

ĐẠI HỌC QUỐC GIA HÀ NỘI
TRƯỜNG 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**

LUẬN ÁN TIẾN SĨ HỆ THỐNG THÔNG TIN

Hà Nội - 2025

ĐẠI HỌC QUỐC GIA HÀ NỘI
TRƯỜNG 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

LUẬN ÁN TIẾN SĨ HỆ THỐNG THÔNG TIN

NGƯỜI HƯỚNG DẪN KHOA HỌC:

- 1. PGS. TS NGUYỄN HÀ NAM**
- 2. TS. NGUYỄN THỊ HẬU**

Hà Nội - 2025

LỜI CAM ĐOAN

Tôi xin cam đoan Luận án Tiến sĩ, với tiêu đề “Xây dựng mô hình dự báo GDP và một số chỉ tiêu kinh tế vĩ mô khác dựa trên dữ liệu đa nguồn” là công trình nghiên cứu khoa học của tôi, được thực hiện dưới sự hướng dẫn của PGS.TS. Nguyễn Hà Nam và TS. Nguyễn Thị Hậu. Nội dung của luận án không có sự sao chép và không được sử dụng để bảo vệ cho bất kỳ học vị nào khác.

Các kết quả được viết chung với các tác giả khác đều được sự đồng ý của các đồng tác giả trước khi đưa vào luận án. Luận án này được hoàn thành trong thời gian tôi làm Nghiên cứu sinh tại Đại học Công nghệ, Đại học Quốc gia Hà Nội.

NGƯỜI CAM ĐOAN

Nghiên cứu sinh

Đỗ Thị Ngọc Lan

LỜI CẢM ƠN

Thời gian học tập, nghiên cứu sinh và thực hiện luận án tại Bộ môn Các Hệ thống thông tin - Khoa Công nghệ thông tin - Trường Đại học Công nghệ - Đại học Quốc gia Hà Nội dưới sự hướng dẫn khoa học tận tâm của PGS.TS. Nguyễn Hà Nam và TS. Nguyễn Thị Hậu là một giai đoạn có ý nghĩa quan trọng trong quá trình hình thành năng lực nghiên cứu của tôi.

Trước hết, tôi xin bày tỏ lòng biết ơn sâu sắc đến PGS.TS. Nguyễn Hà Nam và TS. Nguyễn Thị Hậu, người đã tận tình hướng dẫn, truyền đạt kiến thức và luôn đồng hành cùng tôi trong suốt quá trình thực hiện luận án. Sự tận tâm, trách nhiệm và những ý kiến đóng góp quý báu của Thầy/Cô là nền tảng quan trọng giúp tôi hoàn thiện công trình nghiên cứu này.

Tôi xin chân thành cảm ơn các thầy cô Trường Đại học Công nghệ, Đại học Quốc Gia Hà Nội, đặc biệt là các thầy cô trong bộ môn Các hệ thống thông tin, các thầy cô trong Ban chủ nhiệm Khoa, văn phòng Khoa Công nghệ thông tin, phòng Đào tạo và các phòng ban chức năng đã tạo điều kiện thuận lợi, cung cấp môi trường học thuật chuyên nghiệp với nhiều cơ hội học hỏi, nghiên cứu trong suốt thời gian tôi học tập và thực hiện luận án.

Tôi cũng bày tỏ lòng cảm ơn sâu sắc tới Ban giám đốc Học viện Tài Chính; Ban lãnh đạo Khoa Hệ thống thông tin kinh tế, Bộ môn Tin học TCKT đã tạo mọi điều kiện thuận lợi cho tôi trong quá trình nghiên cứu; Cảm ơn các đồng nghiệp trong Khoa, bộ môn đã luôn ủng hộ, quan tâm và động viên tôi.

Tôi cũng xin gửi lời cảm ơn đến các anh chị em, bạn bè, các nghiên cứu sinh và các anh chị em nhóm nghiên cứu của thầy PGS.TS. Nguyễn Hà Nam đã hỗ trợ, động viên và chia sẻ những kinh nghiệm quý báu trong suốt quá trình nghiên cứu.

Đặc biệt, tôi xin tri ân gia đình thân yêu, cha mẹ, người bạn đời và các con, vì đã luôn bên cạnh, ủng hộ và tiếp thêm sức mạnh tinh thần cho tôi vượt qua mọi khó khăn, thử thách trong hành trình học thuật đầy gian nan nhưng cũng rất đáng tự hào này.

Hà Nội, ngày 25 tháng 06 năm 2025

NGHIÊN CỨU SINH

Đỗ Thị Ngọc Lan

MỤC LỤC

LỜI CAM ĐOAN	I
LỜI CẢM ƠN II	
DANH MỤC KÝ HIỆU VÀ CHỮ VIẾT TẮT	VI
DANH MỤC BẢNG BIỂU	IX
DANH MỤC HÌNH VẼ, ĐỒ THỊ	XI
MỞ ĐẦU	1
Lý do lựa chọn đề tài và mục tiêu nghiên cứu	1
Đối tượng và phạm vi nghiên cứu	3
Phương pháp tiếp cận	4
Phương pháp nghiên cứu	4
Ý nghĩa khoa học và thực tiễn	5
Đóng góp của luận án	6
Cấu trúc của luận án	7
CHƯƠNG 1. TỔNG QUAN NGHIÊN CỨU	9
1.1. Một số chỉ tiêu kinh tế vĩ mô tiêu biểu	9
1.2. Lý thuyết dự báo kinh tế vĩ mô – đặc điểm chuỗi thời gian	15
1.2.1. Khái niệm và vai trò của chuỗi thời gian trong kinh tế	15
1.2.2. Đặc điểm của chuỗi thời gian kinh tế vĩ mô	16
1.2.3. Phân tích và xử lý chuỗi thời gian	17
1.3. Một số mô hình thông dụng trong dự báo KTVM	18
1.3.1. Các mô hình hồi quy	18
1.3.2. Mô hình tự hồi quy tích hợp trung bình trượt	20
1.3.3. Mô hình tự hồi quy vector	22
1.4. Chỉ số đánh giá hiệu suất của mô hình	24
1.5. Tình hình nghiên cứu trong và ngoài nước về dự báo GDP và một số CTKTVM khác	26
1.5.1. Tình hình nghiên cứu trên thế giới	26
1.5.2. Tình hình nghiên cứu ở Việt Nam	33

1.6. Khoảng trống nghiên cứu và đề xuất giải pháp	39
1.6.1. Khoảng trống nghiên cứu	39
1.6.2. Đề xuất giải pháp	40
1.7. Kết luận chương	41
CHƯƠNG 2. XÂY DỰNG BỘ DỮ LIỆU KINH TẾ VĨ MÔ ĐA NGUỒN	42
2.1. Giới thiệu	42
2.2. Quy trình xây dựng bộ dữ liệu kinh tế vĩ mô đa nguồn	44
2.3. Các CTKTVM sử dụng trong nghiên cứu	47
2.4. Nguồn dữ liệu	49
2.5. Hợp nhất và chuẩn hóa dữ liệu	52
2.6. Kiểm tra chất lượng và chuẩn bị dữ liệu	60
2.6.1. Phân tích đặc điểm bộ dữ liệu	60
2.6.2. Phân tích thống kê mô tả	61
2.6.3. Phát hiện ngoại lệ và xử lý dữ liệu thiếu	65
2.6.4. Kiểm định chuỗi thời gian	66
2.6.5. Chuyển đổi dữ liệu	69
2.7. Kết luận chương	70
CHƯƠNG 3. CÁC MÔ HÌNH DỰ BÁO CHỈ TIÊU KINH TẾ VĨ MÔ	71
3.1. Mô hình kinh tế lượng	71
3.2. Mô hình học máy	76
3.3. Mô hình học sâu	83
3.4. Thực nghiệm dự báo GDP và một số CTKTVM khác	88
3.4.1. Dữ liệu	88
3.4.2. Cài đặt và môi trường thực nghiệm	89
3.4.3. Chỉ số đánh giá hiệu suất của mô hình	89
3.4.4. Kết quả và so sánh	90
3.5. Thảo luận về kết quả dự báo GDP và một số chỉ tiêu khác của các mô hình	94
3.6. Kết luận chương	95

CHƯƠNG 4. MÔ HÌNH HỌC SÂU VỚI CƠ CHẾ CHÚ Ý THÍCH ỨNG DỤNG	
BẢO GDP QUỐC GIA	97
4.1. Giới thiệu bài toán	97
4.1.1. Đặt vấn đề	97
4.1.2. Hướng tiếp cận giải quyết vấn đề	99
4.2. Đề xuất kiến trúc mô hình	102
4.3. Mô tả kiến trúc đề xuất	104
4.3.1. Phân đoạn chu kỳ kinh tế	105
4.3.2. Xây dựng cơ chế ánh xạ trọng số attention riêng biệt cho từng pha	107
4.4. Độ phức tạp tính toán của mô hình PAA-LSTM	112
4.5. Thực nghiệm	113
4.5.1. Dữ liệu và Phạm vi	113
4.5.2. Cài đặt môi trường thực nghiệm	114
4.5.3. Chỉ số đánh giá hiệu suất của mô hình	116
4.5.4. Kết quả thực nghiệm và so sánh	116
4.6. Phân tích hiệu năng và ý nghĩa thực tiễn	125
4.7. Kết luận chương	128
KẾT LUẬN	130
Kết quả đạt được	130
Định hướng phát triển	130
DANH MỤC CÁC CÔNG TRÌNH KHOA HỌC CỦA TÁC GIẢ LIÊN QUAN	
ĐẾN LUẬN ÁN	132
TÀI LIỆU THAM KHẢO	133
PHỤ LỤC BỘ DỮ LIỆU KTVM	145

