

HỌC VIỆN CÔNG NGHỆ BƯU CHÍNH VIỄN THÔNG



TRẦN HỒNG NHUNG

**CÁC MÔ HÌNH HỌC SÂU CHO TÍNH TOÁN
ĐỘ TIN CẬY TRONG MẠNG PHỨC TẠP**

Chuyên ngành : Khoa học máy tính

Mã số : 8.48.01.01

TÓM TẮT ĐỀ ÁN TỐT NGHIỆP THẠC SĨ

HÀ NỘI - NĂM 2025

Đề án tốt nghiệp được hoàn thành tại:
HỌC VIỆN CÔNG NGHỆ BƯU CHÍNH VIỄN THÔNG

Người hướng dẫn khoa học: PGS.TS. Trần Đình Quế

Phản biện 1:

Phản biện 2:

Đề án tốt nghiệp sẽ được bảo vệ trước Hội đồng chấm đề án tốt nghiệp thạc sĩ tại
Học viện Công nghệ Bưu chính Viễn thông

Vào lúc: giờ ngày tháng năm

Có thể tìm hiểu đề án tốt nghiệp tại:

- Thư viện của Học viện Công nghệ Bưu chính Viễn thông.

MỞ ĐẦU

1. Lý do chọn đề tài

Mạng phức tạp ngày càng trở nên phổ biến và có tầm ảnh hưởng sâu sắc đến nhiều đối tượng trong đời sống hàng ngày. Các mạng này tồn tại song song với thế giới thực, đồng thời tạo ra những tác động qua lại đáng kể đến đời sống thực tế. Khi con người dành ngày càng nhiều thời gian cho các hoạt động trên mạng phức tạp, các quyết định và công việc thường ngày của họ cũng bị ảnh hưởng không nhỏ bởi các tương tác, thông tin thu thập được từ bạn bè và các kết nối trong mạng.

Trong bối cảnh đó, nhu cầu đánh giá mức độ tin cậy của các đối tượng, cộng đồng, hoặc nội dung trên mạng phức tạp trở thành một yếu tố thiết yếu, giúp người dùng đưa ra quyết định phù hợp và giảm thiểu rủi ro. Tin cậy đã được nghiên cứu trong nhiều lĩnh vực như tâm lý học, triết học, xã hội học và khoa học máy tính, qua đó chỉ ra rằng tin cậy là một khái niệm chủ quan, thay đổi tùy thuộc vào từng cá nhân, tình huống và bối cảnh cụ thể.

Một thách thức lớn đối với các nhà nghiên cứu là việc định nghĩa rõ ràng khái niệm tin cậy, mô tả quá trình hình thành tin cậy và phân tích tác động của nó lên con người. Các nghiên cứu tiếp theo đã tập trung vào việc xác định các yếu tố tiền đề của tin cậy, tức là những yếu tố có khả năng ảnh hưởng đến mức độ tin cậy. Tuy nhiên, việc tổng hợp và đánh giá mức độ ảnh hưởng của các yếu tố này trong những bối cảnh khác nhau vẫn là một thách thức đáng kể.

Trước những khó khăn đó, đề án này tập trung nghiên cứu và làm rõ khái niệm về tin cậy, cụ thể hóa các yếu tố tiền đề ảnh hưởng đến tin cậy. Từ đó, đề án đề xuất các phương pháp đo lường mức độ tin cậy và xây dựng các mô hình tin cậy phù hợp cho mạng phức tạp. Đề án bắt đầu bằng việc tìm hiểu phương pháp mô hình hóa một mạng phức tạp. Theo đó, một mạng phức tạp trực tuyến có thể được mô hình hóa dưới dạng một đồ thị có hướng, trong đó các nút biểu thị người dùng và các cạnh biểu thị mối quan hệ giữa chúng. Hướng của cạnh sẽ chỉ ra người nào được xác định tin cậy.

2. Mục tiêu nghiên cứu

Mục đích của nghiên cứu là phát triển và áp dụng các mô hình học sâu, đặc biệt là CNN và LSTM, để tính toán độ tin cậy trong mạng phức tạp. Mục tiêu bao gồm:

- **Thiết kế và huấn luyện các mô hình học sâu** để trích xuất đặc trưng không gian và thời gian từ dữ liệu mạng, nhằm dự đoán độ tin cậy chính xác.
- **Tích hợp yếu tố không chắc chắn và mối quan hệ phức tạp** để cải thiện độ chính xác khi dữ liệu không hoàn chỉnh hoặc có biến động.

- **Tối ưu hóa mô hình học sâu** để cải thiện tốc độ huấn luyện, giảm yêu cầu dữ liệu và nâng cao khả năng dự đoán độ tin cậy.

3. Đối tượng và phạm vi nghiên cứu

Nghiên cứu tập trung vào các mạng có cấu trúc phức tạp, đặc trưng bởi mật độ kết nối cao, tính phân tán, và sự biến động động lực học của các yếu tố trong mạng. Những đặc điểm này sẽ được phân tích và tích hợp vào mô hình đánh giá độ tin cậy nhằm phản ánh chính xác bản chất phi tuyến và khó dự đoán của hệ thống.

4. Phương pháp và phạm vi nghiên cứu

Phạm vi nghiên cứu của đề tài bao gồm các khía cạnh sau:

- Các loại mạng phức tạp: Nghiên cứu sẽ tập trung vào các mạng phức tạp như mật độ kết nối, tính phân tán và sự biến động của các yếu tố trong mạng sẽ được phân tích và đưa vào mô hình.
- Các mô hình học sâu sử dụng trong nghiên cứu: Nghiên cứu sẽ khảo sát và sử dụng một số mô hình học sâu:
 - Mạng học sâu tích chập (CNN)
 - Mạng hồi tiếp dài ngắn (LSTM)
 - Mạng nơ-ron hồi tiếp (RNN)
 - Dữ liệu nghiên cứu:

- Dữ liệu nghiên cứu:

Dữ liệu sử dụng trong nghiên cứu này sẽ được thu thập từ các mô hình mạng phức tạp được mô phỏng hoặc từ các cơ sở dữ liệu thực tế có sẵn trong các lĩnh vực mạng viễn thông, mạng Internet và giao thông. Việc lựa chọn dữ liệu từ các nguồn này nhằm mục đích tạo ra một bức tranh toàn diện về các yếu tố ảnh hưởng đến độ tin cậy mạng, đặc biệt là trong các môi trường mạng có tính phức tạp cao và thay đổi nhanh chóng.

Nghiên cứu sẽ sử dụng bộ dữ liệu từ tập Stanford Large Network Dataset Collection <https://snap.stanford.edu/data>. Bộ sưu tập này cung cấp các đồ thị mạng thực tế với các đặc trưng cấu trúc mạng phong phú, từ đó giúp mô phỏng các tình huống thực tế hơn trong nghiên cứu về độ tin cậy mạng. Việc sử dụng dữ liệu thực tế từ Stanford sẽ giúp kiểm tra độ chính xác và tính khả thi của các mô hình trong môi trường thực tế, nơi mà các yếu tố phức tạp và không thể dự đoán trước (như sự cố mạng, lưu lượng giao thông thay đổi, v.v.) có thể xảy ra bất ngờ.

5. Phương pháp lý thuyết và phân tích

- **Phân tích đặc điểm mạng phức tạp:** Nghiên cứu sẽ bắt đầu bằng việc phân tích các đặc điểm của mạng phức tạp, bao gồm cấu trúc đồ thị của mạng, các yếu tố ảnh hưởng đến độ tin cậy (như độ bền của kết nối, mật độ mạng, tính phân tán của các thành phần trong mạng, v.v.). Việc phân tích này sẽ giúp xác định các yếu tố quan trọng cần được đưa vào mô hình học sâu.
- **Xây dựng mô hình độ tin cậy mạng:** trên cơ sở các lý thuyết đồ thị và lý thuyết về độ tin cậy, nghiên cứu sẽ xây dựng các mô hình cơ bản để tính toán độ tin cậy mạng, làm cơ sở để so sánh với kết quả thu được từ các mô hình học sâu sau này
- **Nghiên cứu các mô hình học sâu:** nghiên cứu sẽ phân tích các mô hình học sâu CNN, RNN và LSTM. Các mô hình này sẽ được chọn lựa dựa trên khả năng xử lý dữ liệu phức tạp và các yếu tố không chắc chắn trong mạng phức tạp.

Phương pháp mô phỏng và thực nghiệm

- Sử dụng các bộ dữ liệu thực tế từ trang <https://snap.stanford.edu/data>, bao gồm thông tin về cấu trúc mạng, kết nối, tỷ lệ hỏng hóc và yếu tố môi trường.
- Mô hình học sâu được huấn luyện: CNN, RNN và LSTM.
- Đánh giá mô hình: So sánh kết quả tính toán độ tin cậy giữa các mô hình dựa trên các chỉ số: độ chính xác, tốc độ huấn luyện và khả năng tổng quát hóa.

Phương pháp phân tích kết quả

- Phân tích hiệu quả mô hình: Đánh giá chi tiết các mô hình CNN, RNN, LSTM dựa trên độ chính xác và độ tin cậy.
- Đánh giá khả năng áp dụng và tổng quát hóa: Kiểm tra hiệu quả của mô hình trên các mạng phức tạp khác nhau, như mạng có cấu trúc biến động hoặc không đồng nhất, để xác minh khả năng mở rộng ra các bối cảnh ứng dụng khác.

6. Cấu trúc đề án

Đề án tốt nghiệp có cấu trúc như sau:

- Chương 1: Tổng quan về mô hình tin cậy và học sâu
- Chương 2: Phân tích đặc trưng dữ liệu và các mô hình học sâu
- Chương 3: Thử nghiệm và đánh giá

CHƯƠNG 1: TỔNG QUAN VỀ MÔ HÌNH TIN CẬY VÀ HỌC SÂU

1.1 Giới thiệu

1.1.1 Khái niệm và các thuộc tính của mạng phức tạp

Mạng phức tạp là một dạng mô hình đồ thị gồm các nút (đại diện cho thực thể như người dùng, thiết bị, tổ chức) và cạnh (đại diện cho mối quan hệ hoặc tương tác). Mạng có thể có hướng hoặc vô hướng, có trọng số hoặc không trọng số, và có thể chứa đa quan hệ giữa các nút.

Mạng phức tạp có vai trò quan trọng trong việc mô hình hóa các hệ thống trong mạng xã hội, viễn thông, IoT, v.v. Với hơn 5 tỷ người dùng Internet toàn cầu (2023), mạng phức tạp ngày càng ảnh hưởng sâu rộng đến đời sống thực tế và các hoạt động trực tuyến.

Các thuộc tính nổi bật của mạng phức tạp gồm:

- **Tính liên kết cao:** Một số nút có nhiều kết nối vượt trội, làm tăng vai trò trung tâm trong mạng.
- **Tính động:** Mạng thay đổi liên tục do thêm hoặc xóa nút và cạnh.
- **Tính phân cụm:** Các nút có xu hướng hình thành nhóm với mật độ liên kết nội bộ cao.
- **Tính ảnh hưởng và lan truyền:** Một số nút có ảnh hưởng lớn, có thể khuếch đại thông tin (influencer).
- **Tính phi tuyến và không đồng nhất:** Mạng phân phối không đều; một số nút chiếm phần lớn liên kết, tuân theo phân phối lũy thừa.

Những đặc điểm này khiến mạng phức tạp trở thành đối tượng lý tưởng để phân tích hành vi, cá nhân hóa nội dung, phát hiện cộng đồng và đánh giá độ tin cậy bằng các công cụ như mô hình học sâu.

1.1.2 Phân tích mạng phức tạp

Phân tích mạng phức tạp là phương pháp sử dụng lý thuyết đồ thị để nghiên cứu cấu trúc và hành vi của các hệ thống gồm nhiều thực thể tương tác, trong đó các thực thể được mô hình hóa thành nút và các mối quan hệ là cạnh. Việc phân tích tập trung vào các chỉ số như tính trung tâm (đo mức độ ảnh hưởng của một nút), mật độ (mức độ liên kết của toàn mạng), và phát hiện cộng đồng (nhóm các nút liên kết chặt chẽ). Một nội dung quan trọng khác là dự đoán liên kết – nhằm ước lượng khả năng hình thành kết nối mới, ứng dụng trong gợi ý bạn bè, đối tác hoặc thông tin.

Phân tích mạng phức tạp có vai trò thiết yếu trong nhiều lĩnh vực như y tế (mô hình lây lan dịch bệnh), kinh doanh (xếp hạng trang web bằng PageRank), và an ninh mạng (phát hiện hành vi bất thường). Nhờ đó, phương pháp này cung cấp nền tảng vững chắc để hiểu và khai thác các đặc trưng động, phi tuyến của hệ thống mạng – từ đó hỗ trợ hiệu quả cho các mô hình học sâu trong việc đánh giá độ tin cậy.

