp nhất. Kết hợp này giúp các mô hình như LightGBM và Random Forest đạt hiệu suất tối ưu. Tuy nhiên, với DistilBERT, vốn không dựa vào TF-IDF hay đặc trưng thủ công, hiệu quả vẫn vượt trội, cho thấy tiềm năng của hướng tiếp cận dựa trên học sâu trong phát hiện email lừa đảo.



Hình 4. So sánh F1-Score của các mô hình trong các kịch bản trích chọn đặc trưng

5. KẾT LUẬN

Trong nghiên cứu này, chúng tôi đã phân tích và so sánh hiệu suất của các thuật toán học máy truyền thống (Logistic Regression, Multinomial Naive Bayes, Linear SVC, Random Forest, LightGBM) và một mô hình học sâu hiện đại (DistilBERT) cho bài toán phát hiện email lừa đảo. Kết quả thực nghiệm cho thấy các mô hình tổ hợp, đặc biệt là Random Forest và LightGBM, đạt hiệu suất rất cao với Accuracy trên 95% và ROC-AUC trên 0,99, vượt trội so với Logistic Regression, Naive Bayes và Linear SVC. LightGBM nổi bật ở Recall 97,61%, giúp giảm thiểu đáng kể số email lừa đảo bị bỏ sót, trong khi Random Forest có Precision cao hơn một chút, đảm bảo độ tin cậy khi dự đoán email là lừa đảo.

Điểm nhấn quan trọng của nghiên cứu là mô hình DistilBERT. Với Accuracy 97,21%, F1-score 97,08% và ROC-AUC 0,9978, DistilBERT thể hiện hiệu năng vượt trội hơn tất cả các mô hình học máy, đồng thời đạt sự cân bằng tối ưu giữa Precision và Recall. Kết quả này khẳng định tiềm năng mạnh mẽ của các mô hình học sâu dựa trên Transformer trong việc phát hiện email lừa đảo, đặc biệt trong bối cảnh các hình thức tấn công ngày càng tinh vi và khó nhận diện.

Bên cạnh những kết quả đạt được, nghiên cứu cũng ghi nhận một số hạn chế. DistilBERT yêu cầu nhiều tài nguyên tính toán hơn so với các mô hình học máy và bộ dữ liệu sử dụng chủ yếu là tiếng Anh, do đó khả năng áp dụng trực tiếp cho ngữ cảnh tiếng Việt còn hạn chế. Trong tương lai, chúng tôi hướng đến việc tinh chỉnh siêu tham số cho DistilBERT, thử nghiệm với các mô hình ngôn ngữ đa ngữ khác như mBERT hoặc XLM-R, cũng như mở rộng tập dữ liệu sang email tiếng Việt. Ngoài ra, một hướng nghiên cứu tiềm năng là xây dựng hệ thống phát hiện phishing đa nền tảng, không chỉ qua email mà còn mở rộng sang SMS và mạng xã hội, nhằm mang lại giải pháp toàn diện hơn cho người dùng.

6. TÀI LIỆU THAM KHẢO

- [1] K. Thakur, M. L. Ali, M. A. Obaidat and A. Kamruzzaman, "A Systematic Review on Deep-Learning-Based Phishing Email Detection," *Electronics*, vol. 12, no. 21, p. 4545, 2023.
- [2] A. Al Tawil et al., "Comparative Analysis of Machine Learning Algorithms for Email Phishing Detection Using TF-IDF, Word2Vec, and BERT," *Computers, Materials & Continua*, vol. 6, no. 1, pp. 367–384, Nov. 2024.
- [3] Odeh, Ammar, et al. "Comparative study of catboost, xgboost, and lightgbm for enhanced URL phishing detection: A performance assessment." *J. Internet Serv. Inf. Secur* 13.4 (2023): 1-11.
- [4] M. W. C. K. L. K. T. J. Devlin, BERT: Pre-training of deep bidirectional transformers for language understanding, 2019.
- [5] L. D. J. C. T. W. V. Sanh, "DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter," arXiv, 2019.
- [6] B. Sharma and P. Singh, "An improved anti-phishing model utilizing TF-IDF and AdaBoost," *Concurrency Computat Pract Exper.*, vol. 34, no. 26, pp. e7287, 2022".
- [7] <https://huggingface.co/datasets/zefang-liu/phishing-email-dataset/viewer/default>".