DANH MỤC KÝ HIỆU VÀ CHỮ VIẾT TẮT

Từ viết tắt	Diễn giải	
	Tiếng Anh	Tiếng Việt
ADC	Analog-to-Digital Converter	Bộ chuyển tín hiệu tương tự
ADF	Augmented Dickey-Fuller	Kiểm định tính dừng ADF
AI	Artificial Intelligence	Trí tuệ nhân tạo
ANN	Artificial Neural Network	Mạng nơ-ron
API	Application Programming Interface	Giao diện lập trình ứng dụng
BD	Big data	Dữ liệu lớn
BEA	Bureau of Economic Analysis	Cục phân tích kinh tế Hoa Kỳ
BLUE	Best Linear Unbiased Estimator	Ước lượng tuyến tính tốt nhất và không chệch
BLS	Bureau of Labor Statistics	Cục Thống kê Lao động Hoa Kỳ
CTKTVM	Macroeconomic Indicators	Chỉ tiêu kinh tế vĩ mô
CPI	Consumer Price Index	Chỉ số giá tiêu dùng
DAS	Data Acquisition Systems	Hệ thống thu thập dữ liệu tự động
DL	Deep learning	Học sâu
DRL	Deep reinforcement learning	Học tăng cường sâu
ECB	European Central Bank	Ngân hàng Trung ương châu Âu
FDI	Foreign Direct Investment	Đầu tư trực tiếp nước ngoài
FED	Federal Reserve	Cục dự trữ liên bang Hoa Kỳ
FRED	Federal Reserve Economic Data	Dữ liệu kinh tế của cục dự trữ liên bang Hoa Kỳ
GDP	Gross Domestic Product	Tổng sản phẩm Quốc nội
GSO	General Statistics Office of Vietnam	Cục Thống kê Việt Nam

Từ viết tắt	Diễn giải	
	Tiếng Anh	Tiếng Việt
IPI	Industrial Production Index	Chỉ số sản xuất công nghiệp
IMF	International Monetary Fund	Quỹ tiền tệ Quốc tế
KPSS	Kwiatkowski–Phillips–Schmidt–Shin	Kiểm định KPSS
KTVM	Macroeconomic	Kinh tế vĩ mô
LSTM	Long Short-Term Memory	Mạng nơ-ron bộ nhớ ngắn dài hạn
ML	Machine Learning	Học máy
MPI	Ministry of Planning and Investment	Bộ kế hoạch và đầu tư Việt Nam (Nay Bộ này sáp nhập, hợp nhất với Bộ Tài chính)
MAR	Missing at Random	Thiếu dữ liệu ngẫu nhiên
MCAR	Missing Completely at Random	Thiếu dữ liệu hoàn toàn ngẫu nhiên
MNAR	Missing Not at Random	Thiếu dữ liệu không ngẫu nhiên
NLP	Natural Language Processing	Xử lý ngôn ngữ tự nhiên
OECD	Organisation for Economic Co-operation and Development	Tổ chức Hợp tác và Phát triển Kinh tế
OLS	Ordinary Least Squares	Bình phương tối thiểu thông thường
PWT	Penn World Table	Bảng dữ liệu kinh tế thế giới
RF	Random Forest	Rừng ngẫu nhiên
RMSE	Root Mean Squared Error	Sai số quân phương trung bình
SBV	State Bank of Vietnam	Ngân hàng nhà nước Việt Nam
SMOTE	Synthetic Minority Over-sampling Technique	Kỹ thuật tăng mẫu thiểu số tổng hợp
SNA	System of National Accounts	Hệ thống tài khoản quốc gia
SVM	Support Vector Machine	Máy vector tựa

Từ viết tắt	Diễn giải	
	Tiếng Anh	Tiếng Việt
TVP	Time-Varying Parameter	Mô hình hệ số thay đổi theo thời gian
WB	World Bank	Ngân hàng thế giới
XGBoost	Extreme Gradient Boosting	Thuật toán tăng cường độ dốc cực đại

DANH MỤC BẢNG BIỂU

Bảng 2. 1 Các chỉ tiêu kinh tế vĩ mô cần thu thập-----	48
Bảng 2. 2 Tổng hợp các chỉ tiêu KTVM sau hợp nhất -----	56
Bảng 2. 3 Bảng tham số thống kê các chỉ số KTVM Việt Nam-----	63
Bảng 2. 4 Các ngoại lệ -----	65
Bảng 2. 5 Bảng kiểm định tính dừng của một số biến KTVM -----	67
Bảng 3. 1 So sánh các mô hình học máy trong dự báo kinh tế vĩ mô -----	81
Bảng 3. 2 So sánh các mô hình học sâu trong dự báo kinh tế vĩ mô-----	86
Bảng 3. 3 MAE, RMSE, R^2 của mỗi mô hình -----	90
Bảng 4. 1 Kết quả của điều chỉnh tham số trong mô hình theo tập dữ liệu -----	115
Bảng 4. 2 Kết quả dự báo GDP của Việt Nam -----	117
Bảng 4. 3 Kết quả dự báo GDP của Ấn Độ -----	118
Bảng 4. 4 Kết quả dự báo GDP của Trung Quốc-----	120
Bảng 4. 5 Kết quả dự báo GDP của Nga -----	121
Bảng 4. 6 Kết quả dự báo GDP của Mỹ -----	122
Bảng 4. 7 Kết quả dự báo GDP của Canada -----	124
Bảng 4. 8 Bảng RMSE của các mô hình với 6 quốc gia-----	125
Bảng PL. 1 Các chỉ tiêu đo lường quy mô nền kinh tế, lực lượng lao động và năng suất theo thời gian -----	145
Bảng PL. 2 Các chỉ tiêu đo lường sản lượng và tích lũy vốn -----	146
Bảng PL. 3 Các chỉ tiêu dựa trên tài khoản quốc gia -----	147
Bảng PL. 4 Các chỉ tiêu điều chỉnh sức mua-----	148

Bảng PL. 5 Nhóm các chỉ tiêu mức giá theo PPP của các cấu phần chi tiêu và vốn -----	148
Bảng PL. 6 Các biến tỷ trọng trong GDP theo PPP-----	149
Bảng PL. 7 Các chỉ tiêu giá tiêu dùng và lạm phát-----	150

DANH MỤC HÌNH VẼ, ĐỒ THỊ

Hình 1. 1 Các bước phân tích và xử lý dữ liệu chuỗi thời gian-----	17
Hình 2. 1 Quy trình xây dựng bộ dữ liệu KTVM đa nguồn -----	45
Hình 2. 2 Các nguồn dữ liệu -----	50
Hình 2. 3 Quy trình hợp nhất dữ liệu -----	52
Hình 2. 4 Tỷ lệ thiếu dữ liệu giai đoạn 1950-2019-----	54
Hình 2. 5 Tỷ lệ thiếu dữ liệu giai đoạn 1980-2019-----	55
Hình 3. 1 MAE của các mô hình -----	91
Hình 3. 2 RMSE của các mô hình -----	92
Hình 3. 3 R^2 của các mô hình -----	92
Hình 4. 1 Kiến trúc của mô hình PAA-LSTM-----	104
Hình 4. 2 Phân pha chu kỳ kinh tế-----	106
Hình 4. 3 Kiến trúc Attention thích ứng theo pha (p)-----	107
Hình 4. 4 Giải mã của thuật toán attention trong PAA-LSTM-----	110
Hình 4. 5 So sánh RMSE của các mô hình khi dự báo GDP của Việt Nam-----	117
Hình 4. 6 Kết quả dự báo GDP của Việt Nam-----	117
Hình 4. 7 So sánh RMSE của các mô hình khi dự báo GDP của Ấn Độ-----	119
Hình 4. 8 Kết quả dự báo GDP của Ấn Độ -----	119
Hình 4. 9 So sánh RMSE của các mô hình khi dự báo GDP của Trung Quốc -----	120
Hình 4. 10 Kết quả dự báo GDP của Trung Quốc -----	120
Hình 4. 11 So sánh RMSE của các mô hình khi dự báo GDP của Nga -----	121
Hình 4. 12 Kết quả dự báo GDP của Nga -----	122
Hình 4. 13 So sánh RMSE của các mô hình khi dự báo GDP của Mỹ-----	123

Hình 4. 14 Kết quả dự báo GDP của Mỹ -----	123
Hình 4. 15 So sánh RMSE của các mô hình khi dự báo GDP của Canada -----	124
Hình 4. 16 Kết quả dự báo GDP của Canada. -----	124
Hình 4. 17 So sánh RMSE của các mô hình với 6 quốc gia -----	126

MỞ ĐẦU

Lý do lựa chọn đề tài và mục tiêu nghiên cứu

Dự báo các chỉ tiêu kinh tế vĩ mô (CTKTVM) đóng vai trò quan trọng trong nghiên cứu và hoạch định chính sách kinh tế ở hầu hết các quốc gia. 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]. 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 CTKTVM 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 CTKTVM 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]. Những cải tiến như mô hình DSGE, Markov-switching hay mô hình hệ số thay đổi theo thời gian (TVP) đã được đưa ra nhằm mô tả tốt hơn động học chuỗi thời gian, nhưng vẫn gặp trở ngại trong việc xử lý dữ liệu có số chiều cao và không có cấu trúc [80].

Sự phát triển mạnh mẽ của trí tuệ nhân tạo và các kỹ thuật học máy (ML) đã mở ra những hướng tiếp cận mới trong lĩnh vực dự báo kinh tế. 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.

