

ĐẠI HỌC QUỐC GIA HÀ NỘI
ĐẠI HỌC CÔNG NGHỆ

ĐỒNG THỊ NGỌC LAN

NGHIÊN CỨU XÂY DỰNG MÔ HÌNH DỰ BÁO GDP VÀ
MỘT SỐ CHỈ TIÊU KINH TẾ VĨ MÔ KHÁC DỰA TRÊN
DỮ LIỆU ĐA NGUỒN

Ngành: Hệ thống Thông tin
Mã số: 9480104

TÓM TẮT LUẬN ÁN TIẾN SĨ HỆ THỐNG THÔNG TIN

Hà Nội - 2025

Công trình được hoàn thành tại: Trường Đại học Công
nghệ, Đại học Quốc gia Hà Nội

Người hướng dẫn khoa học: PGS. TS. Nguyễn Hà Nam
TS. Nguyễn Thị Hậu

Phản biện: PGS.TS. Nguyễn Linh Giang

Phản biện: TS. Hoàng Đỗ Thanh Tùng

Phản biện: PGS.TS. Vũ Việt Vũ

Luận án sẽ được bảo vệ trước Hội đồng cấp
.....chấm luận án tiến sĩ tại
.....
vào hồi ... giờ ...phút, ngày ... thángnăm

Có thể tìm hiểu luận án tại:

- Thư viện Quốc gia Việt Nam

MỞ ĐẦU

Các chỉ tiêu kinh tế vĩ mô như tổng sản phẩm quốc nội (GDP), lạm phát, tỷ lệ thất nghiệp, đầu tư trực tiếp nước ngoài (FDI), v.v., được xem là những đại lượng phản ánh rõ nét trạng thái hiện tại và xu hướng vận động của nền kinh tế [97]. Dự báo các chỉ tiêu kinh tế vĩ mô (CTKTVT) đóng vai trò quan trọng trong nghiên cứu và hoạch định chính sách kinh tế ở hầu hết các Quốc Gia. Với tầm quan trọng đó, các tổ chức quốc tế như Quỹ Tiền tệ Quốc tế (IMF), Tổ chức Hợp tác và Phát triển Kinh tế (OECD), Ngân hàng Thế giới (WB), Ủy ban Châu Âu (European Commission), Ngân hàng Trung ương châu Âu (ECB), và Hội nghị Liên hợp quốc về Thương mại và Phát triển (UNCTAD) thường xuyên công bố các số liệu thống kê định kỳ và thực hiện dự báo cũng như phân tích chuyên sâu nhằm cung cấp luận cứ cho việc hoạch định chính sách tài khóa và tiền tệ [46], [47], [75], [105], [129], [135].

Tuy nhiên, trong thực tiễn, việc dự báo các CTKTVT thường gặp phải những thách thức đáng kể do đặc tính biến động phức tạp và khó lường của nền kinh tế toàn cầu. Một minh chứng rõ nét là sự điều chỉnh mạnh trong dự báo tăng trưởng GDP toàn cầu năm 2020 của Quỹ Tiền tệ Quốc tế (IMF). Cụ thể, trong Báo cáo Triển vọng Kinh tế Thế giới (WEO) công bố tháng 01/2020, IMF dự báo GDP toàn cầu sẽ tăng trưởng 3,3% [76]. Tuy nhiên, đến tháng 6/2020, dự báo này đã được điều chỉnh giảm xuống còn -4,9%, tương ứng mức sai lệch lên tới 8,3 điểm phần trăm [77]. Sự thay đổi này phản ánh tính bất định cao trong môi trường kinh tế toàn cầu và đặt ra yêu cầu cấp thiết trong việc cải tiến các phương pháp dự báo để thích ứng với dữ liệu biến động nhanh và phức tạp.

Trên thế giới, dự báo các CTKTVT đã được nghiên cứu và công bố trong nhiều luận án tiến sĩ [21], [23], [61], [63], [79], [86]. Các nghiên cứu này đề xuất các phương pháp tiếp cận từ truyền thống đến hiện đại. Trong nhiều thập kỷ, các phương pháp truyền thống như ARIMA, VAR hay hồi quy tuyến tính được sử dụng rộng rãi nhờ tính minh bạch và nền tảng lý thuyết vững chắc [65]. Tuy nhiên, các mô hình này bộc lộ hạn chế trong việc mô hình hóa quan hệ phi tuyến và khả năng phản ứng với các cú sốc bất ngờ [65]. Các mô hình hiện đại như Random Forest, XGBoost, LSTM hay Transformer cho thấy hiệu quả cao trong mô hình hóa các mối quan hệ phi tuyến và phức tạp trong dữ liệu chuỗi thời gian [42], [58]. Chẳng hạn, nghiên cứu của Chen và cộng sự công bố năm 2025 [38] đã chứng minh mô hình LSTM có khả năng

dự báo GDP bình quân đầu người với độ chính xác vượt trội so với các mô hình tuyến tính khi sử dụng CPI và tỷ lệ thất nghiệp làm biến đầu vào.

Bên cạnh đó, trong lý thuyết kinh tế học, chu kỳ kinh tế được xem là một yếu tố nền tảng cấu thành nên động thái vận động của nền kinh tế vĩ mô [34]. Mặc dù các mô hình kinh tế lượng như Markov-switching (Hamilton, 1989) [66], Bry-Boschan (1971) [33] đã có tiếp cận định lượng cho việc phân pha chu kỳ kinh tế, thì các mô hình học sâu hiện đại vẫn chủ yếu coi dữ liệu là chuỗi liên tục và đồng nhất.

Xuất phát từ những lí do nêu trên, luận án lựa chọn đề tài “Nghiên cứu xây dựng mô hình dự báo GDP và một số chỉ tiêu kinh tế vĩ mô khác dựa trên dữ liệu đa nguồn” nhằm hiện thực hóa **ba mục tiêu chính**: (1) phát triển một bộ dữ liệu kinh tế vĩ mô đa nguồn có tính chuẩn hóa cao phục vụ dự báo; (2) Triển khai và phân tích đánh giá sự tương thích của các mô hình dự báo kinh tế vĩ mô khác nhau đối với bộ dữ liệu đa nguồn đã được xây dựng từ các mô hình truyền thống tới các mô hình hiện đại; và (3) đề xuất một mô hình học sâu (DL) có khả năng tích hợp thông tin chu kỳ thông qua cơ chế attention thích ứng theo pha nhằm nâng cao hiệu quả dự báo trong môi trường dữ liệu phi tuyến và biến động cao.

Đối tượng nghiên cứu:

GDP và một số CTKTVM khác – là những chỉ số phản ánh trực tiếp trạng thái và xu thế của nền kinh tế.

Các mô hình dự báo GDP và một số CTKTVM khác.

Phạm vi nghiên cứu:

Về không gian: Nghiên cứu tập trung vào dữ liệu KTVM của Việt Nam, các nước Đông Nam Á và một số quốc gia tiêu biểu thuộc nhóm phát triển (G7), mới nổi (BRICS) và đang phát triển.

Về thời gian: Dữ liệu nghiên cứu bao phủ từ năm 1950 đến nay, tùy theo mức độ sẵn có của nguồn dữ liệu (ví dụ như dữ liệu của Nga từ năm 1991, Việt Nam từ 1970).

Về nội dung: Nghiên cứu tập trung vào việc xây dựng, thử nghiệm và đánh giá mô hình dự báo GDP và một số chỉ tiêu kinh tế vĩ mô khác.

Phương pháp nghiên cứu: kết hợp định tính và định lượng theo hướng thực nghiệm.

Cấu trúc của luận án

Cấu trúc luận án bao gồm bốn chương ngoài phần Mở đầu và Kết luận.

Chương 1 Tổng quan nghiên cứu: Trình bày tổng quan về các chỉ tiêu

KTVM phổ biến; đặc điểm của dữ liệu chuỗi thời gian kinh tế; các mô hình thông dụng trong dự báo GDP và một số CTKTVM khác. Chương này cũng tổng hợp các nghiên cứu tiêu biểu trong và ngoài nước, từ đó xác định khoảng trống nghiên cứu làm cơ sở cho định hướng mô hình hóa trong các chương tiếp theo.

Chương 2 Xây dựng bộ dữ liệu kinh tế vĩ mô đa nguồn: Trình bày chi tiết quy trình thu thập, hợp nhất, phân tích đặc trưng, phân tích thống kê mô tả, kiểm định tính dừng và xử lý dữ liệu kinh tế vĩ mô từ nhiều nguồn khác nhau. Bao gồm các bước tiền xử lý như làm sạch, mã hóa, chuẩn hóa, xử lý giá trị thiếu,... Kết quả nghiên cứu được sử dụng trong thực nghiệm của chương 3 và chương 4.

Chương 3 Các mô hình trong dự báo chỉ tiêu kinh tế vĩ mô: Đánh giá 3 nhóm mô hình là mô hình kinh tế lượng, học máy và học sâu trong dự báo GDP và một số chỉ tiêu kinh tế vĩ mô khác. Việc kiểm chứng hiệu suất của các nhóm mô hình đã được chúng tôi triển khai trong các công trình [CT1]–[CT4] trên các chỉ tiêu đa dạng khác nhau.