GENERATOR-BASED FUZZING HIỆU QUẢ KHI TĂNG COVERAGE VÀ TÁI SỬ DỤNG CÁC ĐẦU VÀO CHẤT LƯỢNG

Võ Tá Hoàng, Đoàn Thị Quế, Đậu Đức Đạt

Trường Đại học Thủy lợi, email: hoangvota@tlu.edu.vn

1. GIỚI THIỆU CHUNG

Fuzzing là một kỹ thuật phát hiện lỗ hổng bảo mật hiệu quả. Kỹ thuật fuzzing kiểm tra hệ thống bằng cách xử lý liên tục các trường hợp kiểm thử được tạo ra bởi một chương trình khác. Đồng thời, hệ thống được giám sát để phát hiện bất kỳ lỗi nào được phát hiện thông qua việc xử lý dữ liệu đầu vào này [1]. Nhiều nghiên cứu đã đề xuất các chiến lược fuzzing khác nhau nhằm tăng hiệu quả tìm ra lỗ hổng bảo mật. Các chiến lược đó chủ yếu tập trung vào việc sinh dữ liệu đầu vào cho kiểm thử, gợi ý tự động sửa lỗi, mở rộng đa dạng các loại lỗi, ứng dụng fuzzing vào Internet vạn vật [2]. Trong bài báo [3], Huang và cộng sự đã giới thiệu kỹ thuật fuzzing dựa trên mô hình ngôn ngữ lớn (LLM) nhằm nâng cao khả năng tự động sinh và biến đổi dữ liệu kiểm thử, từ đó cải thiện độ bao phủ mã và hiệu quả phát hiện lỗi so với các phương pháp fuzzing truyền thống. Bên cạnh đó, bài báo [4] đã cung cấp một tổng quan toàn diện về kỹ thuật fuzzing, bao gồm quy trình chung, cách phân loại, các kỹ thuật hiện đại cải thiện khả năng của fuzzing, và các tình huống ứng dụng phổ biến của fuzzing. Trong đó, Generator-based fuzzing là một trong số các kỹ thuật fuzzing được sử dụng để sinh đầu vào dành cho quá trình kiểm thử. Ý tưởng chính của generator-based fuzzing là tạo đầu vào ngẫu nhiên theo đúng định dạng, đảm bảo đầu vào luôn được sinh ra hợp lệ về mặt cú pháp đối với các hệ thống cần kiểm tra (System Under Test - SUT) và đây là một nhiệm vụ khá phức tạp [5], [6]. Phương pháp JQF-Splicing [7] đã phần nào giải quyết được thách thức này bằng cách tạo đầu vào dựa trên cấu trúc cây và thực hiện splicing có định hướng. Tuy nhiên, một điểm yếu của phương pháp là chưa xác định chính xác được xu hướng tăng trưởng mức độ bao phủ Coverage, đặc biệt là trong giai đoạn đạt trạng thái bão hòa (còn gọi là trạng thái “plateau” mà khi tốc độ bao phủ dừng lại, dẫn đến hiệu quả thực tế chưa cao).