Gần đây, học tăng cường sâu (Deep Reinforcement Learning – DRL) cũng đang dần được ứng dụng trong kinh tế học, đặc biệt trong các bài toán ra quyết định và điều khiển chính sách kinh tế vĩ mô [17]. Mặc dù còn mới mẻ, nhưng các mô hình như DQN, PPO đang cho thấy tiềm năng mạnh mẽ trong việc thích nghi với dữ liệu thay đổi theo thời gian và môi trường không chắc chắn. Tuy nhiên, việc áp dụng các mô hình này trong kinh tế vĩ mô đặt ra nhiều thách thức về tính giải thích, khả năng tổng quát hóa và yêu cầu dữ liệu đầu vào phù hợp [17].

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. Theo khảo sát của nghiên cứu sinh thì việc phân pha chu kỳ kinh tế mới được nghiên cứu trong lý thuyết kinh tế học và ứng dụng trong các mô hình kinh tế lượng mà chưa được nghiên cứu đưa vào mô hình học máy, bên cạnh đó việc sử dụng các mô hình học máy hiện đại vào dự báo CTKTVM cũng còn rất hạn chế.

Từ những vấn đề đã phân tích, có thể khẳng định sự cần thiết của việc phát triển một mô hình học máy hiện đại, có khả năng mô hình hóa các mối quan hệ phức tạp, thích ứng với sự biến động của chu kỳ kinh tế và đồng thời đảm bảo hiệu quả trong dự báo trên dữ liệu kinh tế vĩ mô. Tuy nhiên, quá trình này đối mặt với nhiều thách thức, bao gồm: (i) sự thiếu hụt của một bộ dữ liệu kinh tế vĩ mô có tính nhất quán ảnh hưởng đến chất lượng huấn luyện mô hình; (ii) hạn chế trong việc tích hợp thông tin pha chu kỳ kinh tế vào kiến trúc mô hình học sâu nhằm khai thác đặc trưng động học thay đổi theo thời gian; và (iii) thiếu cơ chế đánh giá hệ thống hiệu quả của các mô hình dự báo trong các bối cảnh kinh tế khác nhau.

Trên cơ sở đó, 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. Cách tiếp cận này kết hợp giá trị khoa học và tiềm năng ứng dụng thực tiễn trong bối cảnh kinh tế toàn cầu biến động mạnh mẽ.

Đối tượng và phạm vi nghiên cứu

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 quốc gia tiêu biểu thuộc các nền kinh tế khác nhau bao gồm: nhóm các nước phát triển (G7), nhóm các nước mới nổi (BRICS) và nhóm 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 tiếp cận

Nghiên cứu áp dụng các hướng tiếp cận sau:

1. **Tiếp cận định hướng dữ liệu:** Nghiên cứu ưu tiên việc khai thác các đặc trưng tự sinh từ dữ liệu. Các mô hình học máy sẽ học trực tiếp từ dữ liệu đầu vào, không áp đặt cấu trúc mô hình trước, giúp khám phá các quan hệ phi tuyến và ẩn trong dữ liệu.
2. **Tiếp cận học máy:** Nghiên cứu các mô hình hồi quy, cây quyết định, boosting, mạng nơ-ron sâu, v.v. để xây dựng mô hình dự báo có khả năng học từ dữ liệu lịch sử.
3. **Tiếp cận so sánh – thực nghiệm:** So sánh hiệu năng dự báo giữa mô hình đề xuất và các mô hình khác như ARIMA, VAR, Randomforest, XGBoost, SVM, LSTM, Bi-LSTM, v.v. trên tập dữ liệu thực tế của nhiều quốc gia để kiểm chứng độ chính xác, tính khái quát và tính ổn định của mô hình.

Phương pháp nghiên cứu

Để đạt được các mục tiêu nghiên cứu đã đề ra, luận án sử dụng kết hợp các phương pháp nghiên cứu định tính và định lượng theo hướng thực nghiệm, bao gồm:

1. **Phương pháp tổng quan tài liệu:** Tiến hành rà soát, phân tích có hệ thống các nghiên cứu trong và ngoài nước liên quan đến dự báo kinh tế vĩ mô, mô hình học máy thống kê, dữ liệu lớn, và các phương pháp phân tích chuỗi thời gian.
2. **Phương pháp phân tích – tổng hợp:** Được sử dụng để đánh giá đặc điểm, cấu trúc và tính chất của dữ liệu; đồng thời phân tích logic và đối chiếu các mô hình truyền thống và hiện đại để rút ra tiêu chí lựa chọn mô hình phù hợp.

3. **Phương pháp xử lý và khai thác dữ liệu:** Bao gồm các bước thu thập, làm sạch, chuẩn hóa, phân tích đặc trưng, tạo tập dữ liệu huấn luyện – kiểm tra từ các nguồn dữ liệu truyền thống như WB, GSO, PWT,... Kỹ thuật xử lý dữ liệu thời gian, phân tích tương quan, kiểm định tính dừng, v.v. sẽ được sử dụng.
4. **Phương pháp mô hình hóa và huấn luyện mô hình học máy:** Nghiên cứu tiến hành xây dựng, huấn luyện và đánh giá các mô hình học máy như Random Forest, XGBoost, SVM, LSTM, Bi-LSTM, Transformer. Quá trình mô hình hóa bao gồm lựa chọn kiến trúc mạng và tối ưu siêu tham số.
5. **Phương pháp đánh giá thực nghiệm:** Các mô hình sẽ được kiểm định hiệu suất dự báo trên bộ dữ liệu thực tế đồng thời so sánh với các mô hình khác để kiểm tra mức độ cải thiện.
6. **Phương pháp đối sánh và kiểm chứng khả năng tổng quát:** Mô hình đề xuất sẽ được áp dụng thử nghiệm trên dữ liệu của nhiều quốc gia để kiểm chứng khả năng mở rộng, ứng dụng và tính ổn định của mô hình theo thời gian và không gian.

Ý nghĩa khoa học và thực tiễn

Về mặt khoa học, đề tài góp phần mở rộng nền tảng lý luận trong lĩnh vực dự báo kinh tế vĩ mô thông qua việc tích hợp tư duy kinh tế học với phương pháp phân tích dự báo dữ liệu hiện đại. Thay vì dựa trên các giả định tuyến tính nhất quán của các mô hình truyền thống, nghiên cứu hướng tới một cách tiếp cận linh hoạt hơn, sử dụng mô hình ML/DL để xử lý dữ liệu đa chiều, có tính chu kỳ và biến động cao. Việc xây dựng mô hình ML/DL có khả năng tích hợp thông tin chu kỳ kinh tế được kỳ vọng sẽ nâng cao hiệu quả mô hình hóa chuỗi thời gian kinh tế vĩ mô, từ đó cải thiện độ chính xác và tính ổn định trong dự báo. Đây là đóng góp khoa học có giá trị trong bối cảnh học máy đang là xu thế trong nghiên cứu kinh tế hiện đại.

Về mặt thực tiễn, trong bối cảnh nền kinh tế toàn cầu liên tục đối mặt với những biến động khó lường như dịch bệnh, căng thẳng địa chính trị và biến đổi khí hậu, nhu cầu dự báo các chỉ tiêu kinh tế vĩ mô như GDP, lạm phát hay thất nghiệp ngày càng

trở nên quan trọng. Mô hình được đề xuất trong nghiên cứu có thể đóng vai trò tham khảo đối với các cơ quan hoạch định chính sách như Bộ Tài chính, Ngân hàng Nhà nước, cục thống kê, v.v. trong việc nâng cao khả năng phân tích và hỗ trợ ra quyết định kịp thời, linh hoạt. Đồng thời, kết quả nghiên cứu cũng có thể hữu ích cho doanh nghiệp và tổ chức tài chính trong việc nhận diện xu hướng vĩ mô, từ đó phục vụ quá trình lập kế hoạch và điều chỉnh chiến lược phù hợp với bối cảnh kinh tế. Ngoài ra, việc tích hợp dữ liệu đa nguồn và đưa yếu tố chu kỳ kinh tế vào kiến trúc mô hình là một hướng tiếp cận có tiềm năng, góp phần cải thiện chất lượng dự báo và mở rộng khả năng ứng dụng trong các môi trường kinh tế có đặc điểm khác biệt.

Đóng góp của luận án

Những đóng góp cụ thể của luận án bao gồm:

Thứ nhất, xây dựng và chuẩn hóa được một cơ sở dữ liệu kinh tế vĩ mô đa quốc gia tích hợp từ nhiều nguồn đáng tin cậy như World Bank, Penn World Table và các cơ quan thống kê quốc gia. Dữ liệu được phát triển thông qua quy trình bao gồm: thu thập và hợp nhất từ nhiều định dạng và tần suất khác nhau; làm sạch và xử lý thiếu hụt dữ liệu; chuẩn hóa đơn vị, mã hóa biến và cấu trúc định danh quốc gia – thời gian; kiểm định tính dừng, đặc điểm chuỗi thời gian; và biến đổi đặc trưng phù hợp. Kết quả là một bộ dữ liệu KTVM có khả năng so sánh quốc tế, đảm bảo tính nhất quán, khả năng khái quát hóa và phù hợp cho các mô hình học.