Chương 4 Mô hình học sâu với cơ chế thích ứng dự báo GDP Quốc Gia: Trình bày mô hình đề xuất, thuật toán huấn luyện, môi trường triển khai, chiến lược tối ưu hóa siêu tham số, và kỹ thuật chống quá khớp. Chương này cũng thực hiện so sánh kết quả dự báo của mô hình đề xuất với các mô hình kinh tế lượng, học máy và học sâu, đồng thời phân tích ý nghĩa thực tiễn của kết quả đạt được. Các kết quả nghiên cứu được công bố trong công trình [CT5].

Cuối cùng, phần kết luận tóm tắt những kết quả đạt được đồng thời đề xuất hướng nghiên cứu và giải pháp bổ sung trong tương lai.

Chương 1. TỔNG QUAN NGHIÊN CỨU

Chương này được xây dựng với mục tiêu cung cấp cơ sở lý luận về GDP và một số CTKTVM tiêu biểu và các mô hình dự báo CTKTVM, tập trung vào ba trục nội dung chính: (i) Một số chỉ tiêu kinh tế vĩ mô tiêu biểu và đặc điểm của dữ liệu kinh tế; (ii) một số mô hình dự báo thông dụng trong dự báo GDP và một số CTKTVM khác; (iii) phân tích các nghiên cứu tiêu biểu trong và ngoài nước nhằm xác định khoảng trống nghiên cứu mà luận án hướng đến.

1.1. Các chỉ tiêu kinh tế vĩ mô tiêu biểu

Định nghĩa 1: Tổng sản phẩm quốc nội (Gross Domestic Product – GDP) là tổng giá trị thị trường của tất cả hàng hóa và dịch vụ cuối cùng được sản xuất trong phạm vi lãnh thổ một quốc gia trong một khoảng thời gian nhất định thường là một quý hoặc một năm

Định nghĩa 2 : Lạm phát là tình trạng mức giá chung của hàng hóa và dịch vụ trong nền kinh tế tăng lên liên tục trong một khoảng thời gian, làm giảm sức mua của đồng tiền. Tỷ lệ lạm phát là phần trăm thay đổi của mức giá chung trong một khoảng thời gian nhất định (thường là hàng năm).

Định nghĩa 3: Chỉ số giá tiêu dùng (Consumer Price Index – CPI) đo lường sự thay đổi trung bình của giá cả một "giỏ hàng hóa và dịch vụ" tiêu biểu mà người tiêu dùng mua sắm trong một khoảng thời gian nhất định, so với một năm cơ sở.

Định nghĩa 4: Tỷ lệ thất nghiệp là tỷ lệ phần trăm số người trong lực lượng lao động (những người đủ tuổi lao động, có khả năng và mong muốn làm việc nhưng không tìm được việc làm) so với tổng lực lượng lao động.

Định nghĩa 5: Đầu tư trực tiếp nước ngoài (Foreign Direct Investment – FDI) là hình thức đầu tư xuyên biên giới trong đó một cá nhân hoặc tổ chức ở một quốc gia đầu tư vào một doanh nghiệp ở quốc gia khác nhằm mục tiêu sở hữu lâu dài và kiểm soát hoạt động sản xuất, kinh doanh của doanh nghiệp đó.

Định nghĩa 6: Tiêu dùng là tổng giá trị hàng hóa và dịch vụ cuối cùng được các hộ gia đình và cá nhân mua để phục vụ mục đích sinh hoạt, giải trí và duy trì cuộc sống. Đây là thành phần lớn nhất trong GDP theo phương pháp chi tiêu.

Định nghĩa 7: Tổng cầu là tổng chi tiêu cho hàng hóa và dịch vụ cuối cùng trong nền kinh tế, bao gồm: Tiêu dùng (C), Đầu tư (I), Chi tiêu Chính phủ (G) và Xuất khẩu ròng (NX). Tổng cầu phản ánh toàn bộ nhu cầu hàng hóa dịch vụ trong một nền kinh tế tại một mức giá nhất định.

Định nghĩa 8: Cán cân thương mại là chênh lệch giữa giá trị xuất khẩu và giá trị nhập khẩu hàng hóa của một quốc gia trong một khoảng thời gian nhất định, thường tính theo tháng, quý hoặc năm.

Định nghĩa 9: Tỷ giá hối đoái là giá của một đồng tiền này được biểu thị bằng một đồng tiền khác. Ví dụ, 1 USD = 25.000 VND.

Định nghĩa 10: Năng suất lao động là chỉ tiêu đo lường sản phẩm (giá trị gia tăng hoặc sản lượng) mà một lao động tạo ra trong một đơn vị thời gian nhất định, thường là một năm. Chỉ tiêu này phản ánh hiệu quả sử dụng lao động trong quá trình sản xuất.

Định nghĩa 11: Cung tiền (Money Supply) là tổng lượng tiền đang lưu

thông trong nền kinh tế tại một thời điểm nhất định, bao gồm tiền mặt và các khoản tiền gửi có thể chuyển đổi nhanh thành tiền mặt

1.2. Lý thuyết dự báo kinh tế vĩ mô – đặc điểm chuỗi thời gian

Dự báo kinh tế vĩ mô chủ yếu dựa trên dữ liệu chuỗi thời gian, nên việc hiểu rõ bản chất và đặc điểm kỹ thuật của loại dữ liệu này là điều kiện tiên quyết. Chuỗi thời gian kinh tế vĩ mô thường mang tính không dừng, chu kỳ, phi tuyến và có độ trễ, phản ánh mối quan hệ phức tạp giữa các biến số. Việc xử lý chuỗi bao gồm kiểm định tính dừng, tách xu hướng và chu kỳ, phân tích tự tương quan nhằm đảm bảo phù hợp với giả định mô hình. Trong khi các mô hình tuyến tính như ARIMA, VAR phù hợp với chuỗi ổn định, thì các mô hình học sâu như LSTM, GRU, Transformer tỏ ra hiệu quả hơn trong xử lý dữ liệu phi tuyến và dài hạn.

1.3. Một số mô hình dự báo KTVM thông dụng

Trước khi học máy và AI được áp dụng rộng rãi trong kinh tế học, các mô hình dự báo truyền thống như hồi quy, ARIMA và VAR là công cụ chủ đạo, với nền tảng lý thuyết vững chắc, dễ diễn giải và phù hợp với dữ liệu kinh tế dạng chuỗi thời gian. Mô hình hồi quy mô tả mối quan hệ nhân quả giữa biến phụ thuộc và các biến giải thích; ARIMA tập trung vào dự báo chuỗi đơn biến; còn VAR mở rộng dự báo đa biến với các mối liên hệ trễ giữa các chỉ tiêu kinh tế.

Dù hiệu quả trong nhiều thập kỷ, các mô hình này gặp hạn chế khi đối mặt với dữ liệu phi tuyến, cú sốc cấu trúc và yêu cầu xử lý lượng lớn biến số. Tuy nhiên, việc hiểu rõ các mô hình truyền thống vẫn là nền tảng cần thiết để đánh giá và phát triển các phương pháp học máy tiên tiến trong dự báo kinh tế vĩ mô.

1.4. Chỉ số đánh giá hiệu suất của mô hình

Để đánh giá hiệu suất của các mô hình dự báo kinh tế vĩ mô, đặc biệt là mô hình chuỗi thời gian hoặc học máy, các chỉ số sai số (error metrics) được sử dụng nhằm đo lường mức độ chênh lệch giữa giá trị dự báo và giá trị thực tế là: MAE, RMSE, MAPE, R^2 .

1.5. Tình hình nghiên cứu trong và ngoài nước về dự báo GDP và một số CTKTVM khác

1.5.1. Một số nghiên cứu trên thế giới

Tổng quan các nghiên cứu trên thế giới cho thấy sự phát triển đa dạng và tiến bộ vượt bậc trong việc dự báo các chỉ tiêu kinh tế vĩ mô. Trước đây, với đặc điểm kinh tế vĩ mô tương đối ổn định và dữ liệu hạn chế, các nghiên cứu dự báo chủ yếu áp dụng các mô hình kinh tế lượng truyền thống như ARIMA, VAR, VECM, DSGE với nền tảng lý thuyết vững chắc và khả năng diễn giải tốt. Tuy nhiên, trong những năm gần đây, cùng với sự gia tăng mạnh mẽ của dữ liệu và độ phức tạp của các yếu tố kinh tế, các nghiên cứu

trên thế giới đã tiếp cận các mô hình học máy và học sâu hiện đại như Random Forest, LSTM, Transformer nhờ ưu thế về khả năng xử lý phi tuyến và dữ liệu lớn, bổ sung cho các hạn chế của mô hình truyền thống.