1.1.3 Cộng đồng người dùng trên các trang mạng phức tạp

Trong các mạng phức tạp như mạng xã hội, cộng đồng người dùng được hình thành từ các nhóm cá nhân có mức độ tương tác cao và chia sẻ những đặc điểm chung như sở thích, lĩnh vực chuyên môn hoặc quan điểm xã hội. Những cộng đồng này không tồn tại một cách ngẫu nhiên, mà được cấu trúc dựa trên các mối liên kết hình thành qua thời gian và được củng cố nhờ lịch sử tương tác, mức độ tin tưởng và sự tương đồng giữa các thành viên.

Tính chất mở và linh hoạt của các cộng đồng trực tuyến cho phép người dùng dễ dàng gia nhập, rời bỏ hoặc tương tác chéo giữa các nhóm, tạo nên một hệ sinh thái đa tầng, phức tạp và năng động. Đặc biệt, hiện tượng homophily – tức xu hướng kết nối với những người có đặc điểm giống mình – đóng vai trò quan trọng trong việc duy trì độ gắn kết của cộng đồng, đồng thời ảnh hưởng mạnh đến quá trình hình thành niềm tin, chia sẻ thông tin và khuếch tán ảnh hưởng.

Trong bối cảnh đó, các yếu tố như lịch sử tương tác (ví dụ: số lần trao đổi, bình luận tích cực), vai trò của cá nhân trong mạng (ví dụ: người quản trị, thành viên chủ chốt), và mức độ tương đồng về giá trị hay hành vi là những cơ sở chính để đánh giá độ tin cậy giữa các thành viên. Một thành viên thường xuyên chia sẻ nội dung chất lượng, nhận được phản hồi tích cực, hoặc có độ trung tâm cao trong mạng sẽ được cộng đồng mặc định xem là đáng tin cậy hơn.

Sự tồn tại và đặc trưng của các cộng đồng người dùng trên mạng không chỉ phản ánh hành vi xã hội số, mà còn đặt nền tảng cho việc xây dựng các mô hình học sâu nhằm dự đoán hoặc suy diễn độ tin cậy trong các hệ thống mạng phức tạp.

1.1.4 Các vấn đề liên quan đến sự quan tâm người dùng

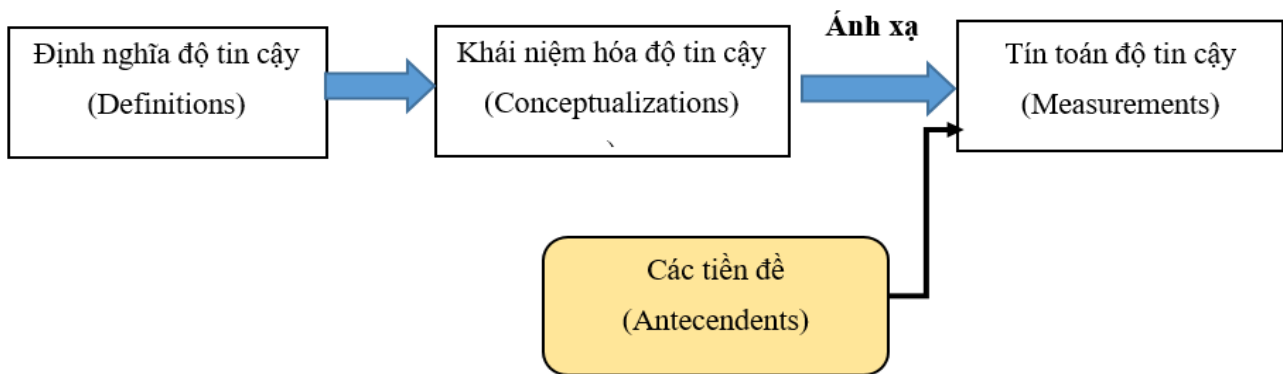
Sự quan tâm của người dùng trong mạng phức tạp không chỉ bị chi phối bởi nội dung hiển thị mà còn phụ thuộc vào hành vi tương tác và các thuật toán cá nhân hóa. Hiệu ứng “bong bóng lọc” khiến người dùng bị giới hạn trong phạm vi thông tin quen thuộc, làm giảm khả năng tiếp cận đa chiều. Bên cạnh đó, yếu tố tin cậy đóng vai trò then chốt trong việc người dùng lựa chọn và tiếp nhận thông tin. Niềm tin này có thể được xây dựng từ trải nghiệm, lịch sử tương tác hoặc các tín hiệu gián tiếp như lượt theo dõi hay đánh giá tích cực. Đề án đề xuất tích hợp các yếu tố tương đồng và tin cậy vào mô hình học sâu nhằm nâng cao chất lượng đề xuất nội dung và giảm thiểu rủi ro từ thông tin sai lệch.

1.2 Tin cậy

Tin cậy là một khái niệm đa chiều, phức tạp và phụ thuộc mạnh vào ngữ cảnh, đặc biệt trong

môi trường mạng phức tạp. Trong các mạng này, độ tin cậy không chỉ dựa vào trải nghiệm trực tiếp mà còn được hình thành qua lịch sử tương tác, vai trò xã hội, và sự tương đồng giữa các thành viên. Các định nghĩa học thuật đều thừa nhận rằng tin cậy bao gồm yếu tố kỳ vọng tích cực từ phía người tin tưởng trong bối cảnh có thể bị tổn thương.

Việc đo lường độ tin cậy đòi hỏi tiếp cận đa phương diện, thông qua các chỉ số như độ trung tâm, mức độ tương tác, tính đồng thuận và uy tín cộng đồng. Bên cạnh đó, các yếu tố như phi đối xứng, phụ thuộc ngữ cảnh, biến đổi theo thời gian và tính lan truyền khiến độ tin cậy trở nên động và khó mô hình hóa bằng các phương pháp tuyến tính truyền thống. Do vậy, việc ứng dụng mô hình học sâu như CNN và LSTM giúp khai thác tốt hơn đặc trưng thời gian, cấu trúc và mối quan hệ phức hợp của độ tin cậy trong mạng, góp phần nâng cao độ chính xác trong dự đoán và đánh giá.



Hình 1.4 Quy trình đánh giá độ tin cậy

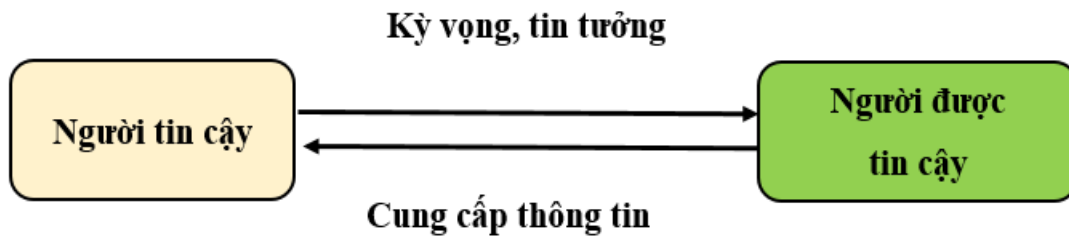
1.2.1 Định nghĩa độ tin cậy

Độ tin cậy là trạng thái tâm lý phản ánh kỳ vọng tích cực của một cá nhân hoặc hệ thống vào hành vi của bên khác trong bối cảnh có rủi ro hoặc sự dễ bị tổn thương. Cốt lõi của khái niệm này nằm ở việc người tin tưởng chấp nhận khả năng bị tổn hại dựa trên niềm tin rằng bên còn lại sẽ hành động vì lợi ích hoặc không gây hại.

Trong các lĩnh vực như triết học, tâm lý học và khoa học máy tính, độ tin cậy được tiếp cận từ các góc nhìn khác nhau nhưng đều nhấn mạnh vai trò của niềm tin, kinh nghiệm quá khứ, uy tín và sự không chắc chắn trong tương tác. Đặc biệt trong môi trường số, độ tin cậy không chỉ hình thành từ trải nghiệm trực tiếp mà còn từ các tín hiệu gián tiếp như đánh giá, phản hồi, hoặc xác thực từ cộng đồng.

Do đó, định nghĩa độ tin cậy trong mạng phức tạp cần được hiểu như một hiện tượng động, đa chiều, phụ thuộc ngữ cảnh và có thể được mô hình hóa thông qua dữ liệu tương tác,

lịch sử hoạt động và vị trí cấu trúc trong mạng.



Hình 1.5 Sơ đồ trung gian của người tin cậy và người được tin cậy

1.2.2 Các tiền đề tính toán độ tin cậy

Độ tin cậy trong mạng phức tạp không hình thành một cách ngẫu nhiên, mà được xác lập dựa trên nhiều yếu tố nền tảng, được gọi là **các tiền đề**. Các tiền đề này có thể chia thành ba nhóm chính:

- **Tiền đề nhân khẩu học:** Bao gồm các yếu tố như tuổi, giới tính, trình độ học vấn, văn hóa hoặc tôn giáo. Những yếu tố này ảnh hưởng đến cách mỗi cá nhân đánh giá và đặt niềm tin.
- **Tiền đề tương tác:** Liên quan đến mức độ thân quen, lịch sử giao tiếp, tần suất tương tác, và đặc biệt là **tính tương đồng (homophily)** – xu hướng con người tin tưởng những người có đặc điểm giống mình về sở thích, giá trị hay hành vi.
- **Tiền đề thuộc về người được tin cậy:** Gồm uy tín xã hội (dựa trên lượt thích, chia sẻ, theo dõi), phong cách thể hiện và nội dung được chia sẻ. Những cá nhân có ảnh hưởng cao hoặc được cộng đồng đánh giá tích cực thường được xem là đáng tin hơn.

Việc nhận diện và tích hợp các tiền đề này là cơ sở quan trọng để mô hình hóa và dự đoán độ tin cậy một cách chính xác, đặc biệt khi ứng dụng trong các hệ thống học sâu hoặc mạng phức tạp có nhiều yếu tố biến động.

1.2.3 Các giá trị của độ tin cậy

Trong mạng phức tạp, độ tin cậy có thể được biểu diễn dưới nhiều hình thức khác nhau, tùy thuộc vào mục tiêu phân tích và phương pháp mô hình hóa. Ba dạng phổ biến nhất gồm:

- **Độ tin cậy nhị phân:** Biểu diễn mối quan hệ tin tưởng dưới dạng có/không (1 hoặc 0). Đây là cách đơn giản, phù hợp với các bài toán phân loại, nhưng thiếu tính linh hoạt trong thể hiện mức độ niềm tin.

- **Độ tin cậy có trọng số:** Biểu diễn niềm tin dưới dạng liên tục (ví dụ từ 0 đến 1), cho phép mô hình học sâu như LSTM hoặc GAT học được cường độ khác nhau trong quan hệ tin cậy, phản ánh sát hơn thực tế.
- **Độ tin cậy theo thời gian:** Là sự thay đổi giá trị tin cậy theo chuỗi thời gian, ghi nhận sự gia tăng hoặc suy giảm niềm tin giữa các thực thể. Đây là dạng biểu diễn quan trọng trong các mạng động, cần đến các mô hình như LSTM để khai thác quan hệ phụ thuộc theo thời gian.

Việc lựa chọn dạng biểu diễn phù hợp giúp mô hình học sâu hiểu rõ hơn về bản chất của niềm tin, từ đó nâng cao khả năng dự đoán chính xác trong các tình huống thực tế có nhiều biến động và tính bất định cao.

1.2.4 Các thuộc tính của tin cậy

Độ tin cậy trong mạng phức tạp không chỉ là một giá trị định lượng, mà còn mang nhiều thuộc tính đặc trưng phản ánh bản chất động và phi tuyến của mối quan hệ giữa các thực thể. Các thuộc tính chính bao gồm:

- **Tính phi đối xứng:** Mức độ tin tưởng giữa hai thực thể không nhất thiết phải bằng nhau ($A \text{ tin } B \neq B \text{ tin } A$), đòi hỏi mô hình hóa theo cặp quan hệ riêng biệt.
- **Phụ thuộc ngữ cảnh:** Niềm tin thay đổi theo từng hoàn cảnh hoặc lĩnh vực cụ thể, ví dụ một người có thể đáng tin trong chuyên môn nhưng không trong tài chính.
- **Tính biến động theo thời gian:** Độ tin cậy không cố định mà thay đổi theo lịch sử tương tác; vì vậy cần mô hình chuỗi như LSTM để phản ánh xu hướng tăng – giảm của niềm tin.
- **Tính lan truyền (transitivity):** Niềm tin có thể lan từ người này sang người khác qua các đỉnh trung gian, làm nảy sinh chuỗi tin cậy gián tiếp trong mạng.