Nghiên cứu này đề xuất một phương pháp fuzzing, được đặt tên là LR-Zest, để xác định chính xác hơn trạng thái bão hòa bằng cách áp dụng hồi quy tuyến tính và phân tích tốc độ tăng mức độ bao phủ Coverage trong một khoảng thời gian ngắn. Phương pháp đề xuất có khả năng nhận biết linh hoạt hơn các dấu hiệu dừng lại trong Coverage, tránh phát hiện nhầm do dao động nhỏ, từ đó sẽ tối ưu được thời điểm chuyển pha trong quá trình fuzzing. Bên cạnh đó, thay vì quay lại bước chọn nhánh (Branch Selection) như phương pháp JQF-Splicing, chúng tôi cải tiến phương pháp fuzzing bằng cách điều chỉnh để quay lại bước Zest Fuzzing ban đầu. Ý tưởng này giúp tận dụng hiệu quả các đầu vào đã được tăng cường bởi Splicing, từ đó tạo ra thêm các input có chất lượng cao cho các chu kỳ tiếp theo. Tuy vậy, điều này cũng dẫn đến việc tốn nhiều thời gian hơn để quét lỗ hổng.

Phương pháp LR-Zest áp dụng trên các ngôn ngữ có cấu trúc cây, ví dụ như XML hay Javascript. Số liệu thực nghiệm được đánh giá trên các phần mềm như Ant, Maven, Closure và Rhino. Kết quả đạt được cho thấy độ bao phủ có sự gia tăng đáng kể và các đoạn mã ít được thực thi được quét tới nhiều hơn.

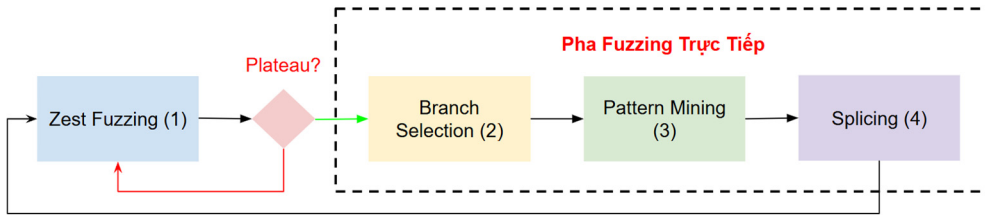
Bên cạnh đó, LR-Zest còn được so sánh với hai đại diện tiêu biểu của nhóm fuzzer dựa trên đột biến là AFL++ và libFuzzer. Cả AFL++ và libFuzzer đều là công cụ được dẫn hướng bởi độ coverage: chúng cung cấp cơ chế chen mã đo phủ hiệu quả, quản lý tập trường hợp kiểm thử, các chiến lược tối ưu hóa phép đột biến (lịch trình năng lực, từ điển mẫu, phản hồi từ so sánh, v.v.) và báo cáo các chỉ số thời gian chạy cơ bản như độ coverage, tốc độ thực thi, số trường hợp kiểm thử duy nhất, số lỗi sập và số lỗi treo. Nhờ tốc độ cao và độ ổn định, AFL++ và libFuzzer rất phù hợp cho các kịch bản fuzzing thông thường và là nền tảng thực nghiệm đáng tin cậy khi so sánh hiệu

năng. Tuy nhiên, các công cụ này thường chỉ cung cấp các số liệu thô về độ coverage mà không tích hợp cơ chế phân tích xu hướng tăng trưởng độ coverage theo thời gian (ví dụ xác định hệ số độ dốc hay phát hiện trạng thái bão hòa), đây chính là điểm LR-Zest hướng tới để tự động điều chỉnh chiến lược kiểm thử. Do đó, LR-Zest có thể đóng vai trò bổ trợ cho AFL++ và libFuzzer: vừa tận dụng hạ tầng đo lường và tốc độ của chúng, vừa thêm lớp phân tích theo thời gian thực để quyết định khi nào nên chuyển sang pha sinh đầu vào có cấu trúc hoặc ghép nối nhằm tiếp cận các nhánh ít được khám phá và gia tăng độ coverage một cách hiệu quả hơn.