1.5.2. Một số nghiên cứu ở Việt Nam

Trái ngược với tình hình nghiên cứu sôi động của các Quốc gia trên thế giới về ứng dụng học máy và học sâu trong dự báo kinh tế vĩ mô thì các nghiên cứu tại Việt Nam vẫn còn ở giai đoạn khởi phát. Trong hơn hai thập kỷ qua, các nghiên cứu trong nước có thể thấy hầu hết tiếp cận theo hướng sử dụng thống kê kinh tế lượng, phần lớn các nghiên cứu dự báo kinh tế vĩ mô ở Việt Nam sử dụng các mô hình kinh tế lượng như ARIMA, VAR, mô hình cấu trúc đồng thời hoặc hồi quy tuyến tính. Các nghiên cứu áp dụng ML/DL ở Việt Nam hiện nay vẫn còn rời rạc, quy mô nhỏ và thiếu hệ thống, chưa hình thành khung lý thuyết mạnh hoặc nền tảng dữ liệu đủ rộng để triển khai mô hình dự báo tích hợp và có tính mở rộng cao. Đặc biệt, việc áp dụng các mô hình hiện đại như LSTM, GRU, Transformer, TCN cho dự báo GDP và một số CTKTVM khác gần như chưa được đề cập trong bối cảnh Việt Nam.

1.6. Khoảng trống nghiên cứu và đề xuất giải pháp

1.6.1. Khoảng trống nghiên cứu

Trên cơ sở khảo sát các công trình trong nước và quốc tế, có thể nhận diện một số khoảng trống chính mà luận án này hướng đến giải quyết:

Thứ nhất, việc xây dựng một bộ dữ liệu kinh tế vĩ mô tiêu chuẩn, đồng bộ và tích hợp từ nhiều nguồn đáng tin cậy vẫn chưa được chú trọng đúng mức,

Thứ hai, thiếu các mô hình dự báo có khả năng học linh hoạt và thích ứng theo pha chu kỳ kinh tế,

Thứ ba, thiếu đánh giá khả năng tổng quát hóa của mô hình trong môi trường liên quốc gia.

1.6.2. Đề xuất giải pháp

Để khắc phục các khoảng trống nêu trên, luận án đề xuất một hướng tiếp cận tích hợp dựa trên ba trụ cột chính:

Thứ nhất, tích hợp dữ liệu đa nguồn. Dữ liệu sẽ được xử lý và chuẩn hóa để đưa về cấu trúc phù hợp cho mô hình học máy/học sâu;

Thứ hai, đánh giá hiệu suất các mô hình hiện có để xác định hướng phát triển phù hợp;

Thứ ba, xây dựng mô hình học sâu linh hoạt và thích ứng, sử dụng các kiến trúc hiện đại có khả năng học các quan hệ phi tuyến, chu kỳ và thích ứng theo thời gian.

1.7. Kết luận chương

Chương 1 đã cung cấp nền tảng lý luận cho luận án thông qua việc hệ

thống hóa các chỉ tiêu kinh tế vĩ mô tiêu biểu, tổng quan lý thuyết dự báo chuỗi thời gian, và tổng quan các mô hình thông dụng trong dự báo chỉ tiêu kinh tế vĩ mô. Bên cạnh đó, chương 1 cũng đã tổng hợp và đánh giá một cách hệ thống các nghiên cứu liên quan trong và ngoài nước, qua đó chỉ ra những khoảng trống học thuật đáng chú ý.

Chương 2. XÂY DỰNG BỘ DỮ LIỆU KINH TẾ VĨ MÔ ĐA NGUỒN

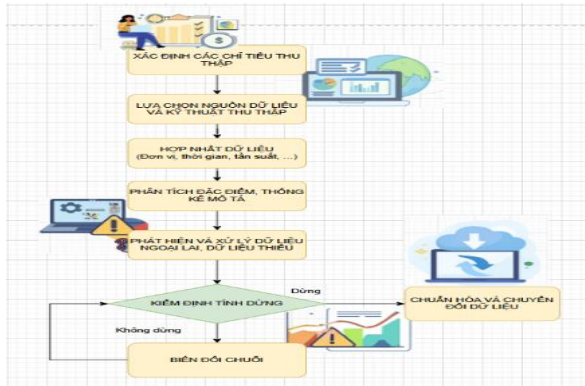
Mục tiêu của chương này là xây dựng bộ dữ liệu kinh tế vĩ mô đa nguồn đáng tin cậy, đáp ứng đồng thời yêu cầu của các mô hình thống kê và học sâu hiện đại. Nội dung gồm: (i) Quy trình xây dựng bộ dữ liệu KTVM đa nguồn, (ii) Xác định các CTKTVM mô cần thu thập, (iii) Xác định các nguồn dữ liệu KTVM và kỹ thuật thu thập, (iv) Hợp nhất dữ liệu; (v) Phân tích đặc điểm dữ liệu, (vi) phát hiện xử lý ngoại lệ và gán giá trị bị thiếu, (vii) kiểm định chuỗi thời gian, và (viii) chuẩn hóa và chuyển đổi dữ liệu hướng tốt nhất cho việc xây dựng bộ dữ liệu đầu vào chất lượng cao cho mô hình học máy.

2.1. Giới thiệu

Một trong những thách thức lớn khi làm việc với dữ liệu kinh tế vĩ mô là sự không nhất quán giữa các nguồn dữ liệu về khái niệm chỉ tiêu, đơn vị đo lường, tần suất thu thập và mức độ đầy đủ theo thời gian [49], [131]. Tuy nhiên, mỗi nguồn dữ liệu chỉ tập trung vào một số lĩnh vực hoặc quốc gia nhất định. Không có một nguồn đơn lẻ nào cung cấp đầy đủ và toàn diện tất cả các chỉ tiêu cần thiết với mục tiêu và độ phủ thống nhất [131]. Do đó, việc xây dựng bộ dữ liệu đáp ứng các tiêu chí của mô hình dự báo không chỉ là mở rộng độ phủ thông tin, mà còn là đảm bảo tính nhất quán, khả năng so sánh và đồng bộ hóa dữ liệu giữa các quốc gia – yếu tố quyết định để huấn luyện và đánh giá mô hình một cách khách quan, hiệu quả trong bối cảnh liên quốc gia.

2.2. Quy trình xây dựng bộ dữ liệu kinh tế vĩ mô đa nguồn

Quy trình xây dựng bộ dữ liệu KTVM đa nguồn gồm các bước trong hình 2.1.



Hình 2. 1 Quy trình xây dựng bộ dữ liệu KTMV đa nguồn

2.3. Xác định CTKTMV cần thu thập

Luận án lựa chọn các chỉ tiêu dựa trên ba tiêu chí: (1) tính đại diện cho các lĩnh vực cốt lõi của nền kinh tế như tăng trưởng, giá cả, lao động, thương mại và tài chính – tiền tệ; (2) mức độ sẵn có và độ tin cậy của dữ liệu từ các nguồn chính thống như World Bank, IMF, PWT; và (3) mức độ phù hợp với mục tiêu nghiên cứu cụ thể như dự báo GDP, lạm phát hay thất nghiệp. Các chỉ tiêu thu thập thuộc các nhóm chỉ tiêu: Nhóm các chỉ tiêu nhóm đo lường quy mô nền kinh tế, lực lượng lao động và năng suất theo thời gian; Nhóm đo lường tổng sản lượng và tích lũy vốn theo mức giá hiện hành đã điều chỉnh theo PPP; Nhóm các biến dựa trên tài khoản Quốc gia; Nhóm chỉ số mức giá/tỷ giá điều chỉnh sức mua theo PPP; Nhóm các chỉ tiêu mức giá theo PPP của các cấu phần chi tiêu và vốn; Nhóm các biến tỷ trọng trong GDP theo PPP hiện hành; Nhóm các chỉ tiêu tiêu dùng và lạm phát.

2.4. Lựa chọn nguồn dữ liệu

Căn cứ trên nhóm các chỉ tiêu kinh tế vĩ mô cần thu thập xác định ở mục 2.2, đối với dữ liệu của một số Quốc Gia trên thế giới, WB và PWT được lựa chọn làm hai nguồn dữ liệu chính; đối với Việt Nam luận án lựa chọn cục Thống kê và FiinGroup (FiinTech) làm hai nguồn dữ liệu chính.

2.5. Hợp nhất dữ liệu

Mục này thực hiện hợp nhất bộ dữ liệu bao gồm: chuẩn hóa định danh và đơn vị đo lường cho các CTKTMV, đồng bộ thời gian và tần suất, gộp bảng dữ liệu theo khóa chung, xử lý dữ liệu thiếu và kiểm tra toàn vẹn dữ liệu.

2.6. Phân tích đặc điểm, thống kê mô tả

2.7. Phát hiện ngoại lệ và xử lý dữ liệu thiếu

Tỷ trọng sai lệch thống kê trong GDP chi tiêu) của Việt Nam có một giá trị âm lớn vào năm 1990 (- 0.237363), cho thấy khả năng mất cân đối tạm thời giữa các thành phần chi tiêu hoặc sai số thống kê trong năm này. Tuy nhiên, các giá trị ngoại lệ phản ánh hiện tượng kinh tế quan trọng (ví dụ như GDP âm) cần được giữ lại để không làm mất đi bản chất thực tế của chuỗi dữ liệu.

Theo phân tích đặc điểm bộ dữ liệu mục 2.5, các giá trị bị thiếu chủ yếu nằm ở giai đoạn trước năm 1990, các giá trị này được xử lý thay thế bằng trung bình/trung vị, nội suy tuyến tính với các chỉ tiêu tiêu dùng, đầu tư, chi tiêu chính phủ, xuất nhập khẩu; và ước lượng với chỉ tiêu tỷ giá. Riêng Việt Nam một số chỉ tiêu không có dữ liệu bao gồm: pl_k , $rwtfpna$, $rtfpna$, $rkna$, $labsh$, irr nên các chỉ tiêu được loại bỏ khỏi bộ dữ liệu Việt Nam.