Nhận diện và mô hình hóa các thuộc tính này giúp cải thiện đáng kể khả năng dự đoán độ tin cậy trong các mạng động, phi tuyến và có cấu trúc phức tạp.

1.2.5 Vai trò của một đỉnh đến quá trình lan truyền độ tin cậy

Trong mạng phức tạp, mỗi đỉnh (nút) không chỉ là điểm lưu trữ thông tin, mà còn đóng vai trò trung gian trong quá trình lan truyền độ tin cậy giữa các thực thể. Các đỉnh giữ vị trí quan trọng như **hub** (nút có nhiều liên kết), **bridge node** (nối giữa các cụm cộng đồng) hoặc **nút có tương tác cao theo thời gian** thường có ảnh hưởng lớn đến dòng chảy niềm tin trong mạng.

Sự gián đoạn tại các đỉnh này – do mất kết nối, thiếu dữ liệu, hoặc hành vi bất thường – có thể làm đứt gãy chuỗi lan truyền, gây sai lệch trong việc suy diễn độ tin cậy giữa các nút khác. Đặc biệt trong mạng động hoặc có cấu trúc không đồng đều, việc một đỉnh trung gian bị loại bỏ có thể ảnh hưởng nghiêm trọng đến toàn bộ hệ thống đánh giá.

Do đó, việc nhận diện và mô hình hóa đúng vai trò của từng đỉnh là thiết yếu để đảm bảo tính chính xác, liên tục và khả năng mở rộng của các mô hình học sâu trong bài toán tính toán độ tin cậy mạng.

1.3 Mô hình tin cậy

Mô hình tin cậy là cách tiếp cận khoa học nhằm định lượng và dự đoán mức độ tin tưởng giữa các thực thể trong mạng phức tạp. Độ tin cậy người dùng trong ngữ cảnh đề cập đến mức độ tin tưởng giữa mọi người trong các mạng xã hội khác nhau, được hình thành từ các tương tác, hành vi, và bối cảnh trực tuyến, đồng thời chịu ảnh hưởng trực tiếp từ độ tin cậy của các tài nguyên mà họ sử dụng. Đề án sẽ giới thiệu một vài mô hình tin cậy để có thể đánh giá tổng quan xung quanh việc đánh giá độ tin cậy trong mạng phức tạp.

1.3.1 Phân loại mô hình tin cậy

Mô hình tin cậy được phân loại dựa trên **đối tượng được đánh giá**, phản ánh cách tiếp cận khác nhau trong từng ngữ cảnh ứng dụng. Theo đó, có ba nhóm chính:

- **Độ tin cậy nội dung:** Đánh giá mức độ đáng tin của thông tin được chia sẻ (bài viết, bình luận, đánh giá), thường dựa trên ngữ nghĩa, tính mạch lạc, và sự phản hồi tích cực. Các mô hình NLP như LSTM, BERT được sử dụng để trích xuất đặc trưng và học biểu diễn tin cậy từ văn bản.
- **Độ tin cậy dịch vụ:** Đề cập đến hiệu suất và uy tín của hệ thống hoặc dịch vụ trực tuyến (ví dụ: website, nền tảng thương mại điện tử). Chỉ số như PageRank, tên miền uy tín, và đánh giá từ người dùng thường được dùng để xác định.
- **Độ tin cậy người dùng:** Là mức độ tin tưởng giữa các cá nhân trong mạng, dựa trên lịch sử tương tác, uy tín xã hội và hành vi trong cộng đồng. Đây là loại phổ biến trong các hệ thống mạng xã hội, gợi ý và cộng tác số.

1.3.2 Các nhóm độ đo tin cậy trong mạng xã hội

Trong mạng xã hội, độ tin cậy có thể được ước lượng thông qua nhiều nhóm độ đo khác nhau, phản ánh các khía cạnh cấu trúc, hành vi và cộng đồng của người dùng. Ba nhóm chính bao gồm:

- **Độ đo dựa trên cấu trúc mạng:** Sử dụng đặc điểm kết nối giữa các nút để đánh giá độ tin cậy. Các chỉ số phổ biến gồm:
 - *Degree Centrality* – số lượng liên kết,
 - *Closeness Centrality* – mức độ tiếp cận,
 - *Betweenness Centrality* – vai trò cầu nối,
 - *Eigenvector Centrality* – ảnh hưởng gián tiếp thông qua các nút quan trọng.
 Các chỉ số này giả định rằng nút có vị trí trung tâm và kết nối tốt thì đáng tin cậy hơn.
- **Độ đo dựa trên cộng đồng:** Giả định rằng các nút trong cùng một cộng đồng thường có độ tin cậy cao hơn nhờ chia sẻ mục tiêu, ngữ cảnh và tương tác thường xuyên. Một số cách tiếp cận gồm: đo độ chồng lấp cộng đồng, hoặc đánh giá tin cậy dựa trên bán kính cộng đồng lân cận (*k-hop community*).
- **Độ đo dựa trên lịch sử tương tác:** Khai thác trực tiếp dữ liệu hành vi như tần suất tương tác, phản hồi tích cực, tốc độ phản hồi, hay độ ổn định theo thời gian. Các mô hình thường dùng như trung bình trọng số, danh tiếng tích lũy, hoặc Bayesian trust cho phép ước lượng độ tin cậy từ quá khứ.

Mỗi nhóm độ đo có ưu điểm riêng và thường được kết hợp trong các mô hình học sâu để nâng cao khả năng đánh giá tin cậy trong mạng phức tạp, đặc biệt khi dữ liệu không đầy đủ hoặc có tính động.

1.4 Kết luận chương 1

Chương 1 đã trình bày tổng quan về mạng phức tạp, khái niệm độ tin cậy và các cách tiếp cận mô hình hóa tin cậy trong bối cảnh mạng xã hội và hệ thống phân tán. Qua việc phân tích các đặc tính cấu trúc, hành vi và ngữ cảnh của mạng phức tạp, có thể thấy rằng độ tin cậy là một khái niệm động, phi tuyến và phụ thuộc mạnh vào lịch sử tương tác cũng như vị trí cấu trúc của từng đỉnh trong mạng. Từ đó, chương 1 tạo cơ sở lý thuyết và thực tiễn vững chắc cho các chương tiếp theo, nơi các mô hình học sâu như CNN, LSTM và các biến thể tích hợp sẽ được triển khai nhằm giải quyết bài toán dự đoán độ tin cậy một cách hiệu quả trong môi trường mạng phức tạp.

CHƯƠNG 2: PHÂN TÍCH ĐẶC TRƯNG DỮ LIỆU VÀ CÁC MÔ HÌNH HỌC SÂU

Chương 2 tập trung mở rộng cách tiếp cận tính toán độ tin cậy trong các mạng phức hợp thông qua việc kết hợp mức độ tương tác, cấp độ quen thuộc, dữ liệu phân tán và yếu tố ảnh hưởng. Độ tin cậy được xem như một quan hệ định lượng giữa bên tin tưởng và bên được tin tưởng trong bối cảnh có sự tương tác.

Để thực hiện điều này, chương ứng dụng các mô hình học sâu như **CNN**, **RNN** và **LSTM** nhằm dự đoán độ tin cậy trên nhiều bộ dữ liệu khác nhau. Các thí nghiệm được thiết kế để khảo sát tác động của tương tác và ảnh hưởng đến hiệu quả dự đoán. Kết quả thực nghiệm cho thấy rằng việc tích hợp đồng thời cả tương tác và yếu tố ảnh hưởng giúp mô hình đạt độ chính xác cao hơn so với các phương pháp chỉ dựa trên dữ liệu tương tác thuần túy.

2.1 Phân tích và lựa chọn đặc trưng dữ liệu

Việc lựa chọn đặc trưng dữ liệu là bước quan trọng trong quá trình xây dựng mô hình dự đoán độ tin cậy bằng học sâu. Trong mục này, tác giả phân tích các yếu tố dữ liệu có ảnh hưởng trực tiếp đến độ tin cậy, bao gồm: lịch sử tương tác giữa các cặp người dùng, số lần đánh giá, mức độ quen thuộc, thời gian phản hồi, và hành vi đánh giá.

Các đặc trưng được chọn phải đảm bảo khả năng biểu diễn mối quan hệ tin cậy, đồng thời tương thích với đầu vào của mô hình học sâu như chuỗi thời gian hoặc ma trận đặc trưng. Để đảm bảo tính tổng quát và loại bỏ nhiễu, dữ liệu được tiền xử lý thông qua các bước chuẩn hóa, mã hóa thời gian và tạo cặp người dùng có mối quan hệ tin cậy lẫn không tin cậy.

Việc phân tích kỹ đặc trưng trước khi huấn luyện giúp mô hình học sâu tập trung vào những tín hiệu mang tính phân biệt cao, từ đó nâng cao độ chính xác và khả năng khái quát của mô hình khi áp dụng trên dữ liệu thực tế.

2.2 Công thức tổng quan về mạng phức tạp

2.2.1 Mạng phức tạp (Complex network)

Một mạng phức hợp được biểu diễn dưới dạng đồ thị $G = (V, E)$ trong đó V là tập hợp các nút, và E là tập hợp các cạnh. Số lượng nút là $n = |V|$, và số lượng cạnh là $m = |E|$. Trong mạng, mỗi nút có thể gửi và nhận thông điệp từ các nút khác. Dựa trên công trình trước đây [30], bài báo này không xét đến nội dung cụ thể của các thông điệp mà tập trung vào các loại và số

lượng tương tác hoặc phản hồi, được chuyển tiếp qua các kết nối giữa các nút. Ngoài ra, chúng tôi còn sử dụng các chỉ số trong cấu trúc mạng để xây dựng thước đo mức độ ảnh hưởng, từ đó góp phần vào việc ước lượng giá trị độ tin cậy.

2.2.2 Mức độ quen biết (*Familiarity*)

Familiarity là một đặc trưng cốt lõi trong việc đánh giá độ tin cậy trong mạng phức tạp, đặc biệt trong các mạng xã hội, nơi niềm tin được hình thành chủ yếu từ các tương tác giữa người dùng. Đặc trưng này phản ánh mức độ thân quen giữa hai nút thông qua các hành vi tương tác trực tiếp như số lần nhắn tin, bình luận, hoặc tương tác nội dung.

Familiarity không chỉ cho thấy mối quan hệ cá nhân mà còn là một chỉ báo tiềm năng cho khả năng hợp tác và tin tưởng giữa các thành viên trong mạng. Người dùng có mức độ quen biết cao với nhiều người khác thường được đánh giá là đáng tin cậy hơn nhờ tính minh bạch và sự hiện diện liên tục trong hệ thống.

Về mặt định lượng, Familiarity có thể được tính theo trung bình số lượng tương tác theo thời gian, giúp phản ánh tần suất và tính ổn định của mối quan hệ, từ đó phục vụ làm đầu vào hiệu quả cho các mô hình học sâu trong dự đoán độ tin cậy.

Định nghĩa 2.1: Đặt $I_i = \{ \text{Tất cả người dùng từ } u_j \text{ có tương tác của } u_i \text{ đến } u_j \}$. Mức độ quen biết của hai người dùng u_i và u_j được định nghĩa như sau:

$$\text{familiarity}(i, j) = \frac{\|I_i \cap I_j\|}{\|I_i \cup I_j\|} \quad (2.1)$$

2.2.3 Mức độ phản hồi (*Responds*)

Mức độ phản hồi (Response) Trong môi trường mạng xã hội, mức độ phản hồi được hiểu là tốc độ và chất lượng người dùng phản hồi trước thông tin sai lệch, ví dụ như việc chỉnh sửa hoặc bác bỏ tin giả. Đặc trưng này đóng vai trò quan trọng trong việc xác định độ tin cậy, vì một nút có khả năng phản hồi tốt thường được xem là đáng tin cậy hơn do có thể xử lý tình huống nhanh chóng và chính xác. Ví dụ, trên Twitter, một người dùng nhanh chóng đính chính thông tin sai lệch về một sự kiện có thể nâng cao độ tin cậy của họ trong cộng đồng.

Định nghĩa 2.2: Giả sử $I_{i \leftarrow j}^{\text{resp}}$ là một tập tất cả các phản hồi từ u_j tới người dùng u_i . Trong đó, một phản hồi được tính là một trả lời từ u_j khi nhận được một thông điệp được gửi đi từ u_i . Khi đó mức độ phản hồi của người dùng u_j tới người dùng u_i , ký hiệu $\text{respond}(i, j)$ được định nghĩa như sau:

$$\text{response}(i, j) = \frac{\|I_{i \leftarrow j}^{\text{resp}}\|}{\|U_k I_{k \leftarrow j}^{\text{resp}}\|} \quad (2.2)$$

2.2.4 Mức độ tương tác (Dispatching)

Như đã trình bày ở trên, phương pháp xác định mức độ tin cậy dựa trên tần suất tương tác giữa những người dùng trong mạng phức tạp. Đề án đưa ra công thức tính toán mức độ tương tác giữa người dùng như sau:

$$\text{dispathch}(i, j) = \frac{\|I_{i \rightarrow j}\|}{\sum_{k=1}^n \|I_{i \rightarrow k}\|} \quad (2.3)$$

Nơi $\|I_{i \rightarrow k}\|$ là số lần ảnh hưởng từ u_i đến mỗi $u_k \in V$.