2. PHƯƠNG PHÁP NGHIÊN CỨU

2.1. Tổng thể quá trình fuzzing

Hình 1 mô tả tổng thể quá trình fuzzing của phương pháp LR-Zest mà bài báo này đề xuất với 2 pha gồm pha Zest Fuzzing và pha Fuzzing trực tiếp. Phương pháp fuzzing LR-Zest cải thiện hiệu quả fuzzing bằng cách khai thác các nhánh mã hiếm khi được thực thi trong chương trình. Quá trình bắt đầu bằng chiến dịch fuzzing thông thường sử dụng Zest để giảm chi phí ban đầu. Khi phát hiện coverage không còn tăng đáng kể (rơi vào trạng thái bão hòa) hệ thống sẽ chuyển sang giai đoạn fuzzing có chủ đích. Tại đây, một nhánh mã hiếm sẽ được chọn làm mục tiêu, sau đó thu thập các input đã từng đi qua nhánh này và thực hiện khai phá mẫu dựa trên cây cấu trúc của chúng. Từ đó, hệ thống trích xuất các đặc trưng đầu vào quan trọng, rồi tái tạo các đặc trưng này bằng cách chèn chúng vào các input được sinh ngẫu nhiên. Mục tiêu là tăng xác suất truy cập nhánh mã hiếm để khám phá thêm mã mới phía sau. Nếu sau 10 phút hoặc khi nhánh không còn hiếm, hệ thống sẽ chọn nhánh mới hoặc quay lại quá trình fuzzing ban đầu nếu không còn nhánh nào phù hợp.



Hình 1. Tổng thể quá trình fuzzing

2.2. Heuristic phát hiện trạng thái bão hòa dựa trên hồi quy tuyến tính.

Khi bắt đầu quá trình fuzzing, thuật toán Zest fuzzing vẫn được sử dụng nhằm đảm bảo sự ổn định cho cả toàn bộ quá trình. Nhằm xác định được thời điểm quá trình fuzzing rơi vào trạng thái bão hòa, nghĩa là độ bao phủ không còn tăng đáng kể theo thời gian, thuật toán heuristic dựa trên hồi quy tuyến tính đã được đề xuất như sau

Tại mỗi thời điểm i , hệ thống lấy một tập gồm k giá trị độ bao phủ gần nhất và ánh xạ tập các giá trị này thành các cặp dữ liệu $\{(t_j, c_j)\}_{j=i-k+1}^i$ trong đó $t_j = j \Delta t$ là mốc thời gian tương ứng ($\Delta t = 1000s$) và c_j là giá trị độ bao phủ ngay tại thời điểm t_j .

Mô hình hồi quy tuyến tính được xây dựng với mục tiêu ước lượng độ dốc b_i đại diện cho tốc độ tăng độ bao phủ trong khoảng thời gian gần nhất được xác định bởi phương trình $\hat{c}_j = a_i + b_i t_j$. Trong đó a_i là hệ số chặn, b_i là hệ số góc được tính theo phương pháp sai số bình phương tối thiểu

$$\bar{t} = \frac{1}{k} \sum_{j=i-k+1}^i t_j, \bar{c} = \frac{1}{k} \sum_{j=i-k+1}^i c_j$$

$$b_i = \frac{\sum_{j=i-k+1}^i (t_j - \bar{t})(c_j - \bar{c})}{\sum_{j=i-k+1}^i (t_j - \bar{t})^2}, a_i = \bar{c} - b_i \bar{t} \quad (1)$$

Nếu b_i nhỏ hơn ngưỡng tối thiểu kỳ vọng b_{min} , hệ thống sẽ xem như độ bao phủ đang tăng quá chậm và chuyển sang trạng thái dừng. Ngưỡng b_{min} được tính toán dựa trên mức tăng tối thiểu kỳ vọng của độ bao phủ trong khoảng thời gian $k \Delta t$, theo công thức $b = \frac{\Delta c_{min}}{k \Delta t_{min}}$. Trong đó, Δc_{min}

là độ bao phủ tối thiểu được kỳ vọng trong một khoảng thời gian dài. Việc so sánh trực tiếp tốc độ tăng độ bao phủ thực tế b_i với tốc độ kỳ vọng b_{min} giúp hệ thống phát hiện trạng thái bão hòa một cách linh hoạt, tránh bị ảnh hưởng bởi các nhiễu nhỏ trong dữ liệu.