Thứ hai, triển khai và đánh giá thực nghiệm một số mô hình dự báo GDP và một số chỉ tiêu kinh tế vĩ mô khác trên bộ dữ liệu đã được chuẩn hóa ở trên, bao gồm cả các mô hình truyền thống như hồi quy, ARIMA và hiện đại như Random Forest, XGBoost, SVM và LSTM. Trên cơ sở đó, tiến hành phân tích thống kê so sánh hiệu suất dự báo giữa các mô hình theo các chỉ số đánh giá như RMSE, MAE và R^2 nhằm đánh giá mức độ phù hợp của từng phương pháp trong các bối cảnh dữ liệu kinh tế khác nhau.

Thứ ba, đề xuất mô hình học sâu PAA-LSTM (LSTM networks with a phase adaptive attention mechanism) với nhiều lớp LSTM tích hợp attention thích ứng theo

pha chu kỳ kinh tế. Cách tiếp cận này cho phép mô hình học được các đặc trưng động học khác biệt giữa các giai đoạn của chu kỳ kinh tế (như suy thoái, tăng trưởng,...), từ đó nâng cao hiệu quả mô hình hóa chuỗi thời gian kinh tế vĩ mô. Mô hình được kiểm chứng thực nghiệm trên bộ dữ liệu của nhiều quốc gia thuộc các nhóm thể chế khác nhau cho thấy hiệu năng vượt trội so với các mô hình hiện đại như XGBoost, LSTM, Bi-LSTM và Transformer, đặc biệt trong môi trường dữ liệu biến động cao.

Cấu trúc của luận án

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

Phần Mở đầu giới thiệu bối cảnh nghiên cứu, lý do chọn đề tài, mục tiêu nghiên cứu, ý nghĩa khoa học và thực tiễn, cùng với các đóng góp chính của luận án. Phương pháp nghiên cứu được mô tả khái quát để định hướng cho các chương tiếp theo.

Chương 1 Tổng quan nghiên cứu: Trình bày tổng quan về các CTKTVM 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 phát triển bộ dữ liệu KTVM đa quốc gia từ nhiều nguồn khác nhau. Bao gồm các bước thu thập, hợp nhất, tiền xử lý, mã hóa, xử lý giá trị thiếu,... và chuẩn hóa. 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 CTKTVM khác. Các kết quả nghiên cứu được công bố trong các công trình [CT1], [CT2], [CT3] và [CT4].

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 với các mô hình truyền thống và đánh giá hiệu năng, 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 một số CTKTVM tiêu biểu và các mô hình dự báo KTVM. Nội dung tập trung vào ba trục 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 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. Cụ thể, phần đầu chương tổng quan về một số chỉ tiêu kinh tế then chốt trong hoạch định chính sách vĩ mô. Phần tiếp theo, các đặc trưng kỹ thuật của dữ liệu chuỗi thời gian và các mô hình dự báo thông dụng trong dự báo KTVM như ARIMA, VAR, v.v. được phân tích, làm cơ sở đối chiếu với các mô hình dự báo hiện đại. Cuối cùng, chương tổng hợp và phân tích các nghiên cứu đã công bố trong và ngoài nước, chỉ ra những hạn chế còn tồn tại và cơ hội để đóng góp học thuật thông qua mô hình đề xuất trong luận án.

1.1. Một số chỉ tiêu kinh tế vĩ mô tiêu biểu

Các chỉ tiêu kinh tế vĩ mô đóng vai trò trung tâm trong việc phản ánh trạng thái, xu hướng và chu kỳ vận động của nền kinh tế quốc gia [97]. Chúng là cơ sở không thể thiếu trong quá trình phân tích, hoạch định và điều hành chính sách kinh tế, tài khóa và tiền tệ [78], [129]. Trong phạm vi nghiên cứu, luận án lựa chọn và phân tích một số chỉ tiêu kinh tế vĩ mô tiêu biểu, được xác định dựa trên hệ thống khái niệm, định nghĩa và ý nghĩa được trình bày trong giáo trình *Principles of Macroeconomics* của N. Gregory Mankiw, ấn bản lần thứ tám (2020) [97].

1.1.1. Tổng sản phẩm quốc nội

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

Ý nghĩa: GDP là thước đo tổng quát nhất về quy mô và tốc độ tăng trưởng của nền kinh tế.

- **GDP tăng:** Cho thấy nền kinh tế đang tăng trưởng, sản xuất và thu nhập tăng, thường đi kèm với việc làm tốt hơn và lợi nhuận doanh nghiệp cao hơn;
- **GDP giảm:** Cho thấy nền kinh tế đang suy thoái, sản xuất và thu nhập giảm, có thể dẫn đến thất nghiệp gia tăng.

Các phương pháp tính GDP phổ biến [97]:

- **Phương pháp chi tiêu:** $GDP = \text{Chi tiêu của hộ gia đình (C)} + \text{Đầu tư (I)} + \text{Chi tiêu chính phủ (G)} + \text{Xuất khẩu ròng (NX = \text{Xuất khẩu} - \text{Nhập khẩu})}$;
- **Phương pháp thu nhập:** Tổng các khoản thu nhập được tạo ra trong nền kinh tế (tiền lương, lợi nhuận, tiền thuê, lãi suất...);
- **Phương pháp sản xuất/giá trị gia tăng:** Tổng giá trị tăng thêm của tất cả các ngành kinh tế.

1.1.2. Lạm phát

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).

Ý nghĩa:

Lạm phát phản ánh sức mua của tiền tệ. Khi mức giá chung trong nền kinh tế tăng, giá trị thực của đồng tiền suy giảm, đồng nghĩa với việc người dân có thể mua được ít hàng hóa và dịch vụ hơn với cùng một lượng tiền – đây chính là sự suy giảm sức mua của tiền tệ.

1.1.3. Chỉ số giá tiêu dùng

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.

Ý nghĩa:

Đo lường lạm phát/giảm phát: CPI là thước đo chính của lạm phát (khi giá cả tăng) hoặc giảm phát (khi giá cả giảm). Lạm phát ở mức vừa phải thường là dấu hiệu của một nền kinh tế phát triển, nhưng lạm phát cao có thể làm giảm sức mua của đồng tiền và ảnh hưởng đến đời sống người dân.

Cách tính:

$$CPI = \frac{\sum p_1 q_0}{\sum p_0 q_0} \text{ hoặc } CPI = \frac{\sum p_1 q_1}{\sum p_0 q_1} \quad (1.1)$$

Trong đó:

+ p_1 và p_0 : là giá cả hàng hóa và dịch vụ kỳ nghiên cứu và kỳ gốc trong giỏ hàng hóa

+ q_1 : là số lượng hàng hóa và dịch vụ kỳ nghiên cứu trong giỏ hàng hóa

+ q_0 : là số lượng hàng hóa và dịch vụ kỳ gốc trong giỏ hàng hóa

1.1.4. Tỷ lệ thất nghiệp

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.

Ý nghĩa:

Tỷ lệ thất nghiệp thấp thường cho thấy nền kinh tế đang tăng trưởng tốt, tạo ra nhiều việc làm. Ngược lại, tỷ lệ thất nghiệp cao cho thấy nền kinh tế đang gặp khó khăn.

Công thức:

$$\text{Tỷ lệ thất nghiệp} = (\text{Số người thất nghiệp} / \text{Lực lượng lao động}) \times 100\%$$

$$\text{Lực lượng lao động} = \text{Số người có việc làm} + \text{Số người thất nghiệp}$$

1.1.5. Vốn đầu tư trực tiếp nước ngoài

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.

Theo định nghĩa của Quỹ Tiền tệ¹ Quốc tế và Tổ chức Hợp tác và Phát triển Kinh tế², một khoản đầu tư được coi là FDI khi nhà đầu tư sở hữu ít nhất 10% quyền biểu quyết trong doanh nghiệp nhận đầu tư.

Ý nghĩa:

FDI là một trong những chỉ tiêu vĩ mô quan trọng thể hiện niềm tin của nhà đầu tư nước ngoài vào môi trường kinh doanh và triển vọng tăng trưởng kinh tế của một quốc gia.

Phân loại FDI phổ biến:

FDI dưới hình thức thành lập doanh nghiệp mới: Nhà đầu tư nước ngoài xây dựng cơ sở sản xuất mới từ đầu tại nước nhận đầu tư.

FDI dưới hình thức mua lại và sáp nhập: Nhà đầu tư nước ngoài mua lại toàn bộ hoặc một phần doanh nghiệp trong nước đang hoạt động.

FDI theo hình thức liên doanh: Nhà đầu tư nước ngoài hợp tác với nhà đầu tư trong nước để thành lập một doanh nghiệp chung.

Phương pháp đo lường FDI:

FDI cam kết: Là tổng vốn đăng ký ban đầu khi nhà đầu tư được cấp giấy chứng nhận đầu tư. Chỉ tiêu này phản ánh kỳ vọng và quy mô đầu tư tiềm năng.

FDI thực hiện: Là tổng vốn đầu tư thực tế được giải ngân trong kỳ. Chỉ tiêu này phản ánh mức độ thực hiện và hiệu quả triển khai dự án.

1.1.6. Chỉ tiêu tiêu dùng

¹ <https://www.imf.org/external/pubs/ft/bop/2007/pdf/bpm6.pdf>

² https://www.oecd.org/en/publications/oecd-benchmark-definition-of-foreign-direct-investment-2008_9789264045743-en.html

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.

Theo Hệ thống Tài khoản quốc gia năm 2008 [130], tiêu dùng được chia thành:

- *Tiêu dùng cá nhân (Household Final Consumption Expenditure – HFCE),*
- *Tiêu dùng của Chính phủ (Government Final Consumption Expenditure – GFCE)*

Ý nghĩa kinh tế: Tiêu dùng là yếu tố dẫn dắt tăng trưởng trong nhiều nền kinh tế, đặc biệt là những nền kinh tế có tỷ trọng khu vực hộ gia đình cao. Sự thay đổi trong tiêu dùng phản ánh mức độ lạc quan, thu nhập và hành vi tiết kiệm của dân cư. Tổng tiêu dùng hộ gia đình (% GDP) phản ánh tỷ trọng tiêu dùng cá nhân trong nền kinh tế

1.1.7. Tổng cầu

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.