2.2.5 Mức độ ảnh hưởng (Influence)

Ảnh hưởng của một nút trong mạng phức tạp được xác định bởi mức độ mà nó tương tác với các nút khác trong mạng. Mục đích của nghiên cứu về ảnh hưởng là xác định các nút quan trọng trong việc lan truyền thông tin, và chủ đề này đã trở thành một hướng nghiên cứu tích cực trong những năm gần đây.

Định nghĩa 2.3: Hàm ảnh hưởng $\text{inf}: V \rightarrow [0,1]$ được định nghĩa là một hàm thỏa mãn các điều kiện sau:

(i) Ảnh hưởng của bất kỳ nút nào đều không âm:

$$\text{inf}(i) \geq 0 \forall i \in V$$

(ii) Giá trị ảnh hưởng cần được chuẩn hóa về khoảng $[0,1]: 0 \leq \text{inf}(i) \leq 1 \forall i \in V$

2.2.6 Chỉ số trung tâm

Số lượng chỉ số trung tâm là một trong những công cụ trực quan nhất để xác định nút quan trọng trong mạng. Trung tâm nút phản ánh mức độ lan truyền từ một nút trung tâm ra các nút khác và sức ảnh hưởng của nút đó trong mạng.

Đề án sẽ sử dụng hai chỉ số chính:

Định nghĩa 2.4: Mức độ trung tâm $\text{degree}(i)$ của một đỉnh u_i tỷ lệ thuận với mức độ (hoặc số lượng tiếp cận được kết nối trực tiếp) của đỉnh đó:

$$\text{degree}(i) = \frac{1}{n-1} \sum_{j=1}^n \alpha_{ij} \quad (2.4)$$

Trong đó: $\alpha_{ij} = 1$ nếu tồn tại u_i và u_j , ngược lại $\alpha_{ij} = 0$

Định nghĩa 2.5: Closeness centrality của một nút u_i là nghịch đảo của khoảng cách đường đi ngắn nhất trung bình từ tất cả $n - 1$ nút có thể tiếp cận được đến nút u_i .

$$closeness(i) = \frac{1}{p-1} \sum_{k=1}^{p-1} d(i, k) \quad (2.5)$$

Trong đó: $d(i, k)$ là khoảng cách ngắn nhất giữa u_i và u_k , và $p - 1$ là số lượng nút có thể tiếp cận từ u_i .

Kết hợp các yếu tố trên, mức độ tin cậy giữa hai nút u_i (trustor) và u_j (trustee) được tính như sau:

Mệnh đề 1: Hàm ảnh hưởng được định nghĩa là:

$$influence(i) = w_1 \times closeness(i) + w_2 \times degree(i) \quad (2.6)$$

Trong đó: $w_1 + w_2 = 1$, $w_1, w_2 \geq 0$

Hàm kết hợp kinh nghiệm tương tác và ảnh hưởng

Độ tin cậy dựa trên trải nghiệm tương tác của người dùng u_i đối với người dùng u_j , ký hiệu là $trust^{exp}(i, j)$, được tính bằng công thức:

$$trust^{exp}(i, j) = \omega_1 \cdot famili(i, j) + \omega_2 \cdot response(i, j) + \omega_3 \cdot dispatch(i, j) \quad (2.7)$$

Với các hệ số: $\omega_1, \omega_2, \omega_3 \geq 0$

$$\omega_1, \omega_2, \omega_3 = 1$$

Định nghĩa 2.6: Cho một người đánh giá độ tin cậy u_i và một người được đánh giá u_j , độ tin cậy tổng hợp từ trải nghiệm và ảnh hưởng là:

$$trust^{expInf}(i, j) = \omega_1 \cdot trust^{exp}(i, j) + \omega_2 \cdot influence(j) \quad (2.8)$$

Trong đó: $\omega_1, \omega_2 \geq 0$, $\omega_1 + \omega_2 = 1$

2.3 Tìm hiểu về mô hình mạng nơ ron nhân tạo CNN và LSTM

Trong bối cảnh dữ liệu ngày càng lớn và phức tạp, các mô hình học sâu đã chứng minh hiệu quả vượt trội trong việc xử lý, phân tích và dự đoán các mối quan hệ phi tuyến tính trong mạng. Trong mục này, hai kiến trúc mạng nơ-ron tiêu biểu là **CNN** và **LSTM** sẽ được trình bày như là nền tảng quan trọng cho việc xây dựng mô hình đánh giá độ tin cậy, với khả năng khai thác cả **đặc trưng không gian** và **phụ thuộc thời gian** trong dữ liệu tương tác giữa các thực thể trong mạng.

- **Mạng nơ-ron tích chập (CNN)**

CNN (Convolutional Neural Network) là một kiến trúc học sâu nổi bật với khả năng tự động trích xuất đặc trưng từ dữ liệu đầu vào thông qua các lớp tích chập (convolutional layers). Trong bài toán đánh giá độ tin cậy, CNN được sử dụng để xử lý **ma trận đặc trưng giữa các cặp người dùng**, từ đó phát hiện các mẫu tương tác tiềm ẩn có liên quan đến mức độ tin tưởng.

Ưu điểm chính của CNN là khả năng **giảm số lượng tham số** nhờ sử dụng các bộ lọc chia sẻ trọng số, đồng thời **phát hiện đặc trưng không gian** hiệu quả mà không cần đặc trưng thủ công. Điều này rất hữu ích khi xử lý dữ liệu mạng dạng ma trận, nơi các mối quan hệ giữa các thành phần (đỉnh – đỉnh) có thể được biểu diễn trực quan.

- **Mạng bộ nhớ dài ngắn hạn (LSTM)**

LSTM (Long Short-Term Memory) là một biến thể mạnh mẽ của RNN, được thiết kế để khắc phục hiện tượng mất thông tin trong chuỗi dài bằng cách sử dụng các cổng điều khiển (gồm cổng vào, cổng quên và cổng đầu ra). Mô hình này rất phù hợp để xử lý dữ liệu tuần tự, chẳng hạn như chuỗi tương tác theo thời gian giữa các người dùng.

Trong bối cảnh đánh giá độ tin cậy, LSTM cho phép học và ghi nhớ ảnh hưởng của các tương tác trong quá khứ, phản ánh sự thay đổi niềm tin theo thời gian. LSTM cũng đặc biệt hữu ích khi độ tin cậy không phải là giá trị tĩnh, mà là một quá trình tiến hóa dựa trên hành vi lịch sử.

Việc sử dụng LSTM giúp mô hình học được xu hướng tin cậy dài hạn, phát hiện các mẫu tăng – giảm niềm tin, từ đó cải thiện độ chính xác trong việc dự đoán độ tin cậy tại các thời điểm tương lai.

2.4 Mô hình CNN kết hợp với LSTM

Các nghiên cứu gần đây đã chứng minh hiệu quả vượt trội của mô hình kết hợp CNN và LSTM trong các bài toán xử lý dữ liệu chuỗi phức tạp. Ví dụ, Yoon et al. (2019) áp dụng mô hình CNN-LSTM vào dự báo chuỗi thời gian y tế và cải thiện độ chính xác so với việc sử dụng riêng lẻ CNN hoặc LSTM. Tương tự, Li et al. (2020) cho thấy mô hình này giúp khai thác đồng thời đặc trưng không gian và thông tin thời gian, nâng cao dự báo các chỉ số mạng như trust và influence.

Một nghiên cứu điển hình là bài báo “A Hybrid CNN-LSTM Model for High Resolution Melting Curve Analysis” (2019), trong đó mô hình CNN-LSTM đã được sử dụng để phân tích đường

cong nóng chảy (HRM) với độ phân giải cao. Kết quả cho thấy mô hình này đạt độ chính xác lên đến 96%, vượt trội hơn các phương pháp như CNN thuần túy và SVM.

Mô hình CNN-LSTM kết hợp ưu điểm của CNN trong việc trích xuất đặc trưng cục bộ và khả năng xử lý các mối quan hệ phụ thuộc thời gian dài hạn của LSTM, từ đó cung cấp công cụ mạnh mẽ cho các bài toán dữ liệu phức tạp như phân tích chuỗi thời gian và dự báo.

2.5 Kết luận chương 2

Chương 2 đã trình bày tổng quan lý thuyết liên quan đến bài toán suy diễn độ tin cậy trong mạng phức tạp. Từ cơ sở khái niệm về tin cậy đến các phân loại và đặc trưng của mạng xã hội, chương này cho thấy tính đa chiều, phi tuyến và có yếu tố thời gian trong hành vi tương tác giữa các cá nhân, điều này khiến việc xây dựng mô hình tính toán độ tin cậy trở nên thách thức nhưng cũng giàu tiềm năng.

Các mô hình suy diễn độ tin cậy truyền thống chủ yếu dựa vào cấu trúc mạng tĩnh hoặc các chỉ số mạng chuẩn hóa. Tuy nhiên, sự phát triển của học sâu trong xử lý chuỗi thời gian và mạng phức tạp đã mở ra một hướng tiếp cận mới, nơi dữ liệu tương tác theo thời gian giữa các cặp nút có thể được chuyển hóa thành chuỗi đặc trưng để đưa vào các kiến trúc mô hình mạnh mẽ như CNN, LSTM và đặc biệt là mô hình kết hợp CNN+LSTM.

Đáng chú ý, việc đưa vào các yếu tố như số bước thời gian, tần suất tương tác và chỉ số ảnh hưởng thông qua các phép biến đổi logarit đã góp phần làm nổi bật vai trò của các tín hiệu yếu trong dữ liệu, đồng thời giảm thiểu ảnh hưởng của các ngoại lệ. Đây là bước tiền đề quan trọng cho chương 3, nơi các mô hình học sâu sẽ được huấn luyện và đánh giá hiệu suất thực tế

trên các bộ dữ liệu có yếu tố thời gian. Các mô hình này không chỉ giúp dự đoán độ tin cậy theo hướng hồi quy và nhị phân, mà còn thể hiện tiềm năng mở rộng trong các hệ thống khuyến nghị và phân tích mạng xã hội động.

CHƯƠNG 3: THỬ NGHIỆM VÀ ĐÁNH GIÁ

3.1 Tổng quan về thử nghiệm

Chương này đánh giá hiệu quả của hai mô hình học sâu, CNN và LSTM, trong việc tính toán độ tin cậy của mạng phức tạp, tập trung vào xử lý đặc trưng không gian và thời gian. Thử nghiệm trên nhiều tập dữ liệu kiểm tra tính tổng quát, đặc biệt trong môi trường thay đổi và dữ liệu không hoàn chỉnh. Các chỉ số MSE, MAE, F1,... được sử dụng để so sánh độ chính xác và khả năng xử lý yếu tố không chắc chắn của mô hình. Nghiên cứu nhằm chứng minh các mô hình học sâu cải thiện độ chính xác, đồng thời tối ưu hóa tham số để giảm yêu cầu dữ liệu và tài nguyên tính toán.

3.2 Kịch bản thực nghiệm

Trong quá trình nghiên cứu, đề án triển khai hai hướng tiếp cận khác nhau để dự đoán độ tin cậy giữa các cặp nút trong mạng: (1) phân loại nhị phân (classification) và (2) hồi quy liên tục (regression). Việc áp dụng cả hai hướng không chỉ nhằm tìm ra mô hình có hiệu quả cao hơn, mà còn giúp đánh giá tính phù hợp của từng phương pháp với bản chất dữ liệu độ tin cậy trong mạng phức tạp.

a) Mô hình phân loại (*Binary Classification*)

Ở hướng tiếp cận này, giá trị độ tin cậy được rút gọn về hai mức độ: *cao* và *thấp*, thông qua một ngưỡng xác định trước. Nhãn đầu ra được mã hóa nhị phân (0 hoặc 1), và mô hình được xây dựng trên kiến trúc CNN+LSTM đa đầu vào, gồm:

- Nhánh xử lý chuỗi tuần tự: gồm các đặc trưng *dispatch*, *familiarity*, *response*, đi qua các lớp `Conv1D`, `BatchNormalization`, `LeakyReLU`, `Dropout`, sau đó đưa vào lớp `Bidirectional LSTM`.
- Nhánh xử lý đặc trưng phi tuần tự: *influence*, đi qua mạng `conv dense`.
- Hai nhánh được hợp nhất ở tầng `Concatenate`, rồi đi qua các lớp `Dense` và `Dropout`.
- Tầng đầu ra dùng hàm kích hoạt `sigmoid`, hàm mất mát là `binary_crossentropy`

b) Mô hình hồi quy (*Continuous Regression*)

Trong hướng tiếp cận thứ hai, giá trị độ tin cậy không còn bị phân loại, mà được giữ nguyên dưới dạng liên tục từ $[0 \rightarrow 1]$. Điều này cho phép mô hình học được sự khác biệt mịn giữa các mức độ tin cậy, phản ánh bản chất liên tục của quan hệ mạng.