2.8. Kiểm định chuỗi thời gian

Kết quả kiểm định cho thấy phần lớn các biến kinh tế vĩ mô như GDP thực theo tiêu dùng ($rgdpe$), GDP theo sản lượng ($rgdpo$), tỷ trọng chi tiêu (csh_c , csh_i), và mức giá tương đối (p_pl_c , p_pl_g), hc , $v.v...$ đều không dừng tại cấp độ gốc (p -value ADF > 0.05 và p -value KPSS < 0.05). Kiểu dừng $I(1)$, tức là cần sai phân bậc nhất để trở nên dừng. Một số biến như pop , emp , và avh có ADF p -value nhỏ hơn 0.05 nhưng đồng thời cũng có KPSS p -value nhỏ hơn 0.05. Điều này cho thấy kết quả kiểm định không đồng thuận, làm phát sinh nghi ngờ về tính dừng tại cấp độ gốc. Theo hướng thận trọng và theo thông lệ học thuật, những biến này được phân loại là $I(2)$ – tức là cần sai phân hai lần để đạt trạng thái dừng và có thể sử dụng trong mô hình dự báo. Từ kết quả kiểm định, hầu hết dữ liệu đầu vào cần được tiền xử lý bằng cách sai phân trước khi đưa vào mô hình dự báo để đảm bảo tính dừng, ổn định và hiệu quả trong dự báo. Những biến dừng ở $I(0)$ được giữ nguyên và sử dụng như biến đầu vào tham chiếu hoặc điều khiển trong mô hình dự báo.

2.9. Chuẩn hóa và chuyển đổi dữ liệu

Luận án lựa chọn phương pháp chuẩn hóa Z-score để xử lý các biến đầu vào trước khi đưa vào quá trình huấn luyện. Sau khi huấn luyện mô hình, giá trị đầu ra (ví dụ GDP dự báo) được đưa trở lại thang đo gốc thông qua phép nghịch đảo của quá trình chuẩn hóa để phục vụ cho việc đánh giá và so sánh với dữ liệu thực tế.

2.10. Kết luận chương

Chương 2 đã tập trung xây dựng bộ dữ liệu KTVM từ nhiều nguồn khác nhau phục vụ cho các mô hình dự báo GDP và một số chỉ tiêu kinh tế vĩ mô

khác. Trên cơ sở phân tích các tài liệu thực nghiệm và lý thuyết đã công bố, nghiên cứu đã lựa chọn các chỉ tiêu có ý nghĩa kinh tế học và tiềm năng đóng vai trò biến đầu vào quan trọng trong dự báo GDP và một số CTKTVM khác. Đồng thời, các kỹ thuật xử lý dữ liệu chuỗi thời gian như kiểm tra tính dừng, chuẩn hóa và chuyển đổi đã được áp dụng nhằm đảm bảo tính phù hợp cho việc huấn luyện các mô hình học máy và học sâu. Những đóng góp này đóng vai trò quan trọng trong việc tăng cường hiệu quả và độ chính xác của các mô hình dự báo được đề xuất trong các phần sau.

Chương 3. MÔ HÌNH DỰ BÁO GDP VÀ MỘT SỐ CHỈ TIÊU KINH TẾ VĨ MÔ KHÁC

Chương 3 đánh giá hiệu quả của các mô hình trong dự báo một số chỉ tiêu kinh tế vĩ mô. Các kết quả nghiên cứu là cơ sở để so sánh và bàn luận về khả năng ứng dụng của từng phương pháp.

3.1. Mô hình kinh tế lượng

Nguyên lý hoạt động:

Mô hình hồi quy tuyến tính dựa trên giả định về mối quan hệ tuyến tính giữa biến phụ thuộc và các biến độc lập, với hệ số được ước lượng bằng phương pháp bình phương tối thiểu (OLS), giúp lượng hóa tác động của các yếu tố kinh tế đến biến mục tiêu.

Mô hình ARIMA kết hợp ba thành phần: tự hồi quy (AR), sai phân (I) và trung bình trượt (MA), nhằm mô hình hóa động lực nội tại của chuỗi thời gian đơn biến, đặc biệt thích hợp cho dữ liệu có xu hướng và không cân bằng giải thích bên ngoài.

Mô hình VAR mở rộng AR sang nhiều biến, mô hình hóa quan hệ động giữa các biến kinh tế mà không cần xác định quan hệ nhân quả trước. VAR thường được sử dụng để dự báo và phân tích tác động chính sách thông qua hàm phản ứng xung lực và phân rã phương sai.

Mô hình VECM là biến thể của VAR áp dụng cho các chuỗi có quan hệ đồng liên kết. VECM mô tả đồng thời động học ngắn hạn và điều chỉnh sai số để duy trì cân bằng dài hạn giữa các biến.

Mô hình DSGE là mô hình vĩ mô kết cấu mô phỏng hành vi tối ưu của các tác nhân kinh tế dưới tác động của các cú sốc ngẫu nhiên, được giải thông qua xấp xỉ tuyến tính quanh trạng thái cân bằng, giúp phân tích chính sách và động thái kinh tế trong dài hạn.

Ưu điểm:

Các mô hình kinh tế lượng truyền thống có nền tảng lý thuyết vững chắc, xuất phát từ các nguyên lý kinh tế học vĩ mô, cho phép mô hình hóa quan hệ nhân quả giữa các biến một cách có cấu trúc và dễ diễn giải. Mô

hình như DSGE hay VAR không chỉ phản ánh tốt phản ứng của nền kinh tế trước các cú sốc chính sách, mà còn giúp định lượng và hỗ trợ hoạch định chính sách. Tham số trong mô hình mang ý nghĩa kinh tế rõ ràng, giúp giải thích kết quả theo tư duy học thuật. Các mô hình này cũng được áp dụng rộng rãi trong thực tiễn, điển hình tại FED, ECB hay BOJ. Ngoài ra, hệ thống công cụ thống kê đi kèm như kiểm định đồng liên kết, phân rã phương sai và hàm phản ứng xung động góp phần nâng cao độ tin cậy và chiều sâu phân tích định lượng. Nhờ sự kết hợp giữa lý thuyết, khả năng diễn giải, tính ứng dụng và công cụ hỗ trợ mạnh mẽ, các mô hình này vẫn giữ vai trò quan trọng trong nghiên cứu và điều hành kinh tế.

Hạn chế:

Các mô hình kinh tế lượng truyền thống như ARIMA, VAR hay VECM gặp nhiều hạn chế khi giả định mối quan hệ tuyến tính cố định giữa các biến, trong khi thực tế kinh tế thường phi tuyến và biến động theo chu kỳ hoặc cú sốc bên ngoài. Ngoài ra, các mô hình này khó mở rộng trong bối cảnh dữ liệu lớn, dễ bị quá khớp và mất ổn định khi số biến vượt quá số quan sát. Chúng cũng nhạy cảm với việc lựa chọn mô hình, độ trễ và biến đầu vào, đòi hỏi kinh nghiệm chuyên sâu từ nhà nghiên cứu. Đặc biệt, hiệu suất dự báo của các mô hình này thường kém hơn so với các phương pháp học máy hiện đại như random forest, boosting hay mạng nơ-ron, vốn có khả năng học tốt hơn các mối quan hệ phi tuyến phức tạp. Những hạn chế này đang thúc đẩy xu hướng kết hợp các phương pháp truyền thống với kỹ thuật học máy để nâng cao hiệu quả phân tích và dự báo trong kinh tế hiện đại.

Điều kiện áp dụng các mô hình kinh tế lượng

Việc áp dụng các mô hình thống kê truyền thống như ARIMA hay VAR đòi hỏi một số điều kiện kỹ thuật để đảm bảo độ tin cậy. Dữ liệu chuỗi thời gian cần đủ dài, có tần suất ổn định (quý/năm) và phải đạt tính dừng, nếu không cần được chuyển đổi phù hợp. Số lượng biến đầu vào nên giới hạn để tránh đa cộng tuyến và quá khớp. Các mô hình này đặc biệt phù hợp với các biến mục tiêu phục vụ phân tích chính sách do khả năng diễn giải tốt. Một giả định quan trọng là cấu trúc mối quan hệ giữa các biến phải ổn định theo thời gian; nếu có cú sốc chính sách hoặc thay đổi cơ cấu, hiệu quả dự báo sẽ giảm. Ngoài ra, mô hình truyền thống chỉ xử lý tốt dữ liệu định lượng có cấu trúc, không phù hợp với dữ liệu phi cấu trúc như văn bản hay dữ liệu thời gian thực từ Internet.