Ý nghĩa: Tổng cầu đại diện cho sức mua tổng thể của nền kinh tế. Nếu tổng cầu tăng, doanh nghiệp sẽ mở rộng sản xuất, tuyển dụng lao động, từ đó thúc đẩy tăng trưởng. Ngược lại, tổng cầu suy yếu là tín hiệu của suy thoái.

1.1.8. Cán cân thương mại

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.

Ý nghĩa kinh tế:

Cán cân thương mại phản ánh vị thế cạnh tranh quốc tế và sức mạnh sản xuất nội địa của một quốc gia. Thặng dư (cán cân thương mại >0) cho thấy xuất khẩu vượt nhập khẩu, thường tích cực cho tăng trưởng và dự trữ ngoại tệ. Thâm hụt (cán cân

thương mại <0) biểu thị nhập khẩu cao hơn, có thể do nhu cầu tiêu dùng/đầu tư lớn nhưng tiềm ẩn rủi ro mất cân đối và áp lực tỷ giá.

1.1.9. Tỷ giá hối đoái

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.

Ý nghĩa:

Ảnh hưởng đến xuất nhập khẩu: Tỷ giá hối đoái tăng (đồng nội tệ mất giá) làm hàng hóa xuất khẩu rẻ hơn và hàng hóa nhập khẩu đắt hơn, có lợi cho xuất khẩu và hạn chế nhập khẩu. Ngược lại, tỷ giá hối đoái giảm (đồng nội tệ tăng giá) làm hàng hóa xuất khẩu đắt hơn và hàng hóa nhập khẩu rẻ hơn.

Tác động đến dòng vốn đầu tư: Tỷ giá hối đoái ảnh hưởng đến quyết định của nhà đầu tư nước ngoài khi đầu tư vào một quốc gia.

Kiểm soát lạm phát: Một đồng tiền mạnh có thể giúp kiểm soát lạm phát do hàng hóa nhập khẩu trở nên rẻ hơn.

Chính sách của ngân hàng trung ương: Ngân hàng trung ương thường can thiệp vào thị trường ngoại hối để ổn định tỷ giá, tránh những biến động lớn gây ảnh hưởng đến nền kinh tế.

1.1.10. Năng suất lao động

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.

Ý nghĩa kinh tế:

Đây là chỉ tiêu then chốt phản ánh hiệu quả sử dụng yếu tố lao động trong quá trình sản xuất, đồng thời là nền tảng quan trọng đối với tăng trưởng kinh tế dài hạn. Năng suất cao thúc đẩy tăng thu nhập, cải thiện mức sống và năng lực cạnh tranh; ngược lại, năng suất thấp hoặc tăng trưởng chậm hàm ý sự hạn chế về công nghệ, kỹ năng và phân bổ nguồn lực, có thể cản trở quá trình phát triển bền vững.

1.1.11. Chỉ số cung tiền

Định nghĩa 11: Cung tiền 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. Các chỉ số cung tiền phổ biến bao gồm:

M1: Bao gồm tiền mặt trong lưu thông và tiền gửi không kỳ hạn. Đây là chỉ số phản ánh khả năng thanh toán tức thời của nền kinh tế.

M2: Bao gồm M1 cộng với tiền gửi có kỳ hạn. Đây là chỉ số cung tiền có thể được sử dụng cho tiêu dùng và đầu tư.

Ý nghĩa kinh tế:

Cung tiền và chính sách tiền tệ ảnh hưởng trực tiếp đến tổng cầu, giá cả, đầu tư, tỷ giá và tốc độ tăng trưởng kinh tế. Khi cung tiền tăng nhanh, có thể kích thích tiêu dùng và đầu tư, từ đó thúc đẩy tăng trưởng kinh tế, nhưng đồng thời có thể làm gia tăng áp lực lạm phát nếu vượt quá năng lực sản xuất.

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

Dự báo kinh tế vĩ mô là quá trình ước lượng giá trị tương lai của các biến số kinh tế dựa trên dữ liệu quá khứ, trong đó phần lớn các chỉ tiêu như GDP, CPI, FDI, tỷ lệ thất nghiệp, v.v. đều được thu thập dưới dạng chuỗi thời gian. Do đó, việc hiểu rõ bản chất và đặc điểm kỹ thuật của dữ liệu chuỗi thời gian là điều kiện tiên quyết để xây dựng các mô hình dự báo có hiệu quả thực tiễn và độ chính xác cao.

1.2.1. Khái niệm và vai trò của chuỗi thời gian trong kinh tế

Khái niệm [65], [95]: Chuỗi thời gian (time series) là một tập hợp các điểm dữ liệu theo trình tự thời gian đều đặn, chẳng hạn theo tháng, quý hoặc năm. Trong kinh tế học, chuỗi thời gian phản ánh tiến trình vận động của các hiện tượng kinh tế vĩ mô như tổng sản phẩm quốc nội, chỉ số giá tiêu dùng, đầu tư, thương mại, tỷ lệ thất nghiệp và các chỉ tiêu liên quan đến chính sách tài khóa, tiền tệ.

Phân tích chuỗi thời gian giúp phát hiện cấu trúc nội tại của dữ liệu, từ đó hỗ trợ xây dựng mô hình định lượng và dự báo chính xác các xu hướng kinh tế trong tương lai [95].

1.2.2. Đặc điểm của chuỗi thời gian kinh tế vĩ mô

Chuỗi thời gian kinh tế vĩ mô là một loại dữ liệu đặc thù, phản ánh sự biến động của các chỉ tiêu kinh tế theo thời gian và thường được sử dụng rộng rãi trong các nghiên cứu dự báo và phân tích động học kinh tế. Khác với dữ liệu chéo hay dữ liệu phi cấu trúc, chuỗi thời gian kinh tế vĩ mô sở hữu những đặc điểm thống kê riêng biệt, đòi hỏi phải được xử lý và mô hình hóa một cách cẩn trọng. Một số đặc điểm chính bao gồm :

Thứ nhất, là tính không dừng. Nhiều chuỗi kinh tế vĩ mô thể hiện xu hướng tăng hoặc giảm theo thời gian, tức không có kỳ vọng và phương sai ổn định. Trong các mô hình truyền thống như ARIMA, VAR, v.v. việc dữ liệu không dừng có thể dẫn đến sai lệch nghiêm trọng trong ước lượng và suy luận, do đó yêu cầu kiểm định tính dừng và biến đổi (như sai phân) là bắt buộc [65][87].

Thứ hai, tính chu kỳ và mùa vụ. Các chỉ tiêu như tổng tiêu dùng, chỉ số giá, và sản lượng công nghiệp thường biến động theo chu kỳ kinh tế hoặc theo mùa (ví dụ, theo quý hoặc năm). Đây là đặc tính có thể gây sai lệch nếu không được loại bỏ hoặc đưa vào mô hình dưới dạng biến giải thích [70].

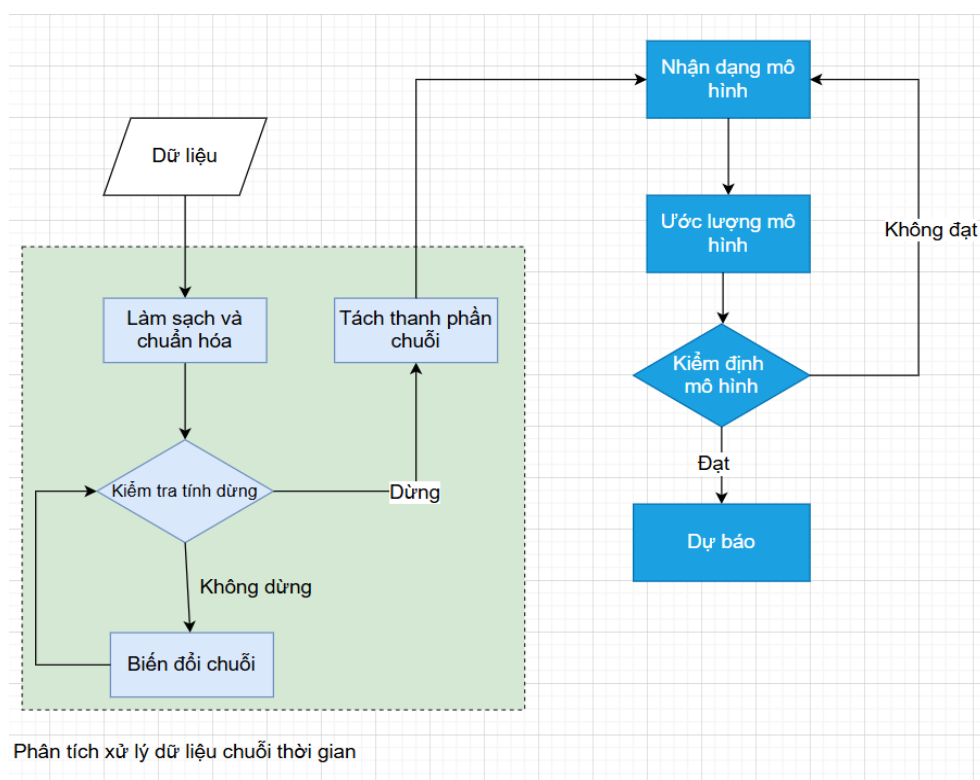
Thứ ba, tính phi tuyến và bất định. Các mối quan hệ giữa các biến kinh tế hiếm khi tuân theo cấu trúc tuyến tính đơn giản. Thay vào đó, chúng thường thể hiện tính phi tuyến, ngưỡng, hoặc hiện tượng phương sai thay đổi theo thời gian (volatility clustering) [65], [28]. Những đặc điểm này làm suy giảm hiệu quả của các mô hình tuyến tính truyền thống và là động lực cho sự phát triển của các mô hình phi tuyến hoặc học sâu [92].