- Cấu trúc mạng sử dụng kiến trúc tương tự mô hình classification ở trên, tuy nhiên thay đổi hàm mất mát sang `mean_squared_error`, vẫn giữ hàm kích hoạt `sigmoid` để chuẩn hóa đầu ra.
- Trong một số thử nghiệm, mô hình chỉ sử dụng CNN hoặc chỉ LSTM được triển khai riêng rẽ để so sánh hiệu suất.

Việc xây dựng song song hai hướng tiếp cận giúp đánh giá toàn diện khả năng học của mô hình trong cả bài toán ranh giới và bài toán giá trị liên tục.

3.3 Thử nghiệm và đánh giá

Danh sách tập dữ liệu đề án sử dụng Stanford Large Network Dataset Collection “<https://snap.stanford.edu/data>”:

Trong nghiên cứu này, đề án sử dụng 5 tập dữ liệu mạng phức tạp, chủ yếu là các mạng xã hội, mạng giao tiếp có yếu tố thời gian. Các tập dữ liệu được lựa chọn nhằm kiểm tra tính hiệu quả và khả năng tổng quát hóa của mô hình học sâu trong nhiều bối cảnh mạng khác nhau.

Tên tập dữ liệu	Dạng dữ liệu	Nodes	Edges	Temporal Edges	Static Edges
email-Eu-core-temporal	Mạng email có thời gian	986		332,334	24,929
CollegeMsg	Mạng nhắn tin có thời gian	1,899		20,296	59,835
sx-mathoverflow	Mạng hỏi đáp StackExchange	24,818		506,550	239,978
Email Enron	Mạng email lớn	36,692	183,831		
Email – EuAll	Mạng email tổng hợp có thời gian	265,214	420,045		

Mục đích thử nghiệm của đề án là đánh giá hiệu quả của các mô hình học sâu (CNN, LSTM, và CNN+LSTM) trong dự đoán độ tin cậy trong mạng xã hội, tập trung vào xử lý dữ liệu thời gian và đặc trưng tương tác mạng. Các mô hình được đánh giá qua chỉ số MSE, MAE, Accuracy, Recall, F1-Score, nhằm chứng minh rằng sự kết hợp CNN+LSTM có thể dự đoán độ tin cậy hiệu quả trong mạng phức tạp.

3.3.1 Dự đoán tin cậy theo hướng phân loại nhị phân (Binary Classification)

Nghiên cứu kiểm nghiệm hiệu quả mô hình phân loại nhị phân trong dự đoán độ tin cậy giữa các đối tượng trong mạng phức tạp, sử dụng các tập dữ liệu **Email Enron** và **Email – EuAll**. Các tập dữ liệu này phản ánh các tương tác trong môi trường tổ chức khác nhau và giúp gán nhãn độ tin cậy thành hai nhãn: **0** (tin cậy thấp) và **1** (tin cậy cao), từ đó đánh giá mô hình trong các mạng xã hội có cấu trúc tương đồng.

- **Chuẩn bị đặc trưng và nhãn**

Dữ liệu được đọc từ tập dữ liệu dưới dạng edgelist, biểu diễn các kết nối giữa các thực thể.

=== 10 dòng đầu tiên của tập dữ liệu gốc ===

```
0      1
0      4
0      5
0      8
0     11
0     20
```

=====

=== 10 dòng đầu tiên của DataFrame sau xử lý ===

	dispatch	familiarity	response	influence	trust
0	0.374540	0.950714	0.731994	0.003616	0.208256
1	0.598658	0.156019	0.155995	0.000826	0.091645
2	0.058084	0.866176	0.601115	0.001109	0.153313
3	0.708073	0.020584	0.969910	0.000166	0.169973
4	0.832443	0.212339	0.181825	0.000449	0.122975
5	0.183405	0.304242	0.524756	0.001218	0.102093
6	0.431945	0.291229	0.611853	0.001644	0.134653
7	0.139494	0.292145	0.366362	0.001093	0.080565
8	0.456070	0.785176	0.199674	0.001569	0.145190
9	0.514234	0.592415	0.046450	0.000373	0.115571

=====

- Mỗi cạnh trong mạng được gán 3 thuộc tính ngẫu nhiên:
 - o *dispatch*: cường độ kết nối
 - o *familiarity*: độ quen thuộc
 - o *response*: khả năng phản hồi
- Mỗi nút được tính *influence* dựa trên độ trung tâm (degree centrality).
- Từ các đặc trưng này, giá trị độ tin cậy (*trust*) được tính theo công thức 2.7.
- Tập giá trị (*trust*) được phân thành hai lớp dựa trên trung vị hoặc trung bình.
- Mỗi mẫu dữ liệu gồm: [*dispatch, familiarity, response, influence*] → nhãn (0/1).

- **Kiến trúc mô hình huấn luyện**

Mô hình sử dụng ba kiến trúc học sâu được xây dựng và thử nghiệm: CNN, LSTM và CNN+LSTM kết hợp. Các mô hình đều sử dụng dữ liệu chuỗi đặc trưng đầu vào với độ dài cố định, trong đó mỗi bước thời gian chứa 4 đặc trưng: số lượng tương tác, mức độ phản hồi, chỉ số lan truyền và ảnh hưởng cục bộ. Tất cả các mô hình đều sử dụng hàm mất mát `binary_crossentropy` với hàm kích hoạt sigmoid tại đầu ra nhằm trả về xác suất thuộc lớp “tin cậy cao”.

- **Mô hình CNN:**

Mô hình CNN được thiết kế để khai thác hiệu quả các mẫu cục bộ trong chuỗi đặc trưng thời gian. Cấu trúc gồm:

- o **Conv1D (32 filters, kernel size = 3)**: trích xuất đặc trưng từ 3 bước thời gian liên tiếp.
- o **GlobalMaxPooling1D**: giữ lại đặc trưng nổi bật nhất trên toàn chuỗi.

- **Dense (32 units, ReLU)**: học quan hệ phi tuyến.
- **Dense (1 unit, Sigmoid)**: xác suất phân lớp.

– Mô hình LSTM

Mô hình LSTM có khả năng ghi nhớ mối quan hệ dài hạn trong chuỗi. Cấu trúc gồm:

- **LSTM (64 units)**: xử lý thứ tự chuỗi và học phụ thuộc dài hạn.
- **Dense (32 units, ReLU)**: tầng ẩn phi tuyến.
- **Dense (1 unit, Sigmoid)**: đầu ra phân lớp.

– Mô hình kết hợp CNN + LSTM

Mô hình lai kết hợp hai ưu thế: trích xuất đặc trưng cục bộ bằng CNN và học phụ thuộc thời gian bằng LSTM:

- **Conv1D (32 filters) + BN + Dropout(0.3)**: lọc nhiễu và ổn định huấn luyện.
- **Conv1D (64 filters) + BN**: tiếp tục tăng chiều sâu trích xuất.
- **Flatten → Reshape**: chuyển đặc trưng thành chuỗi để đưa vào LSTM.
- **BiLSTM (2x32 units)**: học phụ thuộc hai chiều trong chuỗi thời gian.
- **Dense (32 units, ReLU) → Dense (1 unit, Sigmoid)**.

• Thuật toán tối ưu (Optimizer)

Mô hình sử dụng **Adam Optimizer**, một thuật toán hiện đại kết hợp ưu điểm của AdaGrad và RMSProp, giúp tự động điều chỉnh tốc độ học dựa trên gradient từng chiều, đảm bảo hội tụ nhanh và ổn định thích hợp với dữ liệu phi tuyến và huấn luyện trên tập lớn.

• Tỷ lệ tách tập validation

Trong quá trình huấn luyện, **20% dữ liệu huấn luyện được tách ra làm tập kiểm định (validation)**. Tập này để theo dõi hiệu suất (val_loss), điều chỉnh mô hình và kích hoạt cơ chế dừng sớm khi cần.

• Cơ chế dừng sớm (Early Stopping)

Trong quá trình huấn luyện mô hình học sâu, việc mô hình tiếp tục cải thiện hiệu suất trên tập huấn luyện nhưng không còn cải thiện trên tập kiểm định (validation) là một dấu hiệu của quá khớp (overfitting). Để kiểm soát hiện tượng này, nghiên cứu sử dụng kỹ thuật dừng sớm (EarlyStopping) như một cơ chế điều khiển huấn luyện thông minh.

Cụ thể, EarlyStopping được cấu hình theo các tham số sau:

- **Tiêu chí theo dõi (monitor)**: `val_loss` – giá trị hàm mất mát trên tập kiểm định.
- **Ngưỡng kiên nhẫn (patience)**: 10 – nếu sau 10 vòng lặp (epoch) liên tiếp mà `val_loss` không giảm thì dừng huấn luyện.
- **Khôi phục trọng số tốt nhất (restore_best_weights)**: `True` – đảm bảo mô hình sử dụng trọng số tại thời điểm đạt hiệu quả tốt nhất.

- **Hiển thị chi tiết (verbose):** 1 – xuất thông báo khi kích hoạt cơ chế dừng.

Cơ chế này đặc biệt phát huy hiệu quả ở các mô hình có số tham số lớn như CNN+LSTM, nơi mà việc huấn luyện quá lâu có thể làm giảm khả năng tổng quát hóa của mô hình.

- **Đánh giá mô hình:**

Sau huấn luyện, mô hình được đánh giá bằng các chỉ số:

- **Phân loại:** Accuracy, Precision, Recall, F1-score
- **Hồi quy:** Mean Squared Error (MSE), Mean Absolute Error (MAE), Mean Absolute Percentage Error (MAPE)

3.3.2 hình hồi quy (Continuous Regression)

Trong nghiên cứu này, các bộ dữ liệu: **email-Eu-core-temporal**, **CollegeMsg** và **Reddit Hyperlink Network** từ tập Stanford được lựa chọn nhằm phục vụ bài toán phân tích độ tin cậy trong mạng động. Các dữ liệu này đều có đặc điểm nổi bật như: mang yếu tố thời gian (temporal), có hướng tương tác (directed), quy mô lớn, và phản ánh đa dạng các loại mạng xã hội thực tế. Nhờ đó, chúng đáp ứng tốt yêu cầu của các mô hình học sâu như LSTM, CNN+LSTM và cho phép đánh giá hiệu quả suy diễn độ tin cậy trong các bối cảnh khác nhau một cách khách quan và tái lập.

- **Bước 1: Cấu trúc dữ liệu đầu vào**

- Dữ liệu gốc có dạng `node1 - node2 - timestamp` được gom nhóm thành chuỗi tương tác giữa từng cặp nút theo từng ngày.
- Sau khi gom nhóm theo ngày và cặp (`node1`, `node2`) → có số lượng tương tác (interaction count) theo từng ngày.

 10 dòng đầu tiên của DataFrame gốc (df):

	node1	node2	timestamp	datetime	date
0	582	364	0	2022-01-01 00:00:00	2022-01-01
1	168	472	2797	2022-01-01 00:46:37	2022-01-01
2	168	912	3304	2022-01-01 00:55:04	2022-01-01
3	2	790	4523	2022-01-01 01:15:23	2022-01-01
4	2	322	7926	2022-01-01 02:12:06	2022-01-01
5	2	790	8061	2022-01-01 02:14:21	2022-01-01
6	42	402	19403	2022-01-01 05:23:23	2022-01-01
7	870	337	19560	2022-01-01 05:26:00	2022-01-01
8	663	362	21077	2022-01-01 05:51:17	2022-01-01
9	663	410	21280	2022-01-01 05:54:40	2022-01-01

📄 10 dòng đầu tiên của DataFrame đã lọc (df_filtered):

	node1	node2	date	interaction_count
0	0	6	2022-04-05	1
1	0	6	2022-04-06	1
2	0	6	2022-05-10	1
3	0	6	2022-05-17	1
4	0	6	2022-05-24	1
5	0	6	2022-05-30	1
6	0	6	2022-05-31	1
7	0	6	2022-06-06	1
8	0	6	2022-06-07	2
9	0	6	2022-06-08	1

- Các chuỗi có đủ số lượng bước thời gian (ví dụ ≥ 10 ngày) mới được giữ lại để đảm bảo chất lượng huấn luyện.