3.2. Một số mô hình học máy

Nguyên lý hoạt động:

Các mô hình học máy hoạt động dựa trên nguyên tắc học từ dữ liệu mà không giả định trước cấu trúc hàm giữa biến đầu vào và đầu ra, cho phép phát hiện các quan hệ phi tuyến tiềm ẩn. Random Forest (RF) là phương pháp

học tập hợp, kết hợp nhiều cây quyết định huấn luyện trên các tập dữ liệu bootstrap, giúp giảm phương sai và chống overfitting hiệu quả. XGBoost sử dụng cơ chế boosting tuần tự, liên tục cải thiện sai số còn lại của mô hình trước, đồng thời tích hợp các kỹ thuật như regularization, tree pruning sớm và tính toán song song, giúp tối ưu hóa tốc độ và hiệu suất khi xử lý dữ liệu lớn và phức tạp. Support Vector Machine (SVM) hoạt động bằng cách tìm siêu phẳng tối ưu để phân tách hoặc hồi quy dữ liệu trong không gian đặc trưng cao, sử dụng các hàm kernel như RBF, Polynomial để xử lý các quan hệ phi tuyến. Support Vector Regression (SVR) là biến thể của SVM dành cho bài toán dự báo giá trị liên tục. Multilayer Perceptron (MLP) là mạng nơ-ron truyền thẳng gồm nhiều lớp, thường dùng các độ trễ chuỗi làm đầu vào trong các bài toán chuỗi thời gian.

Ưu điểm:

Các mô hình học máy có thể học các mối quan hệ phi tuyến [58] và hiệu ứng tương tác giữa các biến đầu vào mà không cần giả định trước về dạng hàm, khắc phục hạn chế của các mô hình tuyến tính truyền thống. Các thuật toán như XGBoost và Randomforest có thể xử lý tốt dữ liệu thiếu, biến phân loại, và dữ liệu có nhiễu, phù hợp với đặc trưng của dữ liệu kinh tế thực tế. Một số mô hình như Random Forest và Gradient Boosting cung cấp chỉ số tầm quan trọng của biến, hỗ trợ đánh giá mức độ ảnh hưởng của các yếu tố kinh tế đến biến mục tiêu. SVM có thể cung cấp một giải pháp tối ưu toàn cục và duy nhất. Nó đặc biệt trong việc xử lý các mẫu phi tuyến tính phức tạp thường có trong dữ liệu chuỗi thời gian. SVM cũng có thể xử lý các phần nhiễu phân phối chuẩn của chuỗi bằng cách sử dụng hàm mất mát Gaussian [137].

Hạn chế:

Dù các mô hình học máy như Random Forest, XGBoost, SVM hay MLP cho thấy hiệu quả trong xử lý dữ liệu phi tuyến và tương tác phức tạp, chúng vẫn tồn tại một số hạn chế khi dự báo các chỉ tiêu kinh tế vĩ mô. So với mô hình kinh tế lượng, học máy thiếu khả năng diễn giải rõ ràng về mặt kinh tế và quan hệ nhân quả. Ngoài ra, các thuật toán này không tích hợp động học chuỗi thời gian, dẫn đến phụ thuộc vào kỹ thuật tạo đặc trưng thủ công. Một số mô hình như SVM và MLP cũng gặp khó khăn khi mở rộng sang dữ liệu kinh tế vĩ mô lớn, không đồng nhất.

Điều kiện áp dụng:

Các mô hình học máy truyền thống như Random Forest, XGBoost và SVM phù hợp trong bối cảnh dữ liệu kinh tế có quy mô vừa, cấu trúc phi tuyến và nhiễu. Chúng hoạt động hiệu quả với dữ liệu không quá lớn nhưng đủ thông tin để huấn luyện, đặc biệt khi quan hệ giữa biến đầu vào và đầu ra phức tạp mà mô hình tuyến tính không nắm bắt được. Các thuật toán

này hỗ trợ chọn lọc đặc trưng tự động, giúp tránh overfitting và tăng khả năng tổng quát hóa trong các tập dữ liệu nhiều biến nhưng chuỗi thời gian ngắn. Ngoài ra, thời gian huấn luyện ngắn, dễ triển khai và khả năng cung cấp thông tin về mức độ quan trọng của biến đầu vào khiến chúng hữu ích trong phân tích và hỗ trợ chính sách. Nhìn chung, đây là lựa chọn phù hợp trong các bài toán dự báo có yêu cầu cao về hiệu suất nhưng hạn chế về tài nguyên tính toán và cần mức độ diễn giải vừa phải.

3.3. Mô hình học sâu

Nguyên lý hoạt động

Học sâu là một nhánh mở rộng của học máy, trong đó các mô hình mạng nơ-ron nhân tạo nhiều tầng (deep neural networks) được sử dụng để học các đặc trưng có cấp độ trừu tượng cao từ dữ liệu. Mạng nơ-ron nhân tạo (ANN) là một tập hợp con của học máy dựa trên các thuật toán số mô phỏng các nơ-ron sinh học, có khả năng nắm bắt các mẫu quan trọng từ chuỗi thời gian phức tạp và phi tuyến tính [67]. Mạng nơ-ron hồi tiếp (RNN) được thiết kế để xử lý dữ liệu tuần tự bằng cách truyền đầu ra của bước trước đó làm đầu vào của bước hiện tại, cho phép chúng "ghi nhớ hầu hết thông tin quá khứ". Chúng sử dụng cùng các tham số cho tất cả các đầu vào, giảm độ phức tạp của tham số. LSTM và GRU là các biến thể của RNN được phát triển để giải quyết vấn đề gradient biến mất và bùng nổ. LSTM sử dụng các cổng (cổng đầu vào, cổng quên, cổng đầu ra) để chọn lọc ghi nhớ hoặc quên thông tin. Mạng tích chập sử dụng các phép toán tích chập để nắm bắt cả phụ thuộc ngắn hạn và dài hạn. Transformer: Ban đầu được thiết kế cho các tác vụ xử lý ngôn ngữ tự nhiên (NLP), các mô hình Transformer đã được điều chỉnh cho dự báo chuỗi thời gian. Điểm mạnh cốt lõi của chúng là cơ chế tự chú ý (self-attention), cho phép chúng cân nhắc tầm quan trọng của các bước thời gian khác nhau trong một chuỗi và nắm bắt các phụ thuộc tầm xa một cách hiệu quả. Chúng có thể xử lý tất cả các điểm dữ liệu đồng thời, cho phép song song hóa [110].

Ưu điểm:

Mạng nơ-ron là các thiết bị xấp xỉ hàm hiệu quả cho việc học bất kỳ hàm nào và có khả năng nhận dạng mẫu cao. Chúng có thể nắm bắt các mối quan hệ phi tuyến tính phức tạp trong dữ liệu kinh tế. RNN và các biến thể của chúng (LSTM, GRU, TCN, Transformer) đặc biệt hữu ích trong các tình huống cần biểu diễn các mối quan hệ thời gian giữa đầu vào và đầu ra, vượt qua các hạn chế của các mô hình tuyến tính truyền thống. Sự phát triển từ các ANN tổng quát đến các kiến trúc chuyên biệt dựa trên hồi quy (RNN, LSTM, GRU) và chú ý (Transformer) phản ánh nỗ lực liên tục để nắm bắt tốt hơn các phụ thuộc thời gian và các mẫu tầm xa trong dữ liệu chuỗi thời gian. Đây là một phản ứng trực tiếp đối với tính chất tuần tự vốn có của dữ

liệu kinh tế. Transformer có thể xử lý tất cả dữ liệu đầu vào đồng thời, tận dụng xử lý song song trên phần cứng hiện đại, dẫn đến thời gian đào tạo nhanh hơn đáng kể. Mạng nơ-ron Bayesian (BNN) có thể mô hình hóa các phi tuyến tính và biến thiên thời gian cho các tập dữ liệu kinh tế vĩ mô và tài chính. Hiệu suất của NN ít nhạy cảm hơn với việc tiền xử lý dữ liệu và có thể tự điều hòa.

Hạn chế:

Các mô hình mạng nơ-ron, đặc biệt là LSTM, GRU và Transformer, thường yêu cầu sức mạnh tính toán lớn và thời gian huấn luyện dài do cấu trúc phức tạp và nhiều tham số. RNN đơn giản dễ gặp vấn đề gradient biến mất, trong khi LSTM và GRU chỉ giảm thiểu chứ không loại bỏ hoàn toàn hiện tượng này. Ngoài ra, các mô hình học sâu thường mang tính “hộp đen”, gây khó khăn trong việc diễn giải kết quả. Hiệu suất của mô hình nhạy cảm với thiết kế mạng khi dữ liệu huấn luyện hạn chế. Do đó, tồn tại sự đánh đổi giữa độ chính xác, khả năng diễn giải và chi phí tính toán trong việc lựa chọn mô hình học sâu.

Điều kiện áp dụng:

Các mô hình học sâu, đặc biệt là LSTM, phù hợp trong bối cảnh chuỗi thời gian có quan hệ phụ thuộc dài hạn, cấu trúc phi tuyến phức tạp và chịu ảnh hưởng chu kỳ rõ rệt. Chúng tỏ ra hiệu quả khi dữ liệu đa chiều, có nguồn gốc đa dạng và biến động theo thời gian. Tuy nhiên, việc áp dụng học sâu đòi hỏi hạ tầng tính toán mạnh và quy trình tối ưu hóa siêu tham số chặt chẽ. Bên cạnh đó, tính “hộp đen” của mô hình đặt ra thách thức trong giải thích. Do đó, học sâu không thay thế hoàn toàn các phương pháp truyền thống, mà đóng vai trò bổ trợ quan trọng trong các bài toán dự báo kinh tế hiện đại, tùy theo mục tiêu và điều kiện dữ liệu cụ thể.