Thứ tư, tính tương tác đa biến và độ trễ. Các biến kinh tế vĩ mô có mối quan hệ phụ thuộc lẫn nhau, không chỉ tức thời mà còn thông qua các độ trễ khác nhau. Chẳng hạn, tác động của chính sách tiền tệ đến chỉ số giá tiêu dùng thường chỉ thể hiện sau một số quý. Do đó, việc xác định đúng độ trễ và mô hình hóa mối quan hệ động giữa các biến là yếu tố thiết yếu trong phân tích và dự báo [65].

Thứ năm, hạn chế về độ dài và chất lượng dữ liệu. Dữ liệu chuỗi thời gian kinh tế vĩ mô thường có tần suất thấp theo quý hoặc năm, dẫn đến số quan sát hạn chế đặc biệt với các quốc gia đang phát triển. Điều này đặt ra thách thức lớn cho việc huấn luyện các mô hình học sâu, vốn yêu cầu dữ liệu lớn. Bên cạnh đó, dữ liệu có thể bị thiếu, nhiễu, hoặc sai lệch do thay đổi phương pháp thống kê, độ trễ công bố hoặc tác động từ yếu tố ngoại sinh.

1.2.3. Phân tích và xử lý chuỗi thời gian

Trước khi tiến hành mô hình hóa và dự báo, chuỗi thời gian cần được xử lý sơ bộ nhằm đảm bảo tính phù hợp với các giả định mô hình cũng như nâng cao độ chính xác dự báo [70], [87], [95]. Hình 1.1 minh họa các bước phân tích và xử lý dữ liệu chuỗi thời gian gồm các bước chính:



Hình 1. 1 Các bước phân tích và xử lý dữ liệu chuỗi thời gian

Kiểm định tính dừng: Việc xác định chuỗi có đặc tính dừng hay không là bước quan trọng đầu tiên trong xử lý chuỗi thời gian. Các kiểm định phổ biến như Augmented Dickey-Fuller (ADF), Phillips-Perron (PP) và Kwiatkowski–Phillips–

Schmidt–Shin (KPSS) được sử dụng để đánh giá liệu chuỗi có cần sai phân để đạt được tính dừng hay không [87].

Biến đổi chuỗi: Khi chuỗi không đạt tính dừng, cần thực hiện các kỹ thuật biến đổi như sai phân, logarit hóa, biến đổi Box–Cox, v.v. nhằm ổn định phương sai và trung bình chuỗi [87].

Tách thành phần xu hướng và chu kỳ: Để loại bỏ ảnh hưởng của xu hướng dài hạn hoặc biến động chu kỳ, các phương pháp như bộ lọc Hodrick–Prescott, bộ lọc Butterworth, hoặc phép sai phân bậc một có thể được áp dụng. Việc tách các thành phần này giúp làm nổi bật động lực ngắn hạn và cải thiện khả năng dự báo [70].

Với các mô hình học sâu, ngoài các bước tiền xử lý truyền thống như kiểm định tính dừng, loại xu hướng, v.v. còn có thể áp dụng thêm các kỹ thuật đặc thù như chuẩn hóa z-score, scaling Min-Max, tạo đặc trưng phi tuyến hoặc embedding chuỗi thời gian để mô hình học hiệu quả hơn trong không gian biểu diễn trừu tượng [83].

Tóm lại, việc nhận diện đúng đặc điểm chuỗi thời gian là cơ sở để lựa chọn phương pháp dự báo phù hợp. Các mô hình tuyến tính truyền thống như ARIMA, VAR, v.v. chỉ hiệu quả nếu chuỗi ổn định, tuyến tính. Với dữ liệu phi tuyến, dài hạn và phức tạp, các mô hình học sâu như LSTM, GRU, Transformer đã chứng minh hiệu quả vượt trội, đặc biệt trong việc ghi nhớ phụ thuộc dài hạn và trích xuất đặc trưng tự động [93], [112].

1.3. Một số mô hình thông dụng trong dự báo KTVM

1.3.1. Các mô hình hồi quy

Mô hình hồi quy là một trong những công cụ cơ bản và mạnh mẽ nhất trong kinh tế lượng, được sử dụng để mô tả mối quan hệ giữa một biến phụ thuộc (chỉ tiêu cần dự báo) và một hoặc nhiều biến độc lập (các yếu tố tác động) [59], [60], [121].

1.3.1.1. Mô hình hồi quy tuyến tính đơn và bội [60]

Cơ sở của phương pháp này là giả định tồn tại một mối quan hệ tuyến tính giữa các biến.

Mô hình hồi quy tuyến tính đơn có dạng:

$$Y = \beta_0 + \beta_1 X + \epsilon \quad (1.2)$$

Trong đó:

- Y là biến phụ thuộc (ví dụ: tốc độ tăng trưởng GDP);
- X là biến độc lập (ví dụ: vốn đầu tư);
- β_0 là hệ số chặn;
- β_1 là hệ số góc, đo lường mức độ thay đổi của Y khi X thay đổi một đơn vị;
- ϵ là sai số ngẫu nhiên, đại diện cho các yếu tố không được giải thích trong mô hình.

Mô hình hồi quy tuyến tính bội mở rộng từ mô hình đơn bằng cách sử dụng nhiều biến độc lập:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \epsilon \quad (1.3)$$

Trong đó $X_1, X_2, \dots, X_k$ là các biến độc lập khác nhau.

Các tham số β của mô hình thường được ước lượng bằng phương pháp Bình phương tối thiểu thông thường (Ordinary Least Squares - OLS).

1.3.1.2. Các giả định của mô hình hồi quy tuyến tính cổ điển

Để các ước lượng OLS là ước lượng BLUE, một số *giả định Gauss-Markov* cần được thỏa mãn [59]:

1. **Liên hệ tuyến tính:** Mối quan hệ giữa biến phụ thuộc và các biến độc lập là tuyến tính theo tham số.
2. **Không có đa cộng tuyến hoàn hảo:** Không có biến độc lập nào trong mô hình có thể được biểu diễn hoàn toàn bằng tổ hợp tuyến tính của các biến độc lập khác. Giả định này đảm bảo rằng các hệ số hồi quy có thể được ước lượng duy nhất và ổn định bằng phương pháp bình phương tối thiểu.
3. **Giả định về sai số:**
 - Kỳ vọng có điều kiện của sai số bằng 0: $E(\epsilon_i | X_1, \dots, X_k) = 0$;

- Phương sai của sai số không đổi: $\text{Var}(\epsilon_i) = \sigma^2$;
- Các sai số không tự tương quan: $\text{Cov}(\epsilon_i, \epsilon_j) = 0$ với $i \neq j$.

1.3.1.3. Ưu điểm và nhược điểm mô hình hồi quy tuyến tính đơn và bội

- **Ưu điểm** [60], [121]:
 - Các mô hình này đơn giản, dễ xây dựng và diễn giải. Các hệ số hồi quy cung cấp một thước đo trực tiếp về tác động của các biến giải thích;
 - Các mô hình hồi quy tuyến tính đơn và bội là nền tảng cho nhiều mô hình phức tạp hơn.
- **Nhược điểm** [60], [121]:
 - Các mô hình này yêu cầu giả định về mối quan hệ tuyến tính có thể không đúng với thực tế phức tạp của kinh tế vĩ mô;
 - Nhạy cảm với dữ liệu ngoại lai;
 - Yêu cầu các giả định nghiêm ngặt về sai số, và các giả định này thường bị vi phạm trong dữ liệu kinh tế vĩ mô (ví dụ: tự tương quan, phương sai thay đổi), đòi hỏi các kỹ thuật kiểm định và khắc phục phức tạp;
 - Khó nắm bắt được mối quan hệ động và các tương tác phức tạp giữa các biến số theo thời gian.

1.3.2. Mô hình tự hồi quy tích hợp trung bình trượt

Mô hình ARIMA (Autoregressive Integrated Moving Average) [29] là một trong những mô hình chuỗi thời gian đơn biến kinh điển, được hệ thống hóa và phổ biến rộng rãi trong công trình nền tảng của Box và Jenkins công bố năm 1970. Mô hình này giả định rằng giá trị tương lai của một biến có thể được giải thích bằng các giá trị trong quá khứ của chính nó và các sai số dự báo trong quá khứ.

Mô hình ARIMA được ký hiệu là $\text{ARIMA}(p,d,q)$, trong đó:

p (Autoregressive): Bậc của thành phần tự hồi quy.

d (Integrated): Bậc của sai phân để chuỗi thời gian trở nên dừng.

q (Moving Average): Bậc của thành phần trung bình trượt.

1.3.2.1. Các thành phần của mô hình

1. Thành phần tự hồi quy (AR(p)):

$$Y_t = c + \phi_1 Y_{t-1} + \phi_2 Y_{t-2} + \dots + \phi_p Y_{t-p} + \epsilon_t \quad (1.4)$$

Trong đó Y_t là giá trị của chuỗi tại thời điểm t , ϕ_i là các tham số của mô hình và ϵ_t là nhiễu trắng.