- *Bước 2: Tạo chuỗi tương tác*

Nghiên cứu trích xuất chuỗi tương tác giữa các cặp nút (u, v) trong mạng động, ánh xạ thành chuỗi thời gian thể hiện số lần tương tác, nhằm mô hình hóa hành vi và độ tin cậy. Chuỗi được chuẩn hóa độ dài bằng zero-padding và chuyển thành vector đặc trưng (tần suất, mức lan truyền, ảnh hưởng), cải tiến so với các nghiên cứu trước chỉ dùng tần suất thô, hỗ trợ hiệu quả cho hồi quy và phân tích mạng xã hội.

- *Bước 3: Trích xuất đặc trưng*

Quá trình trích xuất đặc trưng đóng vai trò quan trọng trong việc biểu diễn thông tin tương tác giữa các cặp nút dưới dạng chuỗi thời gian có thể xử lý được bằng các mô hình học sâu. Dựa trên cơ sở lý thuyết và thực nghiệm từ các nghiên cứu trước đó, đề án lựa chọn bốn đặc trưng cơ bản nhằm phản ánh đa chiều hành vi tương tác trong mạng.

(1) Tần suất tương tác f_t : Đặc trưng này đại diện cho số lần tương tác giữa hai nút (u, v) tại thời điểm t . Đây là chỉ số cơ bản thường được sử dụng trong các mô hình mạng xã hội nhằm đo lường cường độ liên kết giữa các tác nhân. Giá trị này phản ánh trực tiếp xu hướng cộng tác hoặc giao tiếp của các nút trong quá khứ.

(2) Mức lan truyền chuẩn hóa d_t :

Để tránh hiện tượng các chuỗi có giá trị lớn gây mất cân bằng đầu vào, mức lan truyền được chuẩn hóa theo công thức:

$$d_t = \min \frac{f_t}{\theta}, (1.0) \quad (3.3)$$

Trong đó, θ là ngưỡng được xác định thực nghiệm phép chuẩn hóa ngưỡng này đã được áp dụng trong nhiều nghiên cứu mô hình hóa ảnh hưởng trong mạng phức tạp (Leskovec et al., 2007; Tang et al., 2010)[]. Việc đưa giá trị về trong khoảng $[0,1]$ giúp tăng tính ổn định cho quá trình huấn luyện.

(3) Chỉ số ảnh hưởng tương đối inf_t

$$inf_t = \frac{\log(f_t + 1)}{\log(f_{max} + 1)} \quad (3.4)$$

Chỉ số đánh giá mức độ ảnh hưởng trong mạng xã hội, kế thừa từ Gupta et al. (2014) và Zhang et al. (2018), sử dụng phép biến đổi logarit để giảm tác động của giá trị ngoại lệ và làm nổi bật thay đổi nhỏ trong tương tác, hỗ trợ mô hình học sâu khai thác đặc trưng phi tuyến. Công thức 2.3 biểu thị tỷ lệ ảnh hưởng tương đối ổn định nhưng không phản ánh thay đổi thời gian, không phù hợp với chuỗi ngắn hạn (CNN) hoặc dài hạn (LSTM), và gây mất cân bằng đặc trưng. Để khắc phục, công thức 3.3 mới được đề xuất, chuẩn hóa giá trị trong $[0,1]$, tăng cường ổn định, tránh ngoại lệ, và hỗ trợ tối ưu gradient cho mô hình CNN+LSTM, giúp hội tụ nhanh hơn.

(4) Chuẩn hóa độ dài chuỗi:

Do số lượng bước thời gian giữa các cặp nút có thể khác nhau, các chuỗi được chuẩn hóa về độ dài cố định T thông qua kỹ thuật *zero-padding*. Đây là thực hành phổ biến trong các mô hình học sâu xử lý dữ liệu tuần tự, giúp đảm bảo tính đồng nhất cho quá trình huấn luyện. Việc lựa chọn và xây dựng các đặc trưng nói trên không chỉ giúp mô hình học được mối quan hệ theo thời gian giữa các nút mà còn duy trì được tính khái quát và ổn định trong huấn luyện, đặc biệt trong bài toán tính toán độ tin cậy trên mạng phức tạp có yếu tố thời gian.

3.4 Kết quả thực nghiệm

Sau khi trải qua quá trình huấn luyện đề án thu được các kết quả như sau:

- Đối với mô hình dựa vào phân loại nhị phân đề án sẽ sử dụng hai bộ dữ liệu “*Email Enron*” và “*Email – EuAll*” ta được kết quả như sau (Bảng 3.1):

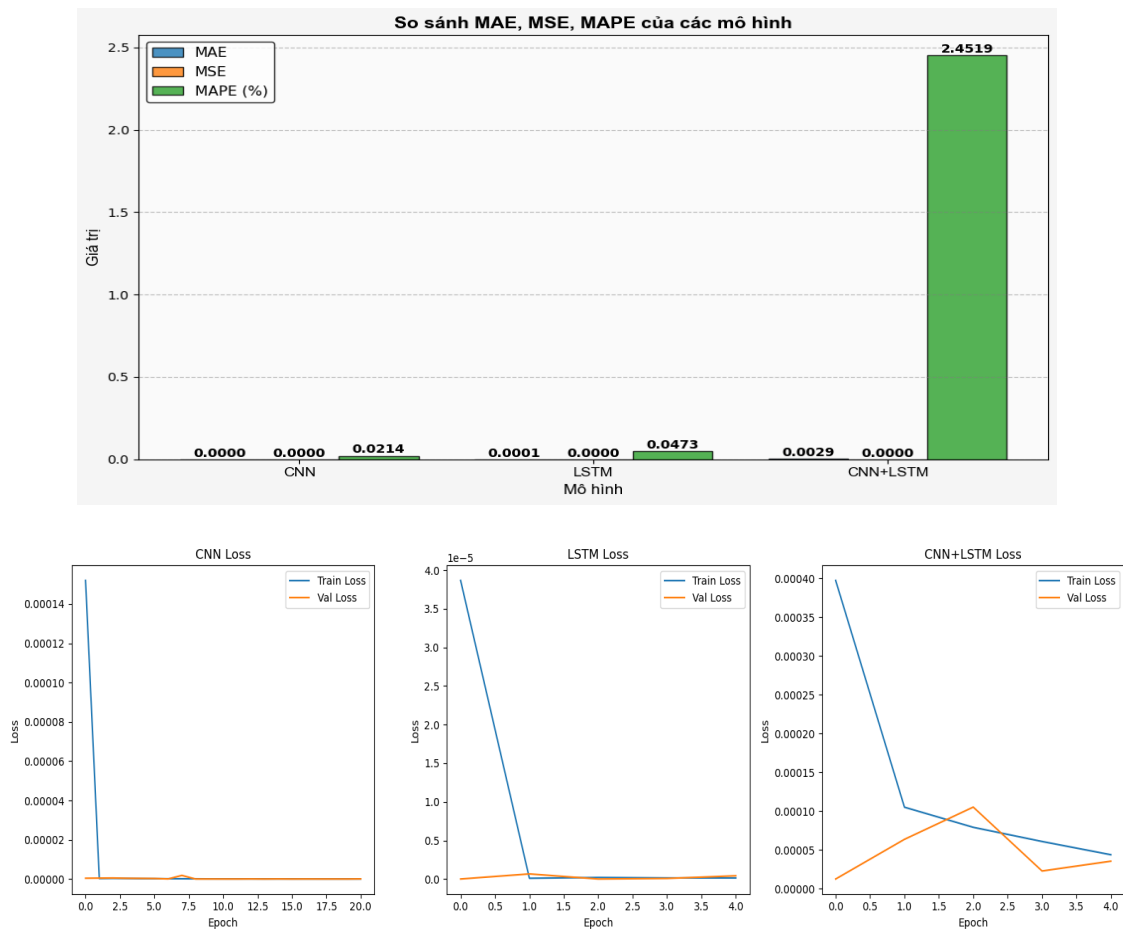
Tên bộ dữ liệu	CNN			LSTM			CNN+LSTM		
	MAE	MSE	MAPE	MAE	MSE	MAPE	MAE	MSE	MAPE
Email Enron	0.00000	0.000023	0.023250	0.00000	0.000198	0.174273	0.000070	0.006618	4.631706
Email – EuAll	0.00000	0.000023	0.021432	0.00000	0.000058	0.047296	0.000013	0.002852	2.451851

Tên bộ dữ liệu	CNN				LSTM				CNN+LSTM			
	A	P	R	F1	A	P	R	F1	A	P	R	F1
Email Enron	0.999973	1.000000	0.999946	0.999973	0.999157	0.999728	0.998585	0.999156	0.961487	1.000000	0.922918	0.959914
Email – EuAll	0.999918	1.000000	0.999836	0.999918	0.999562	0.999863	0.999262	0.999562	0.988703	0.984361	0.993191	0.988756

Ghi chú : A: Accuracy, P = Precision, R = Recall, F1 = F1-score

Bảng 3.1 Bảng so sánh độ đo của bộ dữ liệu “Email enron” và “Email-EuAll”

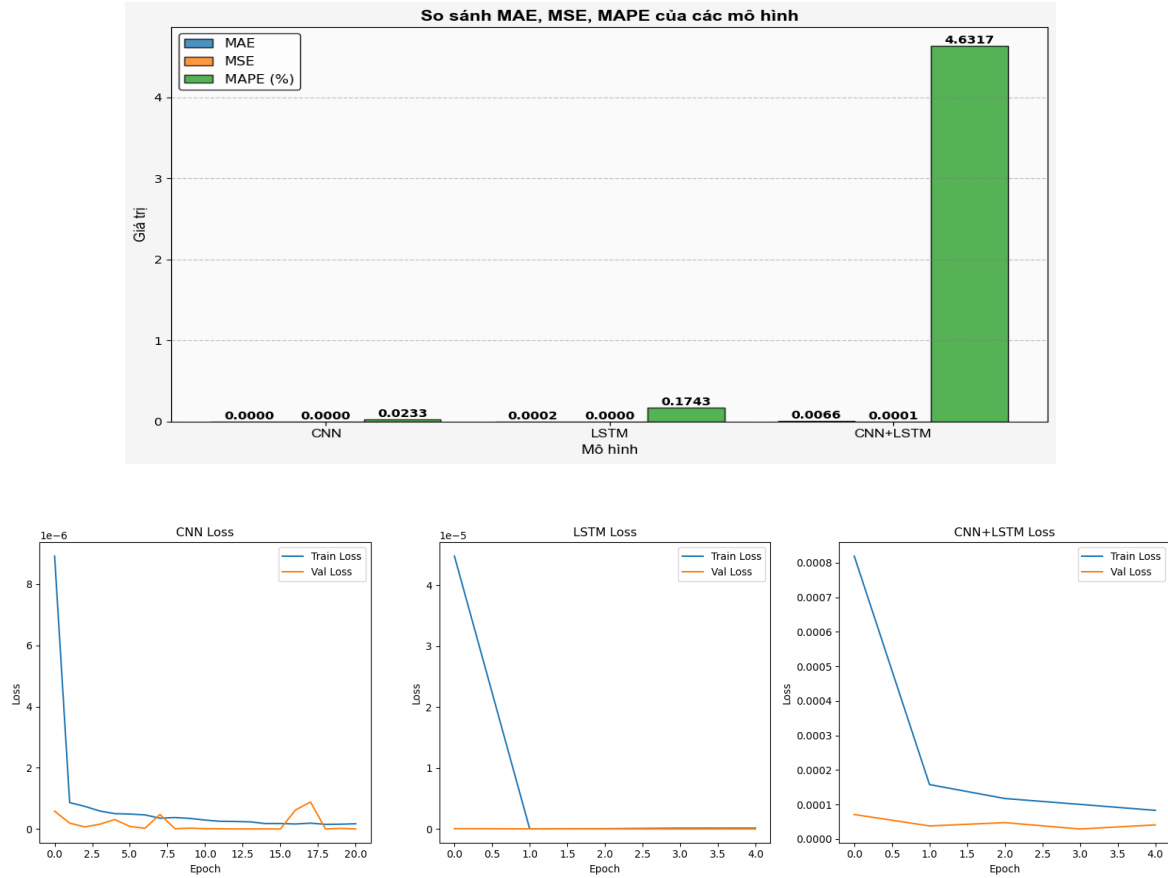
- Bộ dữ liệu Email-Euall



Hình 3.4 Biểu đồ MAE, MSE, MAPE và biểu đồ loss “Email-Euall”

Trên tập dữ liệu Email – EuAll, mô hình CNN vượt trội với $MAE = 0.0000$, $MSE = 0.0000$, $MAPE = 0.0214\%$, cho thấy độ chính xác và ổn định cao. LSTM có MAE tương đương nhưng MAPE cao hơn (0.0473%), thể hiện sai lệch tương đối lớn. Mô hình CNN+LSTM, dù phức tạp, đạt MAPE cao nhất (2.4519%) và biểu đồ loss dao động ở tập validation, cho thấy học không ổn định và có nguy cơ overfitting, do số tham số lớn và tập dữ liệu đơn giản, không đủ để khai thác mô hình lại.