3.4. Thực nghiệm dự báo GDP và một số CTKTVM khác

Chúng tôi đã kiểm chứng hiệu quả của các nhóm mô hình khác nhau trong dự báo kinh tế, bao gồm: (1) mô hình kinh tế lượng trong công trình [CT1] và [CT2]; mô hình học máy và học sâu trong các công trình [CT3] và [CT4]. Những kết quả này được tổng hợp, đối chiếu và mở rộng nhằm xác định mô hình có năng lực cao nhất đối với dự báo GDP và một số chỉ tiêu kinh tế khác. Trong phần này, chúng tôi triển khai các mô hình trên với bộ dữ liệu đa nguồn đã được xây dựng ở chương 2 để tìm hiểu sự phù hợp của các mô hình học máy trong việc xử lý bộ dữ liệu tích hợp đa nguồn.

3.4.1. Dữ liệu

Thực nghiệm sử dụng bộ dữ liệu kinh tế vĩ mô được xây dựng ở chương 2. Phạm vi: Trích xuất dữ liệu của Việt Nam thời gian từ 1980 đến 2019. Các chỉ tiêu dự báo: GDP, lạm phát, tỷ lệ thất nghiệp, tỷ giá hối đoái. Dữ liệu sẽ được chia thành hai tập: tập huấn luyện (80% dữ liệu) và tập kiểm tra

(20% dữ liệu) để đánh giá độ chính xác của các mô hình.

3.4.1. Cài đặt và môi trường thực nghiệm

Toàn bộ thực nghiệm được thực hiện trên một máy trạm với cấu hình phần cứng: CPU: i9 9880H; RAM: 32GB; SSD: 512GB; Card đồ họa: Quadro T2000 4GB.

Sau khi chạy thử nghiệm với nhiều giá trị, luận án đã tìm thấy giá trị cài đặt tốt nhất cho bộ dữ liệu thử nghiệm như sau: Linear Regression: không có siêu tham số cần tinh chỉnh; ARIMA(2,1,0); RandomForest: 50 cây, độ sâu 3; XGBoost: 50 cây, max_depth=3, learning_rate=0.1; SVR chọn C=1, $\epsilon=0.1$ với kernel RBF; LSTM: Số đơn vị ẩn: 10, hàm kích hoạt: tanh, số vòng lặp: 200, số lô: 32, bộ tối ưu: adam, hàm mất mát: mse.

3.4.1. Chỉ số đánh giá hiệu suất của mô hình

Để đánh giá sai số của dự báo thực nghiệm sử dụng chỉ số MAE, RMSE, và chỉ số R^2 .

3.4.2. Kết quả và so sánh

Chúng tôi tổng hợp các kết quả MAE, RMSE và R^2 trong Bảng 3.3

Bảng 3. 1 MAE, RMSE, R^2 của mỗi mô hình

Chỉ tiêu	Model	MAE	RMSE	R^2
GDP	Linear Regression	0.82	1.12	0.89
	ARIMA	0.70	0.95	0.92
	Random Forest	0.68	0.90	0.94
	XGBoost	0.72	1.00	0.93
	SVM	1.75	1.05	0.91
	LSTM	0.60	0.77	0.96
Tỷ lệ lạm phát	Linear Regression	0.65	0.88	0.91
	ARIMA	0.54	0.72	0.94
	Random Forest	0.50	0.68	0.95
	XGBoost	0.48	0.66	0.96
	SVM	0.52	0.70	0.95
	LSTM	0.45	0.63	0.97
Tỷ lệ thất nghiệp	Linear Regression	0.48	0.65	0.87
	ARIMA	0.45	0.60	0.90
	Random Forest	0.42	0.58	0.91
	XGBoost	0.38	0.54	0.93

Chỉ tiêu	Model	MAE	RMSE	R ²
	SVM	0.43	0.59	0.91
	LSTM	0.40	0.56	0.92
Tỷ giá hối đoái	Linear Regression	0.55	0.72	0.92
	ARIMA	0.40	0.55	0.96
	Random Forest	0.42	0.60	0.95
	XGBoost	0.41	0.58	0.95
	SVM	0.43	0.61	0.94
	LSTM	0.40	0.57	0.96

Kết quả thực nghiệm trên bốn chỉ tiêu kinh tế vĩ mô của Việt Nam (1980–2019) cho thấy mô hình học sâu và học máy nhìn chung vượt trội so với kinh tế lượng truyền thống, đặc biệt với các biến phi tuyến mạnh như GDP, CPI và tỷ lệ thất nghiệp. LSTM đạt hiệu suất tốt nhất cho GDP (RMSE = 0.77, R² = 0.96), vượt ARIMA và XGBoost tới 18,9%–23% về RMSE. Với CPI, XGBoost và LSTM có R² ≈ 0.96–0.97, vượt trội so với Linear Regression và ARIMA. Ở tỷ lệ thất nghiệp, XGBoost và LSTM giảm sai số 7–15% so với các mô hình truyền thống. Ngược lại, với tỷ giá hối đoái – biến có tính tuyến tính cao – ARIMA vẫn cạnh tranh ngang ngửa với LSTM và XGBoost. SVM cho kết quả trung bình do nhạy cảm với tham số và kernel. Kết quả khẳng định sự cần thiết lựa chọn mô hình theo đặc tính dữ liệu, đồng thời nhấn mạnh tiềm năng của học sâu khi dữ liệu đủ dài và đa dạng.

3.5. Thảo luận về kết quả dự báo GDP và một số chỉ tiêu khác của các mô hình

Kết quả thực nghiệm cho thấy hiệu suất dự báo của ba nhóm mô hình, thống kê truyền thống (hồi quy, ARIMA), học máy truyền thống (SVR, Random Forest, XGBoost), và học sâu (LSTM), có sự khác biệt tùy theo đặc điểm chuỗi thời gian và chỉ tiêu kinh tế. Các mô hình học máy như XGBoost và Random Forest thể hiện hiệu quả tốt trong các chuỗi ổn định, phi tuyến và trung bình độ dài, đặc biệt với khả năng xử lý dữ liệu nhiễu và phân tích tầm quan trọng của biến đầu vào. Mô hình thống kê như ARIMA vẫn có vai trò nền tảng nhờ tính diễn giải rõ ràng, nhưng hạn chế trong việc xử lý phi tuyến và thích ứng với biến động cấu trúc. Trong khi đó, LSTM nổi bật với khả năng học phụ thuộc thời gian dài và mô hình hóa động học phức tạp, đặc biệt hiệu quả trong dự báo GDP. Nhìn chung, không có mô hình nào tối ưu tuyệt đối, song học sâu cho thấy tiềm năng cao trong bối cảnh dữ liệu kinh tế hiện đại giàu tính phi tuyến, nhiễu nhiễu và biến động theo chu kỳ.

3.7. Kết luận chương

Chương 3 đã đánh giá hiệu quả của các nhóm mô hình dự báo: truyền thống (hồi quy, ARIMA), học máy (RF, XGBoost), và học sâu (LSTM) trong bối cảnh dự báo các chỉ tiêu kinh tế tài chính khác nhau. Việc kiểm chứng hiệu suất của các nhóm mô hình kinh tế lượng, học máy và học sâu đã được chúng tôi triển khai trong các công trình [CT1]–[CT4] trên các chỉ tiêu đa dạng khác nhau. Kết quả thực nghiệm cho thấy không có mô hình duy nhất tối ưu cho mọi bài toán kinh tế vĩ mô. Việc lựa chọn mô hình dự báo cần căn cứ vào đặc tính chuỗi thời gian và mục tiêu ứng dụng. Các kết quả này được xem là kết quả trung gian, góp phần chuẩn hóa nền tảng so sánh và làm rõ giới hạn của từng nhóm mô hình. Trên cơ sở đó, chương 3 đặt nền tảng khoa học cho việc phát triển mô hình học sâu thích ứng theo pha chu kỳ, được trình bày chi tiết trong chương 4 nhằm khắc phục những hạn chế đã được nhận diện và nâng cao độ chính xác dự báo trong thực tiễn.