2. Thành phần tích hợp (I(d)): Nhiều chuỗi thời gian kinh tế vĩ mô là chuỗi không dừng. Để áp dụng mô hình, chuỗi cần được chuyển về dạng dừng thông qua phép sai phân. Phép sai phân bậc 1 được tính bằng:

$$\Delta Y_t = Y_t - Y_{t-1} \quad (1.5)$$

Bậc d chính là số lần lấy sai phân cần thiết.

3. Thành phần trung bình trượt (MA(q)):

$$Y_t = c + \epsilon_t + \theta_1 \epsilon_{t-1} + \theta_2 \epsilon_{t-2} + \dots + \theta_q \epsilon_{t-q} \quad (1.6)$$

Trong đó Y_t phụ thuộc vào các sai số dự báo trong quá khứ.

1.3.2.2. Phương pháp luận Box-Jenkins

Quy trình xây dựng mô hình ARIMA thường tuân theo phương pháp luận Box-Jenkins [29] gồm 3 bước lặp:

1. **Nhận dạng mô hình:** Sử dụng biểu đồ tự tương quan (ACF) và tự tương quan riêng phần (PACF) của chuỗi dữ liệu (sau khi đã lấy sai phân) để xác định các giá trị p và q phù hợp.
2. **Ước lượng tham số:** Ước lượng các tham số ϕ và θ của mô hình, thường bằng phương pháp hợp lý tối đa (Maximum Likelihood Estimation).
3. **Kiểm định chẩn đoán:** Kiểm tra xem phần dư của mô hình có phải là nhiễu trắng hay không. Nếu không, quay lại bước 1 để chọn mô hình khác.

1.3.2.3. Ưu điểm và Nhược điểm

Ưu điểm [73]:

- Rất linh hoạt và hiệu quả trong việc mô hình hóa các cấu trúc tương quan phức tạp trong một chuỗi thời gian đơn biến;
- Không yêu cầu các biến giải thích bên ngoài, chỉ dựa vào thông tin nội tại của chuỗi dữ liệu.

- **Nhược điểm:**

- Bản chất là mô hình đơn biến, không thể hiện được sự tác động qua lại giữa các chỉ tiêu kinh tế. Ví dụ, mô hình ARIMA dự báo lạm phát sẽ không thể hiện được tác động từ chính sách lãi suất của ngân hàng trung ương;
- Giả định rằng cấu trúc của chuỗi thời gian không thay đổi trong tương lai, do đó khó dự báo chính xác khi có các cú sốc cấu trúc (structural breaks) [111];
- Quá trình nhận dạng mô hình (chọn p, q) có thể mang tính chủ quan.

1.3.3. Mô hình tự hồi quy vector

Để khắc phục hạn chế đơn biến của mô hình ARIMA, Năm 1980 Sims đã giới thiệu mô hình tự hồi quy vector (Vector Autoregression - VAR)[118] như một phương pháp thay thế cho các mô hình kinh tế lượng cấu trúc lớn. VAR là một mô hình đa biến, trong đó mỗi biến được hồi quy theo các giá trị trễ của chính nó và của tất cả các biến khác trong hệ thống.

Cấu trúc mô hình VAR

Một mô hình VAR(p) với k biến có thể được viết dưới dạng ma trận:

$$Y_t = C + A_1 Y_{t-1} + A_2 Y_{t-2} + \dots + A_p Y_{t-p} + E_t \quad (1.7)$$

Trong đó:

- Y_t là vector cột ($k \times 1$) chứa k biến tại thời điểm t;
- C là vector ($k \times 1$) các hệ số chặn;
- A_i là các ma trận tham số ($k \times k$) tại độ trễ i;
- E_t là vector ($k \times 1$) các sai số nhiễu trắng.

Các giả định chính

1. **Tính dừng:** Tất cả các biến trong mô hình VAR phải là chuỗi dừng.
2. **Sai số:** Vector sai số E_t có giá trị kỳ vọng bằng 0, không có tự tương quan, nhưng có thể có tương quan đương thời giữa các thành phần của nó.
3. **Độ dài trễ:** Việc lựa chọn độ dài trễ p phù hợp là rất quan trọng và thường dựa trên các tiêu chí thông tin như AIC (Akaike Information Criterion) hoặc SIC (Schwarz Information Criterion) [95].

Ưu điểm và nhược điểm

- **Ưu điểm:**
 - Là mô hình đa biến, cho phép phân tích sự phụ thuộc và tác động lẫn nhau giữa các chỉ tiêu kinh tế;
 - Không yêu cầu các giả định lý thuyết cứng nhắc về mối quan hệ cấu trúc giữa các biến, giảm thiểu sai sót do đặt sai dạng mô hình [118];
 - Bên cạnh dự báo, VAR còn là công cụ mạnh để phân tích chính sách thông qua phản ứng xung lực (Impulse Response Functions - IRF) và phân rã phương sai (Variance Decomposition) (Stock & Watson, 2001)[120].
- **Nhược điểm:**
 - Việc thiếu nền tảng lý thuyết kinh tế chặt chẽ khiến các hệ số của mô hình khó được diễn giải trực tiếp;
 - Tính phức tạp theo chiều không gian (curse of dimensionality): Số lượng tham số cần ước lượng ($c.k+p \times k^2$) tăng rất nhanh khi số biến (k) hoặc độ trễ (p) tăng, làm giảm bậc tự do và hiệu quả của ước lượng [95];
 - Giống như ARIMA, VAR cũng nhạy cảm với các cú sốc cấu trúc [121].

Các mô hình hồi quy, ARIMA và VAR là những công cụ dự báo kinh tế vĩ mô truyền thống, đã được chứng minh hiệu quả trong nhiều thập kỷ. Mỗi mô hình có

những điểm mạnh và hạn chế riêng. Mô hình hồi quy giúp làm rõ mối quan hệ nhân quả nhưng bị giới hạn bởi các giả định chặt chẽ. Mô hình ARIMA mạnh mẽ trong việc nắm bắt động lực của một chuỗi thời gian duy nhất nhưng lại bỏ qua sự tương tác với các biến khác. Mô hình VAR khắc phục nhược điểm này bằng cách xem xét một hệ thống các biến, nhưng lại phải trả giá bằng sự phức tạp và khó diễn giải về mặt lý thuyết. Stock và Watson trong công bố năm 2007 [121] nhận định rằng sự tồn tại của các vấn đề như quan hệ phi tuyến, sự thay đổi cấu trúc đột ngột, và yêu cầu xử lý lượng lớn biến số đã thúc đẩy sự phát triển của các thể mô hình mới hơn.

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 được sử dụng nhằm đo lường mức độ chênh lệch giữa giá trị dự báo và giá trị thực tế.

1. MAE – Mean Absolute Error (Sai số tuyệt đối trung bình)

Công thức:

$$MAE = \frac{1}{n} \sum_{t=1}^n |y_t - \hat{y}_t| \quad (1.8)$$

Trong đó:

y_t : Giá trị thực tế tại thời điểm t

$\hat{y}_t$: Giá trị dự báo tại thời điểm t

n : Số lượng quan sát

Ý nghĩa:

MAE đo lường sai số trung bình theo đơn vị gốc của dữ liệu, phản ánh mức sai lệch trung bình giữa dự báo và thực tế. MAE dễ hiểu, ổn định và không bị ảnh hưởng mạnh bởi các giá trị ngoại lai.

2. RMSE – Root Mean Squared Error (Căn bậc hai sai số bình phương trung bình)

RMSE đo lường căn bậc hai của trung bình bình phương sai lệch giữa giá trị dự báo và giá trị thực tế, cho thấy mức độ phù hợp của mô hình với dữ liệu.

Công thức:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (1.9)$$

Trong đó y_i là giá trị thực tế, $\hat{y}_i$ là giá trị dự báo và n là số quan sát. RMSE càng nhỏ chứng tỏ mô hình dự báo càng chính xác.

Ý nghĩa:

RMSE đánh giá độ lệch trung bình bình phương giữa dự báo và thực tế, nhấn mạnh các sai số lớn hơn nhờ việc bình phương. Do đó, RMSE nhạy với các sai lệch lớn và phản ánh tốt hơn rủi ro dự báo trong các ứng dụng chính sách hoặc tài chính.

3. MAPE – Mean Absolute Percentage Error (Sai số phần trăm tuyệt đối trung bình)

Công thức:

$$MAPE = \frac{100\%}{n} \sum_{t=1}^n \left| \frac{y_t - \hat{y}_t}{y_t} \right| \quad (1.10)$$

Trong đó:

y_t : Giá trị thực tế tại thời điểm t

$\hat{y}_t$: Giá trị dự báo tại thời điểm t

n : Số lượng quan sát

Ý nghĩa:

MAPE biểu thị sai số trung bình dưới dạng phần trăm, rất hữu ích khi so sánh mô hình giữa các biến có đơn vị hoặc quy mô khác nhau. Tuy nhiên, MAPE có thể không ổn định nếu y_t gần bằng 0.

4. Hệ số xác định R^2 (Coefficient of Determination)

R^2 đo lường tỷ lệ phương sai của biến phụ thuộc y được mô hình giải thích. Nó cho biết mức độ “giải thích” hoặc “khớp” của mô hình so với đường ngang trung bình.