- Bộ dữ liệu Email-Enron



Hình 3.5 Biểu đồ MAE, MSE, MAPE và biểu đồ loss “Email-Enron”

Trên tập dữ liệu Email–Enron, mô hình CNN đạt hiệu quả cao nhất với sai số thấp nhất (MAE, MSE, MAPE), vượt trội về độ chính xác và ổn định. LSTM có kết quả chấp nhận được nhưng kém ổn định hơn. Mô hình CNN+LSTM, dù MAE thấp, lại có MAPE cao nhất (4.63%) và biểu đồ loss dao động trên tập validation, cho thấy tổng quát kém và dễ overfitting. Do đó, CNN là mô hình phù hợp nhất cho dự đoán độ tin cậy trên Email–Enron.

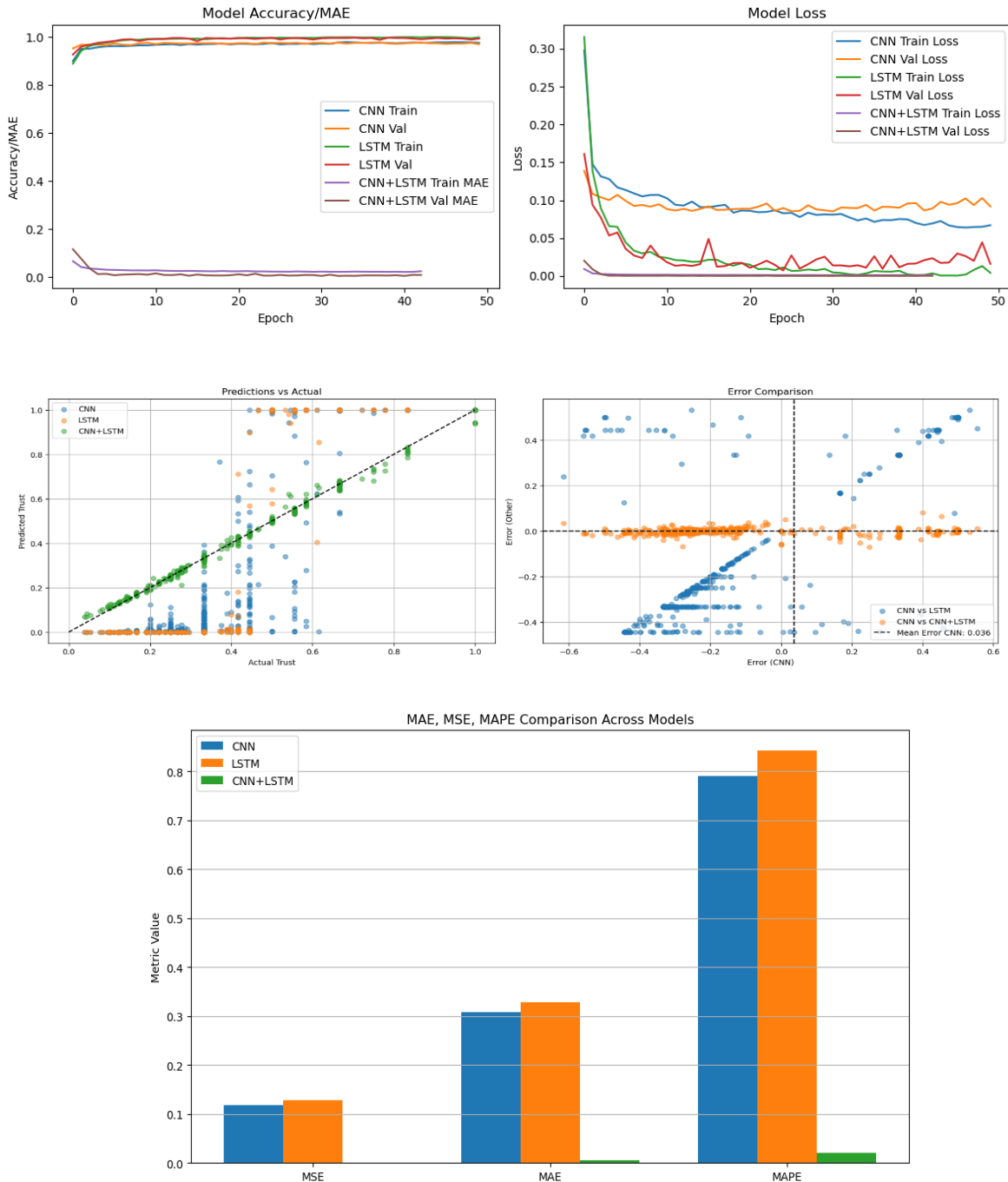
Để cải thiện CNN+LSTM (mục 3.2.2), đề án đề xuất tinh chỉnh kiến trúc để khai thác hiệu quả hơn chuỗi tương tác giữa các cặp nút, nâng cao khả năng học biểu diễn phi tuyến và giảm sai số trong dữ liệu thời gian phức tạp. Kết quả sau tinh chỉnh được ghi nhận trên ba tập dữ liệu: Email-Eu-core-temporal, Email-Enron, CollegeMsg (Bảng 3.2).

Tên dữ liệu	Mô hình	ĐỘ ĐO						
		MAE	MSE	MAPE	A	P	R	F1
email-Eu-core-temporal	CNN	0.3087	0.1173	0.7896	0.96721	0.98068	0.95208	0.96617
	LSTM	0.3278	0.1279	0.8414	0.99590	0.99377	0.99791	0.99584
	CNN+LSTM	0.0001	0.0064	0.0193	0.9375	1.0	0.87291	0.93214
sx-mathoverflow	CNN	0.3058	0.1162	0.7640	0.91463	0.96428	0.88043	0.88043
	LSTM	0.3360	0.1330	0.8498	0.98373	0.99264	0.97826	0.98540
	CNN+LSTM	0.0081	0.0002	0.0229	0.77235	1.0	0.59420	0.74545
CollegeMsg	CNN	0.2868	0.1224	1.1423	0.75	0.66666	0.83333	0.74074
	LSTM	0.2548	0.1158	0.9238	0.85714	0.83333	0.83333	0.83333
	CNN+LSTM	0.1985	0.0522	1.2711	0.75	1.0	0.41666	0.58823

Bảng 0.1 Bảng so sánh độ đo 3 bộ dữ liệu “Email-Eu-core-temporal”, “Sx-mathover”, “CollegeMsg”

Ghi chú : A: Accuracy, P = Precision, R = Recall, F1 = F1-score

- **Bộ dữ liệu email-Eu-core-temporal**



Hình 3.6 Tổng quan hiệu suất huấn luyện và đánh giá mô hình trên tập dữ liệu email-Eu-core-temporal

Trong thực nghiệm, ba mô hình học sâu **CNN**, **LSTM**, và **CNN+LSTM** đã được huấn luyện và đánh giá trên tập dữ liệu **email-Eu-core-temporal**. Kết quả cho thấy:

- **CNN+LSTM** đạt hiệu suất tốt nhất về hồi quy với chỉ số $MAE = 0.0062$, $MSE = 0.0001$ và $MAPE = 0.0205$, cho thấy khả năng học chuỗi thời gian và dự đoán độ tin cậy liên tục chính xác.

- Trong phân loại, **LSTM** đạt hiệu suất cao hơn với $F1 = 0.9958$, nhưng **CNN+LSTM** vẫn mạnh mẽ nhờ kết hợp ưu điểm của cả hai mô hình, đặc biệt là trong bài toán hồi quy độ tin cậy.

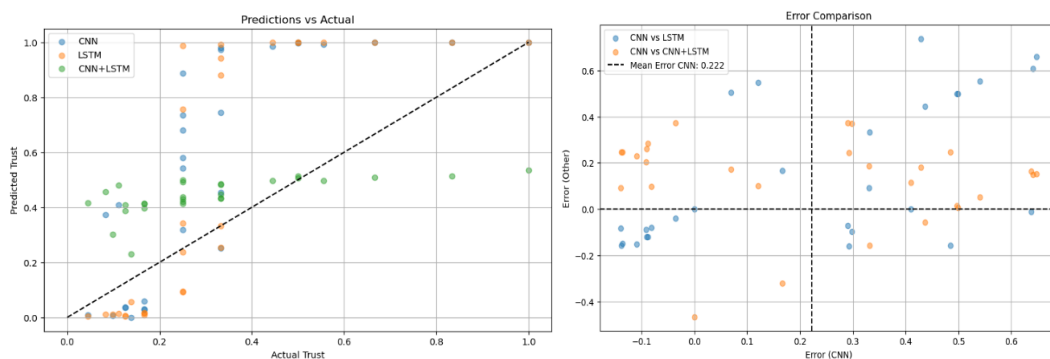
Tóm lại, **CNN+LSTM** phù hợp hơn với bài toán hồi quy, trong khi **LSTM** mạnh mẽ hơn trong phân loại.

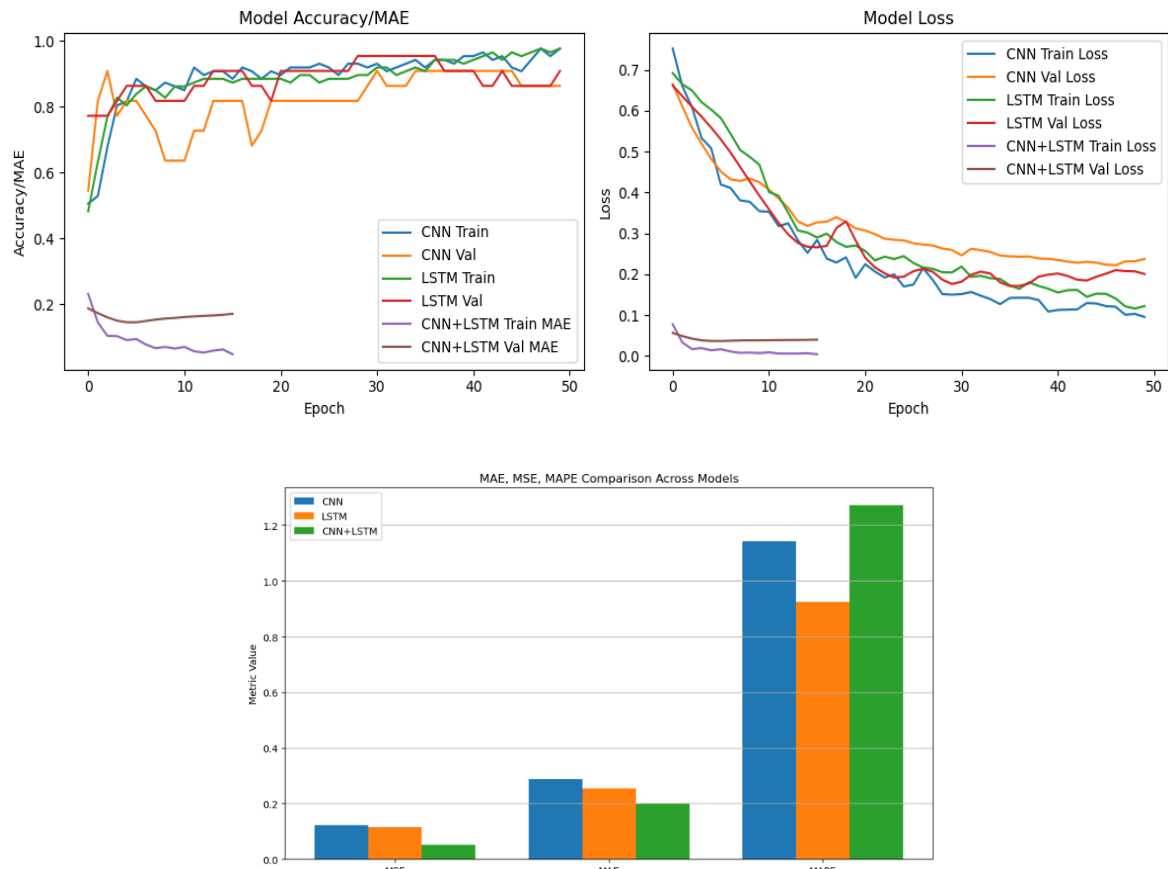
• Bộ dữ liệu CollegeMsg

Tập dữ liệu **CollegeMsg** phản ánh luồng tin nhắn giữa các sinh viên, với các tương tác ngẫu nhiên và bùng phát theo thời gian. Trong thực nghiệm:

- CNN+LSTM** đạt kết quả tốt nhất về MAE và MSE, cho thấy khả năng dự đoán độ tin cậy liên tục chính xác, nhưng có MAPE cao nhất, bị ảnh hưởng bởi ngoại lệ và sai lệch tỷ lệ ở dữ liệu nhỏ.
- LSTM** đạt hiệu suất phân loại cao nhất với $F1 = 0.8333$ và các chỉ số Precision, Recall cân bằng.
- CNN** có kết quả trung bình nhưng ổn định.