2.3. Lựa chọn nhánh mục tiêu

Khi hệ thống phát hiện độ bao phủ bắt đầu rơi vào trạng thái bão hòa, quá trình chuyển sang giai đoạn fuzzing có chủ đích. Bước đầu tiên của giai đoạn này là chọn nhánh mã hiếm làm mục tiêu. Quá trình xác định nhánh mục tiêu dựa trên ý tưởng được trình bày trong nghiên cứu của Caroline Lemieux and Koushik Sen [8]. Nhánh được coi là “hiếm” nếu số lần bị chạm tới thấp hơn một ngưỡng được tính toán động, gọi là cutoff point.

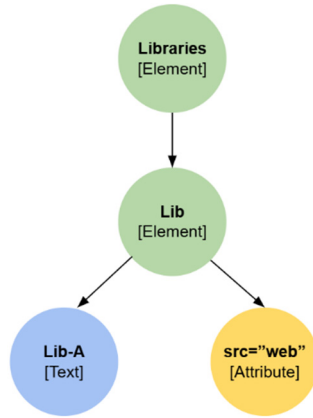
Cụ thể, nếu $min(hits) (> 0)$ là số lần truy cập ít nhất trong tất cả các nhánh tại thời điểm hiện tại, thì cutoff point được định nghĩa là giá trị nhỏ nhất dạng $2^i \geq min(hits)$.

2.4. Khai thác các đặc trưng chung của đầu vào

Sau khi nhánh mục tiêu đã được lựa chọn, các đặc trưng tốt của đầu vào tiếp tục được sàng lọc để có thể phát triển thêm bằng cách sử dụng kỹ thuật khai thác các mẫu tuần tự [9].

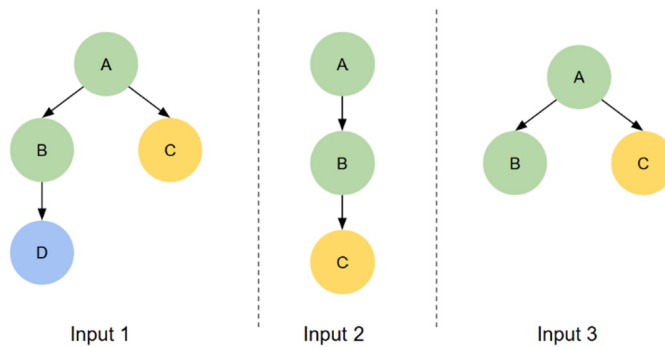
Giả sử đối tượng xét đến là một chuỗi XML có dạng “<libraries><lib src =”web”>”, được biểu diễn dưới dạng cây như Hình 2. Đầu tiên, các đầu vào sẽ được tách làm các đặc trưng, sau đó ghép lại để thành đường đi đặc trưng. Ví dụ như Hình 2, sẽ có 2 đường đi đặc trưng:

- Libraries[E] - Lib[E] - Lib-A[T]
- Libraries[E] - Lib[E] - src=”web”[A]



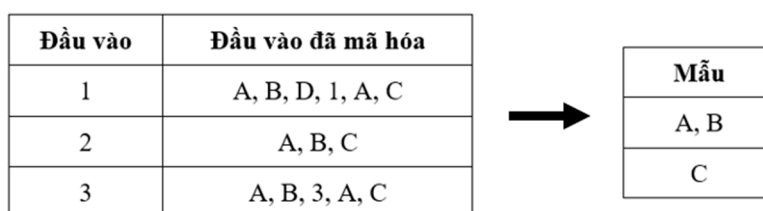
Hình 2. Cấu trúc thành phần của cây XML

Toàn bộ cây sẽ được mã hóa thành một chuỗi ký tự một chiều bằng cách ghép các đặc trưng lại với nhau, như thể hiện trong Hình 3.