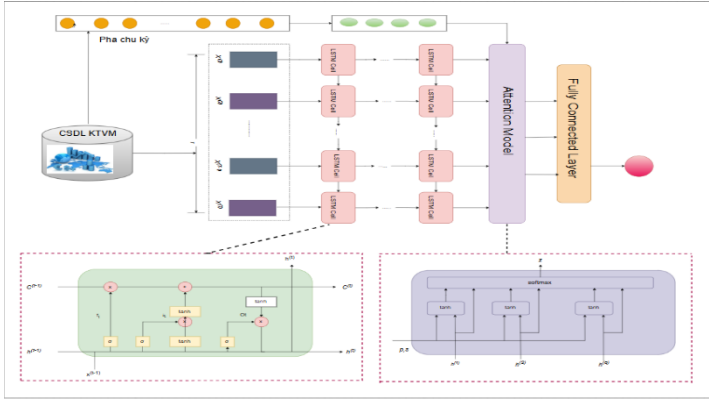
Chương 4. MÔ HÌNH HỌC SÂU VỚI CƠ CHẾ CHÚ Ý THÍCH ỨNG DƯ BẢO GDP QUỐC GIA

4.1. Giới thiệu bài toán

Luận án đề xuất một mô hình mới mang tên PAA-LSTM (Phase-Adaptive Attention Long Short-Term Memory), tích hợp cơ chế Attention thích ứng không chỉ theo thời gian mà còn theo trạng thái chu kỳ nội tại của nền kinh tế. Mô hình này kỳ vọng sẽ không chỉ cải thiện hiệu suất dự báo trong điều kiện kinh tế biến động, mà còn mở ra một hướng tiếp cận mới trong việc kết hợp giữa học sâu và lý thuyết chu kỳ kinh tế trong các hệ thống dự báo vĩ mô hiện đại.

4.2. Đề xuất kiến trúc mô hình

Kiến trúc tổng thể của mô hình PAA-LSTM bao gồm ba thành phần chính: mạng LSTM nhiều tầng, cơ chế attention thích ứng theo pha, và lớp đầu ra kết nối đầy đủ (fully-connected output layer) – minh họa trong Hình 4.1.



Hình 4. 1 Kiến trúc của mô hình PAA-LSTM

4.3. Mô tả kiến trúc và thuật toán

4.3.1. Phân đoạn chu kỳ kinh tế

Để xác định các pha của chu kỳ kinh tế (bao gồm suy thoái, phục hồi, tăng trưởng và tăng trưởng mạnh), luận án phân pha dựa trên thuật toán Bry–Boschan (BB) [33]. Thuật toán này được phát triển bởi Bry và Boschan (1971) trong khuôn khổ nghiên cứu của Cục Nghiên cứu Kinh tế Quốc gia Hoa Kỳ (NBER), nhằm xác định các điểm đảo chiều trong chuỗi thời gian kinh tế theo cách tiếp cận phi tham số và không yêu cầu chuỗi dữ liệu phải dừng với khả năng áp dụng cho các chuỗi thời gian có tần suất thấp như quý hoặc năm.

4.3.2. Xây dựng cơ chế ánh xạ trọng số attention riêng biệt cho từng pha

Các bước chi tiết tính toán các trọng số attention đặc thù theo pha như sau:

Thuật toán PAA-LSTM

Input:

$H = [h_1, h_2, \dots, h_t]$ // Trạng thái ẩn của LSTM

$P = [p_1, p_2, \dots, p_t]$ // Nhãn pha chu kỳ tại mỗi thời điểm

$S = [s_1, s_2, \dots, s_t]$ // Trạng thái hiện tại tại mỗi thời điểm LSTM layer cuối

Tham số mạng

W_h, W_s, W_p : ma trận trọng số cho encoder hidden states, decoder state, vector pha

$v^{(p)}$: vector trọng số cho attention

b : bias vector

Output:

Vecto ngữ cảnh c_t

1: Initialize attention weights $\alpha = []$

2: for $t = 1$ to T do

3: Identify $phase_t \leftarrow P[t]$

4: Retrieve parameters:

$W_h \leftarrow W_h^{(p)}[phase_t]$

$W_s \leftarrow W_s^{(p)}[phase_t]$

$W_p \leftarrow W_p[phase_t]$

$v \leftarrow v[phase_t]$

$b \leftarrow b^p[phase_t]$

5: Compute energy: e_t

6: Compute attention score: $\alpha_t \leftarrow \text{softmax}(e_t)$

7: Append α_t to α

8: end for

9: Compute context vector: $c_t = \sum (\alpha_t \cdot h_t)$

10: Return c_t

4.4. Thực nghiệm

4.4.1. Dữ liệu và phạm vi

Thực nghiệm sử dụng bộ dữ liệu kinh tế vĩ mô được xây dựng ở chương 2. Phạm vi: Trích xuất dữ liệu sáu quốc gia, đại diện cho ba nhóm nền kinh tế khác nhau: Trung Quốc, Nga, Việt Nam, Ấn Độ; Hoa Kỳ, Canada. Với mỗi quốc gia, dữ liệu được chia thành tập huấn luyện và tập kiểm tra theo trình tự thời gian. Cụ thể, 80% quan sát sớm nhất được sử dụng để huấn luyện, 20% còn lại dành cho kiểm tra. Để phục vụ cho huấn luyện mô hình PAA-LSTM, dữ liệu được chuyển đổi thành tập các phân đoạn chuỗi thời gian có độ dài cố định theo cơ chế cửa sổ trượt (sliding window). Mỗi mẫu huấn luyện bao gồm một đoạn dữ liệu đầu vào liên tiếp có độ dài $T_{input} = 5$.

4.4.2. Cài đặt và môi trường thực nghiệm

Toàn bộ thực nghiệm được thực hiện trên một máy trạm chuyên dụng với cấu hình phần cứng như sau: CPU: i9 9880H; RAM: 32GB; SSD: 512GB; Card đồ họa: Quadro T2000 4GB.

Môi trường phần mềm sử dụng bao gồm: Python 3.10, TensorFlow 2.13, Keras, cùng với các thư viện hỗ trợ như Scikit-learn (v1.3), NumPy (v1.24) và Matplotlib (v3.7) [36].

Tham số: Số đơn vị ẩn: 10; hàm kích hoạt: Relu, số vòng lặp: 500; số lô: 64, bộ tối ưu: Adam, hàm mất mát: mse.

4.4.3. Chỉ số đánh giá hiệu suất của mô hình

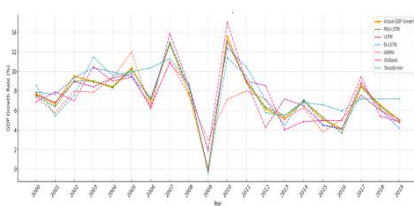
RMSE được lựa chọn vì khả năng phạt mạnh các sai số lớn, điều đặc biệt quan trọng trong dự báo kinh tế vĩ mô, nơi mà các sai lệch lớn có thể dẫn đến hậu quả nghiêm trọng trong hoạch định chính sách. Ngoài ra để đảm bảo tính toàn diện luận án sử dụng chỉ số MAE theo công thức (1.8) và chỉ số R^2 theo công thức (1.11).

4.4.4. Độ phức tạp tính toán của mô hình PAA-LSTM

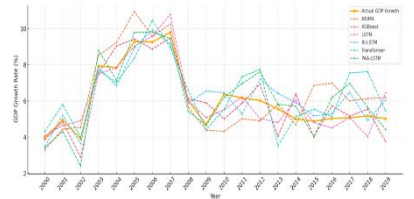
Với số đặc trưng đầu vào d , số đơn vị ẩn h , độ dài chuỗi T , số lớp L và kích thước lô B , độ phức tạp tính toán của PAA-LSTM biểu diễn như sau:

$$O(BTL(dh + h^2 + h)) \quad (4.9)$$

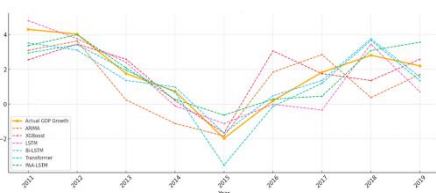
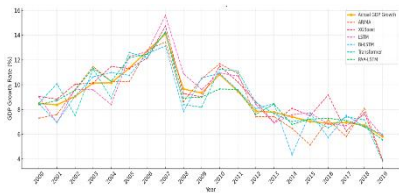
4.4.5. Kết quả thực nghiệm và so sánh



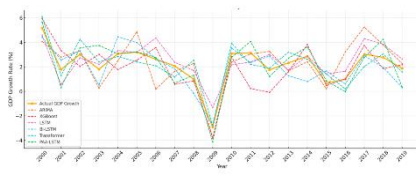
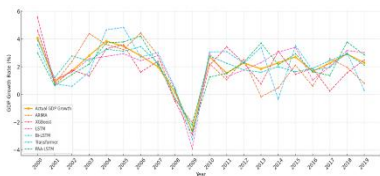
Hình 4.6 Kết quả dự báo GDP của VN



Hình 4.8 Kết quả dự báo GDP của Ấn Độ



Hình 4. 10-12 Kết quả dự báo GDP của Trung Quốc và Nga

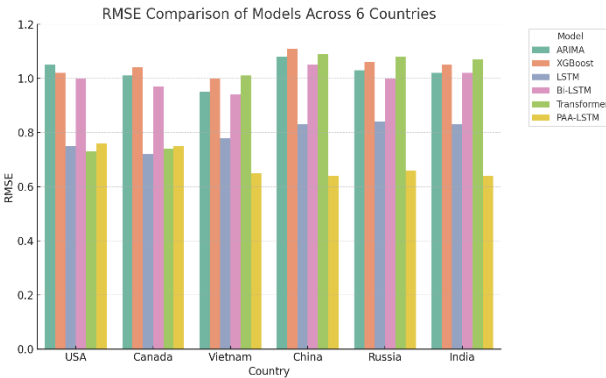


Hình 4. 14-16 Kết quả dự báo GDP của Mỹ và Canada

4.5. Phân tích hiệu năng và ý nghĩa thực tiễn

Thực nghiệm trên sáu quốc gia cho thấy mô hình PAA-LSTM đạt hiệu suất cao nhất tại các nền kinh tế biến động mạnh (Việt Nam, Trung Quốc, Ấn Độ), nhờ khả năng tích hợp thông tin chu kỳ kinh tế vào cơ chế attention. So với các mô hình như ARIMA, XGBoost hay Transformer, PAA-LSTM cho RMSE thấp hơn rõ rệt và diễn giải tốt hơn trong các pha chuyển tiếp kinh tế.