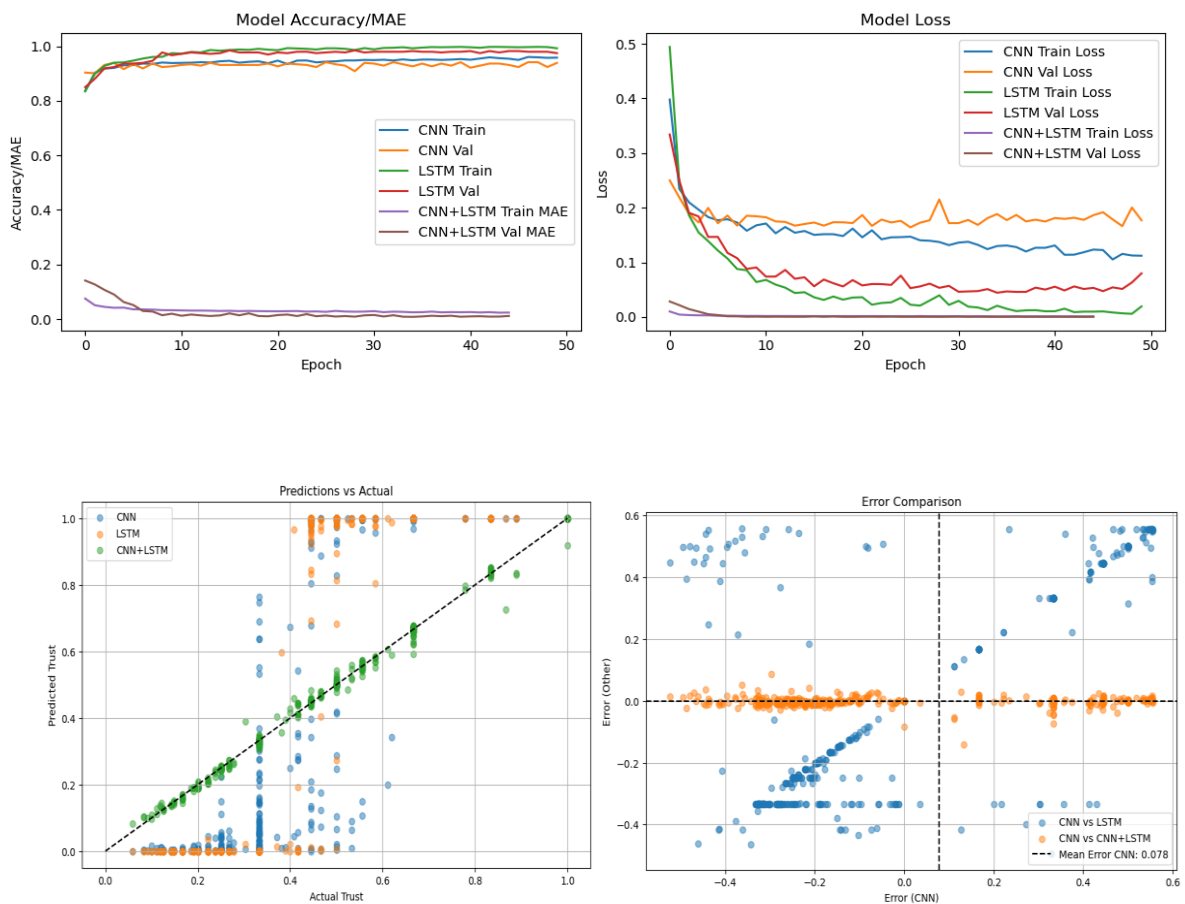
Biểu đồ phân tích cho thấy **CNN+LSTM** có độ lệch dự đoán cao khi Actual Trust > 0.6 và tập trung sai số ở các giá trị thấp, trong khi **LSTM** ổn định hơn với sai số thấp. **CNN+LSTM** hội tụ nhanh và duy trì độ chính xác huấn luyện cao, nhưng cần cải thiện xử lý ngoại lệ.

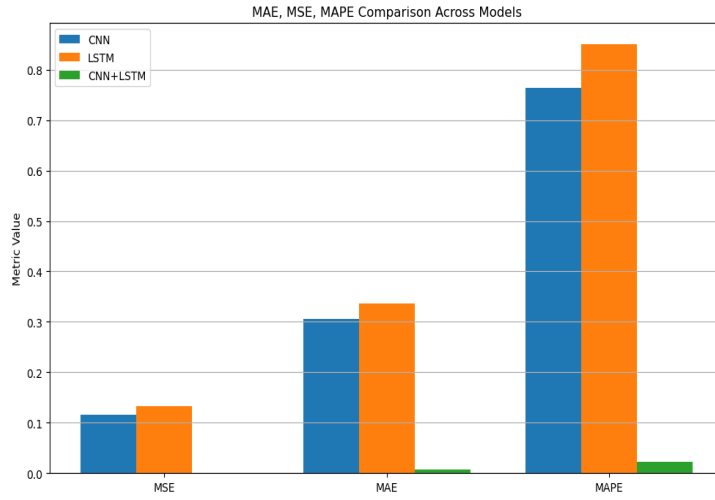




Hình 3.7 Tổng quan hiệu suất huấn luyện và đánh giá mô hình trên tập dữ liệu CollegeMsg

- Bộ dữ liệu sx-mathoverflow**





Hình 3.8 Tổng quan hiệu suất huấn luyện và đánh giá mô hình trên tập dữ liệu *sx-mathoverflow*

- CNN+LSTM đạt được MAE, MSE, MAPE thấp nhất, chứng tỏ khả năng học sâu hiệu quả và mô hình hóa tốt chuỗi tương tác.
- Tuy nhiên, về độ chính xác phân loại (Accuracy, Recall, F1), LSTM đơn thuần lại chiếm ưu thế.
- Điều này cho thấy rằng CNN+LSTM tối ưu rất tốt sai số liên tục (regression) nhưng gặp giới hạn trong phân loại nhị phân đối với tập này, có thể do ảnh hưởng của chuỗi dữ liệu ngắn và độ phân tán cao trong mạng.
- Biểu đồ thể hiện rằng các dự đoán của CNN+LSTM phân bố sát với giá trị thực tế, trong khi sai số trung bình của CNN cao hơn đáng kể. Tuy nhiên, xét theo độ chính xác phân loại (Accuracy, Recall, F1), LSTM vẫn chiếm ưu thế, cho thấy tính ổn định trong bài toán phân lớp.

Tổng thể, mô hình CNN+LSTM phù hợp cho các bài toán hồi quy liên tục trên mạng có chuỗi tương tác ngắn và phức tạp, như *sx-mathoverflow*, nhờ khả năng giảm sai số và học đặc trưng tốt hơn.

3.5 Kết luận thực nghiệm

Ba mô hình CNN, LSTM, và CNN+LSTM được thử nghiệm để dự đoán độ tin cậy từ dữ liệu tương tác thời gian trên các tập Email-Eu-core-temporal, *sx-mathoverflow*, và CollegeMsg. CNN+LSTM vượt trội về độ chính xác, với MSE = 0.0001, MAE = 0.0062, và Precision = 1.0 trên Email-Eu-core-temporal, nhờ khả năng học chuỗi phức tạp (mục 3.3.1). Trên Email-Enron, LSTM vượt CNN về MAE và F1-score, phù hợp hơn với dữ liệu thời gian, dù CNN+LSTM chưa được thử nghiệm trực tiếp.

Các kỹ thuật tinh chỉnh như zero-padding, biến đổi logarit, BatchNormalization, và Dropout cải thiện ổn định, giảm overfitting và sai số. CNN+LSTM là hướng đi hiệu quả cho suy diễn độ tin cậy trong mạng phức tạp, với tiềm năng ứng dụng trong hệ thống khuyến nghị, giám sát hành vi, và phân tích mạng xã hội động.

3.6 Kết luận chương 3

Chương này trình bày việc triển khai các mô hình học sâu bao gồm CNN, LSTM, và đặc biệt là CNN+LSTM để giải quyết bài toán suy diễn độ tin cậy trong mạng phức tạp. Kết quả thực nghiệm cho thấy CNN+LSTM đạt hiệu suất vượt trội cả về phân loại nhị phân và hồi quy giá trị độ tin cậy, với các chỉ số đánh giá như MAE, MSE, và MAPE được cải thiện đáng kể so với các mô hình đơn lẻ, đồng thời có sự ổn định cao qua các lần huấn luyện.

Các mô hình học sâu, đặc biệt là CNN+LSTM, đã chứng minh khả năng học đặc trưng mạnh mẽ, không bị giới hạn bởi cấu trúc đồ thị tĩnh hay yêu cầu tham số lan truyền. Nhờ tận dụng chuỗi tương tác theo thời gian và biểu diễn đặc trưng phi tuyến, mô hình học sâu cho kết quả chính xác hơn, phản ánh tốt hơn tính động và phức tạp của hành vi tin cậy trong mạng.

Tổng quan, chương này chứng minh tiềm năng ứng dụng của mô hình học sâu trong việc tính toán độ tin cậy, mở ra một hướng tiếp cận hiệu quả hơn trong bối cảnh dữ liệu mạng ngày càng lớn và giàu tính thời gian.

KẾT LUẬN VÀ HƯỚNG PHÁT TRIỂN

Đề án đã triển khai thành công các mô hình học sâu, đặc biệt là sự kết hợp giữa CNN và LSTM, để suy diễn độ tin cậy trong mạng phức tạp. Kết quả thực nghiệm cho thấy CNN+LSTM vượt trội trong cả bài toán hồi quy và phân loại nhị phân, cải thiện đáng kể các chỉ số MAE, MSE, và MAPE, đồng thời phản ánh tính động và phức tạp của hành vi tin cậy trong mạng.

Hạn chế và Hướng phát triển

Hạn chế:

Mặc dù đề án đã đạt được những kết quả thực nghiệm khả quan trong việc dự đoán độ tin cậy trên mạng phức tạp thông qua các mô hình học sâu như CNN, LSTM và CNN+LSTM, tuy nhiên vẫn tồn tại một số hạn chế nhất định. Trước hết, các đặc trưng đầu vào được xây dựng chủ yếu dựa trên chuỗi số lượng tương tác giữa các cặp nút theo thời gian, chưa khai thác sâu các đặc trưng cấu trúc của mạng như bậc nút, độ chồng cộng đồng hay ảnh hưởng lan truyền. Bên cạnh đó, phạm vi dữ liệu sử dụng trong nghiên cứu còn tương đối hạn chế, chủ yếu tập trung vào các mạng email với cấu trúc khép kín, do đó tính tổng quát của mô hình vẫn chưa được kiểm chứng rộng rãi trên các mạng xã hội hoặc hệ thống thực tế khác. Một điểm hạn chế đáng chú ý là khả năng giải thích của các mô hình học sâu còn thấp, gây khó khăn trong việc lý giải cụ thể các quyết định đầu ra của mô hình. Ngoài ra, quá trình tinh chỉnh siêu tham số của mô hình hiện vẫn mang tính thủ công, thiếu các kỹ thuật tối ưu hóa hệ thống như AutoML. Cuối cùng, đề án chưa mở rộng sang các kiến trúc mạng đồ thị hiện đại như GNN hoặc GAT, vốn được đánh giá là phù hợp hơn cho các bài toán học trên cấu trúc mạng phức tạp. Đây là những điểm mà các nghiên cứu tiếp theo có thể cải thiện để nâng cao hiệu quả và độ chính xác của mô hình dự đoán.

Hướng phát triển:

Trong tương lai, có thể mở rộng nghiên cứu để xây dựng mô hình đa chiều cho độ tin cậy, kết hợp dữ liệu định tính như hành vi và cảm xúc. Cần nghiên cứu thêm về động lực học của sự tin cậy và phát triển mô hình lai ghép giữa học sâu và các lý thuyết xã hội học, tâm lý học hành vi. Các ứng dụng trong phát hiện thông tin sai lệch, bảo mật mạng xã hội, và thương mại số sẽ làm nổi bật giá trị của đề tài.