Hình 3. Cấu trúc cây của các đặc trưng

Sau đó, thuật toán khai thác chuỗi tuần tự sẽ được sử dụng nhằm tìm ra các đường đi đặc trưng chung. Cuối cùng, các đường đi đặc trưng chung này được chuyển ngược lại thành cây bằng cách hợp nhất các nút có tiền tố giống nhau.



Hình 4. Quá trình khai thác đặc trưng chung

2.5. Tái tạo lại đầu vào với những đặc trưng tốt

Bước cuối trong chu trình fuzzing là chèn cây đặc trưng đã học được vào các input được sinh ngẫu nhiên. Mục tiêu là tạo thêm các test case chứa đặc trưng đầu vào phù hợp để tăng số lần truy cập nhánh mục tiêu. Quá trình chèn không thực hiện một cách ngẫu nhiên mà sẽ thay thế một nút có cùng kiểu trong cây đầu vào nhằm đảm bảo tính cấu trúc tổng thể. Nếu có nhiều cây được sinh ra từ pattern mining, cây con tạo ra nhiều lượt truy cập nhánh nhất trong 100 lần thử đầu tiên sẽ được chọn lựa.

Việc tạo thêm dữ liệu đầu vào mới (Splice) được thực hiện bằng cách biến đổi các Input hiện tại thành các Input mới trong hàng đợi của Zest. Tuy nhiên, không phải Input nào cũng được chọn để Splice. Xác suất Splice phụ thuộc vào mức độ tương đồng về độ bao phủ giữa các đầu vào đã dùng để học đặc trưng và input hiện tại. Nếu độ bao phủ cao, xác suất Splice cũng sẽ cao hơn, vì vậy khả năng thành công của cả chiến dịch fuzzing sẽ lớn hơn.

Giai đoạn fuzzing có Splice kết thúc sau 10 phút hoặc khi nhánh không còn được xem là hiếm. Nếu việc Splice không tạo ra thêm lượt truy cập nhánh mục tiêu, quá trình cũng bị dừng sớm. Sau đó, một nhánh hiếm mới sẽ được chọn và chu trình lặp lại.

3. KẾT QUẢ THỰC NGHIỆM

Để đánh giá thực nghiệm, bốn hệ thống đã được kiểm tra bởi Zest: Apache Maven (v3.5.2), Apache Ant (v1.10.2), Google Closure (v20180204) và Mozilla Rhino (v1.7.8) vẫn tiếp tục được sử dụng với đúng phiên bản đã được kiểm tra.

Các thực nghiệm được tiến hành trên máy chủ chạy hệ điều hành Ubuntu 22.04.5 LTS Virtual Machine, sử dụng bộ xử lý 13th Gen Intel® Core™ i7-13700 2.10 GHz với bộ nhớ RAM 12GB. Mỗi hệ thống được kiểm thử (System Under Test – SUT) như Ant, Maven, Closure và Rhino được chạy trên một máy ảo chủ độc lập để đảm bảo tính cách ly giữa các chiến dịch fuzzing.

Cả ba phương pháp Zest, JQF-Splice và LR-Zest đều dùng chung một kiểu bộ tạo (XML hoặc Javascript) để đảm bảo rằng các nhánh mã có cùng số hiệu giống nhau. Điều này để tránh sự sai lệch giữa các bộ tạo. Các nội dung trong thẻ XML cũng được dùng bởi các từ điển có sẵn của JQF.

Chỉ số đánh giá được dựa trên số lần chạm tới nhánh “hiếm” và giá trị của độ bao phủ cả quá trình. Bảng 1 trình bày kết quả Rare Branch Hits của 3 phương pháp khi chạy với dữ liệu của Apache Ant, Apache Maven, Google Closure và Mozilla Rhino.