Dù không vượt trội tại các nền kinh tế ổn định như Mỹ, Canada, PAA-LSTM vẫn thể hiện độ chính xác cạnh tranh, tuy hiệu quả giảm dần so với mô hình đơn giản hơn. So sánh với các nghiên cứu trước, kết quả xác nhận rằng attention theo pha giúp mô hình hóa tốt hơn các đặc điểm động của chuỗi thời gian vĩ mô, đồng thời giảm độ phức tạp suy diễn so với các mô hình chuyển chế độ cổ điển.



Hình 4. 17 So sánh RMSE của các mô hình với 6 Quốc Gia

Chương này khẳng định lợi ích của chiến lược chọn mô hình thích ứng theo đặc điểm chu kỳ kinh tế, mở rộng ứng dụng học sâu trong dự báo GDP Quốc Gia.

4.6. Kết luận chương 4

Chương này đã đề xuất và kiểm chứng hiệu quả của mô hình PAA-LSTM, một kiến trúc học sâu kết hợp giữa mạng nơ-ron LSTM nhiều tầng và cơ chế chú ý thích ứng theo pha nhằm nâng cao năng lực đối với dự báo GDP Quốc Gia. Khác biệt trọng yếu của mô hình này so với các kiến trúc trong các nghiên cứu khác nằm ở khả năng tự động điều chỉnh trọng số chú ý theo từng giai đoạn của chu kỳ kinh tế (suy thoái, phục hồi, tăng trưởng,

tăng trưởng mạnh), từ đó giúp mô hình thích ứng tốt hơn với các đặc trưng động của dữ liệu kinh tế vĩ mô.

Kết quả thực nghiệm trên sáu quốc gia đại diện cho ba nền kinh tế khác nhau là phát triển, mới nổi và đang phát triển đã chứng minh tính ưu việt của mô hình đề xuất. PAA-LSTM liên tục đạt chỉ số RMSE thấp nhất trong các nền kinh tế có độ biến động cao như Việt Nam, Trung Quốc và Ấn Độ, đồng thời vẫn duy trì hiệu năng cạnh tranh và khả năng giải thích tốt tại các quốc gia phát triển như Hoa Kỳ và Canada. So sánh với các mô hình truyền thống như ARIMA, XGBoost và các kiến trúc học sâu hiện đại như Transformer, Bi-LSTM cho thấy rằng PAA-LSTM có khả năng mô hình hóa vượt trội, đặc biệt trong việc nắm bắt mối quan hệ phi tuyến và phụ thuộc dài hạn trong chuỗi thời gian kinh tế. Kết quả nghiên cứu của chương này đã được công bố trong công trình [CT5].

KẾT LUẬN VÀ KIẾN NGHỊ

Kết quả đạt được

Luận án “Nghiên cứu xây dựng mô hình dự báo GDP và một số chỉ tiêu kinh tế vĩ mô khác dựa trên dữ liệu đa nguồn” đã tập trung giải quyết bài toán dự báo kinh tế vĩ mô trong bối cảnh dữ liệu ngày càng phức tạp, phi tuyến và biến động mạnh theo thời gian. Trên nền tảng lý luận kinh tế học kết hợp với các kỹ thuật hiện đại của học máy và học sâu, luận án đã hoàn thiện khung phương pháp luận và thực nghiệm hệ thống nhằm nâng cao hiệu quả dự báo trong các môi trường kinh tế đa dạng.

Những kết quả chính đạt được:

Thứ nhất, luận án đã xây dựng được bộ dữ liệu kinh tế vĩ mô đa nguồn, tích hợp từ các tổ chức uy tín như World Bank, PWT, GSO và Fiintech. Dữ liệu được xử lý đảm bảo khả năng học tốt cho các mô hình dự báo chuỗi thời gian.

Thứ hai, nghiên cứu đã triển khai và đánh giá hệ thống các mô hình dự báo kinh tế vĩ mô từ truyền thống đến hiện đại. Kết quả thực nghiệm cho thấy các mô hình học sâu vượt trội hơn trong việc mô hình hóa quan hệ phi tuyến và ghi nhớ phụ thuộc dài hạn.

Thứ ba, luận án đề xuất mô hình học sâu mới có tên PAA-LSTM (Phase-Adaptive Attention LSTM) kết hợp LSTM nhiều tầng với cơ chế chú ý thích ứng theo pha chu kỳ kinh tế để dự báo GDP Quốc Gia. Mô hình cho phép học trọng số attention riêng biệt theo từng giai đoạn (suy thoái, phục hồi, tăng trưởng, tăng trưởng mạnh), từ đó phản ánh sát hơn đặc trưng động

học của nền kinh tế. Mô hình được kiểm chứng thực nghiệm trên dữ liệu GDP của các quốc gia thuộc các nhóm thể chế khác nhau khẳng định tính hiệu quả và khả năng mở rộng của mô hình đề xuất.

Định hướng phát triển

Trên cơ sở các kết quả đạt được, luận án đề xuất một số hướng nghiên cứu và ứng dụng tiếp theo:

Thứ nhất, mở rộng mô hình sang các chỉ tiêu kinh tế khác như lạm phát, cán cân thương mại, tỷ lệ thất nghiệp, nhằm kiểm chứng tính linh hoạt và độ khái quát hóa của mô hình PAA-LSTM.

Thứ hai, tích hợp dữ liệu thời gian thực và phi truyền thống, bao gồm tín hiệu từ Google Trends, mạng xã hội, dữ liệu vệ tinh, để kiểm chứng khả năng thích ứng nhanh của mô hình trước các biến động thị trường.

Thứ ba, kết hợp với các phương pháp giải thích mô hình như SHAP,... để tăng cường tính minh bạch của mô hình và hỗ trợ ra quyết định cho các nhà hoạch định chính sách.

Tổng kết lại, với việc đề xuất mô hình học sâu thích ứng theo pha chu kỳ PAA-LSTM, luận án bước đầu cho thấy việc tích hợp kiến thức kinh tế học – cụ thể là đặc điểm chu kỳ kinh tế – vào thiết kế kiến trúc học sâu có tiềm năng góp phần nâng cao độ chính xác và tính ổn định trong dự báo các chỉ tiêu kinh tế vĩ mô. Mặc dù còn cần được kiểm chứng thêm trong các bối cảnh và phạm vi dữ liệu rộng hơn, hướng tiếp cận này cho thấy tiềm năng ứng dụng tại các nền kinh tế có cấu trúc dữ liệu không ổn định và độ trễ chính sách cao, như Việt Nam.

DANH MỤC CÁC CÔNG TRÌNH KHOA HỌC CỦA TÁC GIẢ LIÊN QUAN ĐẾN LUẬN ÁN

- [CT1] **[Lan Dong Thi Ngoc**, Khai Phan Van, Ngo-Thi-Thu-Trang, Gyoo Seok Choi, Ha-Nam Nguyen, "A Multiple Variable Regression-based Approaches to Long-term Electricity Demand Forecasting", International Journal of Advanced Smart Convergence, Vol.10 No.4, page 59-65, 2021
- [CT2] **Lan Dong Thi Ngoc**, Ha-Nam Nguyen, Khai Phan Van, Nam Nguyen Hoai, Hoa Tran Xuan, Dung Hoang Van, GyooSeokChoi, "Research and application of multi-regression analysis method to forecast the electricity consumption demand for Vietnam's

residential sector", 9th International Symposium, ISAAC 2021 in Conjunction with ICACT, page 99-102, 2021

- [CT3] **Lan Dong Thi Ngoc**, Duy-Linh Bui, Sang Van Ha, Huong Tran Thi, Viet Pham Minh, Ha-Nam Nguyen, "Using Machine Learning to Improve Forecasting Efficiency for the Stock Market", International Conference on Research in Management & Technovation, page 439-447, Springer Nature Singapore, 2023.
- [CT4] **Lan Dong Thi Ngoc** , Ngo-Thi-Thu-Trang, Kyoung-Mok Kim, Vung Pham and Ha-Nam Nguyen, "Study of Machine Learning Models for Forecasting Macroeconomic Indicators," 25th International Conference on Advanced Communication Technology (ICACTION), 2025 (Hội nghị trong danh mục scopus)
- [CT5] **Lan Dong Thi Ngoc**, N. D. Hoan, and Ha.-Nam Nguyen, "Gross Domestic Product Forecasting Using Deep Learning Models with a Phase-Adaptive Attention Mechanism," Electronics, vol. 14, no. 11, Art. no. 2132, May 2025. (Tạp chí thuộc danh mục trong Scopus, SCIE (WoS) Q2)