Bảng 1. Kết quả Rare Branch Hits của 3 phương pháp Zest, JQF-SPLICE, LR-ZEST

SUT	ZEST	JQF-SPLICE	LR-ZEST
Apache Ant	954205	1104983	1385022
Apache Maven	1662875	1683507	2134947
Google Closure	1244954	1795504	1185315
Mozilla Rhino	545661	516945	632064

Đối với Rare Branch Hits, ta định nghĩa mức cải thiện của LR-Zest so với Zest trên từng hệ thống thử nghiệm (SUT) là:

$$\Delta_i = \frac{X_i^{LR-Zest} - X_i^{Zest}}{X_i^{Zest}} \times 100\%, i = 1, \dots, 4 \quad (2)$$

Kết quả thu được như sau: Ant $\Delta = +45.1\%$, Maven $\Delta = +28.4\%$, Rhino $\Delta = +15.8\%$ và Closure $\Delta = -4.8\%$. Trung vị (median) của các Δ_i đạt $+22.1\%$, cho thấy LR-Zest nhất quán mang lại lợi ích trên đa số SUT.

Kết quả thực nghiệm trên cho thấy LZ-Zest có số lần chạm được nhánh hiếm vượt trội hơn Zest và JQF-Splicing trên hầu hết các SUT, ngoại trừ Google Closure. Nguyên nhân có thể do Google Closure có quá nhiều nhánh hiếm hoặc định dạng Javascript phức tạp hơn nhiều XML, dẫn tới quá trình pattern mining và splicing đơn giản không đủ tái tạo chính xác.

Ngoài ra, kết quả trong Bảng 2 cho thấy độ đo Unique Coverage Trace thể hiện mức độ bao phủ duy nhất của 3 phương pháp trên, tức là số nhánh duy nhất được thực thi khi fuzzing.

Tương tự, đối với Unique Coverage Trace, ta có:

$$\Delta_i = \frac{Y_i^{LR-Zest} - Y_i^{Zest}}{Y_i^{Zest}} \times 100\%, i = 1, \dots, 4 \quad (3)$$

với kết quả: Ant $\Delta = +54.5\%$, Maven $\Delta = +46.6\%$, Rhino $\Delta = +12.1\%$, Closure $\Delta = -1.49\%$. Trung vị của $\Delta_i = +29.4\%$, cho thấy LR-Zest đạt độ coverage cao hơn hẳn so với Zest. Khi so sánh với JQF-Splice, mức cải thiện trung vị cũng đạt tới $+27.1\%$, chứng minh LR-Zest giữ ưu thế ngay cả trước đối thủ gần nhất.

LR-ZEST cũng đạt kết quả cao nhất ở cả 3 SUT Apache Ant, Apache Maven, Mozilla Rhino, cho thấy hiệu quả vượt trội trong việc tạo đầu vào có khả năng bao phủ các nhánh chương trình mới. Đối với Google Closure, LR-Zest chỉ thấp hơn 2 phương pháp còn lại không đáng kể.

Bảng 2. Kết quả Unique Coverage Trace của 3 phương pháp Zest, JQF-SPLICE, LR-ZEST

SUT	ZEST	JQF-SPLICE	LR-ZEST
Apache Ant	70374	72533	108713
Apache Maven	231329	245896	339222
Google Closure	109881	112693	108249
Mozilla Rhino	92246	88948	103438

4. KẾT LUẬN

Nghiên cứu đã đề xuất một phương pháp cải tiến quy trình xác định trạng thái bão hòa trong Fuzzing bằng cách áp dụng hồi quy tuyến tính để phân tích xu hướng gia tăng Coverage trong một khoảng thời gian ngắn. Kỹ thuật này giúp hệ thống nhận diện chính xác hơn thời điểm Coverage có dấu hiệu chững lại, từ đó quyết định chuyển sang giai đoạn Fuzzing có chủ đích nhằm tối ưu hóa tài nguyên và thời gian. Đồng thời, chúng tôi áp dụng chiến lược lựa chọn Input để tái sinh đặc trưng, chuyển từ cơ chế lặp lại các Input hiếm gặp nhất sang khai thác trực tiếp các Input sinh ban đầu từ pha Zest Fuzzing giúp khả năng kích hoạt nhánh mã hiếm hiệu quả hơn.

Kết quả thực nghiệm cho thấy phương pháp này mang lại lợi ích rõ rệt về số lần truy cập nhánh hiếm và số lượng trace phủ mã duy nhất. Tuy nhiên, phương pháp chưa vượt trội trong việc cải thiện về tổng độ coverage hoặc khả năng phát hiện lỗi. Một vài khả năng được đặt ra là các vùng mã hiếm mà phương pháp hướng tới có thể không dẫn đến mã mới, hoặc các mẫu được Splice vào chưa phản ánh đầy đủ đặc trưng đầu vào cần thiết để kích hoạt hành vi mới. Các kết quả trên đã gợi mở rằng các định dạng đầu vào phức tạp có thể hưởng lợi nhiều hơn từ phương pháp đề xuất LR-Zest, đặc biệt trong thời gian Fuzzing kéo dài.

Đối với hướng phát triển tương lai, chúng tôi dự định áp dụng phương pháp khai phá mẫu phức tạp hơn, chẳng hạn như dựa trên đồ thị (graph-based mining); Tập trung vào các vùng mã còn nhiều đoạn chưa được bao phủ (uncovered) và Tích hợp kỹ thuật kiểm thử dựa trên ngữ pháp (grammar-based fuzzing), nhằm sinh ra đầu vào có cấu trúc chặt chẽ, phù hợp với các kiểu dữ liệu phức tạp.

5. TÀI LIỆU THAM KHẢO

- [1] Zhao, X., Qu, H., Xu, J., Li, X., Lv, W., and Wang, G. G. 2024. A systematic review of fuzzing. *Soft Computing*, 28(6), 5493-5522.
- [2] Beaman, C., Redbourne, M., Mummery, J. D., and Hakak, S. 2022. Fuzzing vulnerability discovery techniques: Survey, challenges and future directions. *Computers & Security*, 120, 102813.
- [3] Huang, L., Zhao, P., Chen, H., and Ma, L. 2024. Large language models based fuzzing techniques: A survey. *arXiv-2402*.
- [4] Yu, Z., Liu, Z., Cong, X., Li, X., and Yin, L. 2024. Fuzzing: Progress, Challenges, and Perspectives. *Computers, Materials & Continua*, 78(1).
- [5] Kraus, Roman, Hoang Lam Nguyen, and Martin A. Schneider. 2024. Generator-based Fuzzing with Input Features. *Proceedings of the 17th ACM/IEEE International Workshop on Search-Based and Fuzz Testing*.
- [6] Nguyen, Hoang Lam, and Lars Grunske. 2022. Bedivfuzz: Integrating behavioral diversity into generator-based fuzzing. *Proceedings of the 44th international conference on software engineering*.
- [7] Kraus, R., Nguyen, H. L., and Schneider, M. A. 2024, April. Generator-based Fuzzing with Input Features. In *Proceedings of the 17th ACM/IEEE International Workshop on Search-Based and Fuzz Testing* (pp. 13-20).
- [8] Caroline Lemieux and Koushik Sen. 2018. FairFuzz: a targeted mutation strategy for increasing greybox fuzz testing coverage. In *ASE. ACM*, 475–485.
- [9] Fournier-Viger, P., Lin, J. C. W., Kiran, R. U., Koh, Y. S., and Thomas, R. 2017. A survey of sequential pattern mining. *Data Science and Pattern Recognition*, 1(1), 54-77.