Công thức:

$$R^2 = \frac{\sum_{i=1}^n (\hat{y}_i - \bar{y})}{\sum_{i=1}^n (y_i - \bar{y})} \quad (1.11)$$

Trong đó:

y_i : Giá trị thực tế tại thời điểm i

$\hat{y}_i$: Giá trị dự báo tại thời điểm i

$\bar{y}$: là giá trị trung bình của y

n : Số lượng quan sát

Ý nghĩa:

R^2 cho biết tỷ lệ phương sai của biến mục tiêu được mô hình giải thích, nghĩa là bao nhiêu phần trăm biến động của dữ liệu đã được “giải thích” bởi các biến đầu vào (giá trị càng gần 1 thì mô hình càng khớp tốt).

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

1.5.1. Tình hình nghiên cứu trên thế giới

1.5.1.1. Các nghiên cứu dựa trên mô hình kinh tế lượng

Phương pháp kinh tế lượng từ lâu đã là trụ cột trong dự báo các CTKTVM, với các mô hình điển hình như ARIMA [29], VAR [118], VECM [82], GARCH [44], và DSGE [119].

Nghiên cứu của Box và Jenkins [29] năm 1970 mở đầu cho phương pháp dự báo chuỗi thời gian với mô hình ARIMA và các dạng mở rộng như SARIMA, ARIMAX, được sử dụng rộng rãi trong phân tích dữ liệu kinh tế có tính chu kỳ. Nghiên cứu này được coi là một trong những công trình kinh điển nhất trong phân tích chuỗi thời gian, và vẫn là nền tảng giảng dạy trong nhiều chương trình thống kê và kinh tế học ứng dụng.

Năm 1982, nghiên cứu của Robert F. Engle đã giới thiệu mô hình ARCH (Autoregressive Conditional Heteroscedasticity) [44] nhằm mô hình hóa phương sai thay đổi theo thời gian trong chuỗi thời gian tài chính và kinh tế. Thay vì giả định phương sai là hằng số, mô hình ARCH cho phép phương sai của sai số tại mỗi thời điểm phụ thuộc vào các giá trị sai số trong quá khứ. Điều này giúp mô hình phản ánh tốt hơn đặc điểm "cụm biến động" thường thấy trong dữ liệu kinh tế – tài chính. Năm 1986, Bollerslev [28] đã mở rộng mô hình ARCH của R. F. Engle bằng cách đề xuất phát triển mô hình GARCH (Generalized ARCH) cho phép phương sai có điều kiện tại thời điểm hiện tại phụ thuộc không chỉ vào bình phương của sai số trong quá khứ

(giống ARCH) mà còn vào chính phương sai có điều kiện trong quá khứ. Bollerslev đã trình bày chi tiết cấu trúc GARCH(p, q), đưa ra phương pháp ước lượng bằng Maximum Likelihood, và chứng minh rằng GARCH có thể được ứng dụng hiệu quả trong nhiều bối cảnh thực nghiệm như lạm phát, tỷ giá, và thị trường tài chính.

Năm 1980, trong nghiên cứu của mình Sims và cộng sự [118] khởi xướng mô hình VAR và được mở rộng bởi Stock & Watson [122], cho phép dự báo và phân tích mối quan hệ tương hỗ giữa nhiều biến kinh tế như GDP, CPI, lãi suất, v.v. Cũng trong hướng tiếp cận này, Sims và cộng sự tiếp tục phát triển mô hình SVAR nhằm áp dụng các ràng buộc kinh tế học vào hệ phương trình VAR để xác định nhân quả giữa các cú sốc kinh tế. Năm 1989, nghiên cứu của Hamilton [66] đã giới thiệu mô hình Markov-switching để mô hình hóa các giai đoạn chu kỳ kinh tế như suy thoái và tăng trưởng, rất phù hợp cho dự báo các biến động mang tính cấu trúc của GDP. Năm 1991, nghiên cứu của Johansen [82] đề xuất mô hình VECM sử dụng khi dữ liệu có mối quan hệ đồng liên kết dài hạn, chẳng hạn giữa GDP và tiêu dùng hoặc đầu tư. Trong khi đó, năm 2003 Smets và Wouters đã phát triển và ước lượng một mô hình có tên DSGE (Dynamic Stochastic General Equilibrium) [119], sau đó mô hình này được sử dụng phổ biến trong các ngân hàng trung ương cho phép mô hình hóa hành vi kinh tế dựa trên kỳ vọng hợp lý và các cú sốc ngẫu nhiên. Ngoài ra, mô hình Bayesian VAR (BVAR) do Litterman [93] phát triển năm 1986 cũng ngày càng được áp dụng trong môi trường có dữ liệu hạn chế, nhờ khả năng khắc phục tình trạng overfitting.

Một số nghiên cứu điển hình gồm: Nghiên cứu của Stock và Watson năm 2002 [122] áp dụng phương pháp VAR kết hợp phân tích thành phần chính (PCA) để rút gọn hàng trăm chỉ báo kinh tế nhằm dự báo GDP và lạm phát của Hoa Kỳ. Kết quả cho thấy mô hình VAR kết hợp PCA cải thiện đáng kể độ chính xác dự báo so với mô hình VAR đơn biến. Trong nghiên cứu của Ghazo năm 2021 [55], mô hình ARIMA được áp dụng để dự báo GDP và CPI của Jordan giai đoạn 1976–2019. Kết quả cho thấy ARIMA mang lại độ chính xác tương đối trong dự báo ngắn hạn. Adolfson và cộng sự (2007) [14] đã phát triển và ước lượng một mô hình DSGE cho nền kinh tế

mở bằng phương pháp Bayes, trong đó cho phép hiện tượng chuyển giá không hoàn toàn (incomplete exchange rate pass-through) – một đặc trưng quan trọng trong thương mại quốc tế. Mô hình được hiệu chỉnh để phù hợp với dữ liệu kinh tế vĩ mô của Thụy Điển, bao gồm GDP, tiêu dùng, lạm phát, tỷ giá và xuất nhập khẩu. Kết quả cho thấy mô hình DSGE có khả năng dự báo cạnh tranh so với VAR, đặc biệt trong các biến kinh tế có tính chu kỳ cao. Nghiên cứu là một ví dụ tiêu biểu cho việc kết hợp mô hình kinh tế học lý thuyết với kỹ thuật định lượng hiện đại, đóng góp vào xu hướng nâng cao khả năng dự báo của các mô hình cấu trúc trong điều kiện nền kinh tế mở và nhiều bất định. Tuy đạt được độ chính xác dự báo tương đối cao và thể hiện tính khả thi của việc ước lượng mô hình DSGE bằng phương pháp Bayes, nghiên cứu vẫn tồn tại một số hạn chế nhất định. Thứ nhất, việc mô hình hóa quá trình chuyển giá không hoàn toàn chỉ dừng lại ở mức tuyến tính và tĩnh, chưa phản ánh đầy đủ tác động phi tuyến và thời gian thay đổi của các cú sốc bên ngoài. Thứ hai, mô hình vẫn dựa vào giả định cấu trúc cố định và không tích hợp khả năng thích ứng động với dữ liệu mới – điều này làm giảm hiệu quả trong các tình huống có biến động lớn hoặc khủng hoảng. Ngoài ra, như nhiều mô hình DSGE khác, mô hình trong nghiên cứu vẫn giả định kỳ vọng hợp lý và tối ưu hóa đầy đủ, điều này không phản ánh hết hành vi phi lý tính hay sai lệch kỳ vọng thường gặp trong thực tiễn kinh tế vĩ mô.

Nghiên cứu của Canova và Ciccarelli [36] công bố năm 2004 đã đề xuất một mô hình Bayesian Panel VAR (BPVAR) nhằm cải thiện khả năng dự báo và phát hiện điểm đảo chiều (turning points) trong chu kỳ kinh tế. Mô hình được thiết kế để tận dụng thông tin chuỗi thời gian từ nhiều quốc gia bằng cách khai thác cấu trúc bảng, qua đó giảm thiểu độ không chắc chắn của ước lượng khi số quan sát riêng lẻ còn hạn chế. Ưu điểm chính của mô hình là khả năng kết hợp linh hoạt giữa thông tin chung toàn cầu và đặc điểm riêng của từng quốc gia, đồng thời cho phép mô hình hóa sự không đồng nhất trong phản ứng của các nền kinh tế với cùng một cú sốc. Mô hình BPVAR còn hạn chế do độ phức tạp tính toán cao, phụ thuộc vào lựa chọn prior, giả định cấu trúc ổn định và thiếu khả năng thích ứng với dữ liệu mới theo thời gian.

Nghiên cứu của M. Banerjee, M. Marcellino và I. Masten [22] công bố năm 2005 về dự báo các biến kinh tế vĩ mô như GDP, sản lượng công nghiệp và lạm phát trong điều kiện mẫu ngắn và có sự thay đổi cấu trúc. Nghiên cứu cho thấy khả năng dự báo vượt trội trong ngắn hạn, đặc biệt đối với các biến như GDP thực hoặc chỉ số sản xuất công nghiệp, nhờ tận dụng được tính đồng biến của nhiều chỉ số kinh tế vĩ mô quan sát được. Tuy nhiên, nghiên cứu cũng có hạn chế là mô hình được thiết lập trong khuôn khổ tuyến tính, nó khó có thể mô hình hóa các mối quan hệ phi tuyến hoặc các tương tác động phức tạp vốn phổ biến trong hệ thống kinh tế vĩ mô. Bên cạnh đó, mô hình chưa tích hợp khả năng thích ứng theo thời gian hoặc xử lý dữ liệu thời gian thực, dẫn đến hạn chế trong việc theo dõi các biến động nhanh và bất định của nền kinh tế hiện đại.

Nghiên cứu của Stock và Watson [123] năm 2003 là một nghiên cứu tiêu biểu trong việc sử dụng mô hình kinh tế lượng truyền thống (OLS, VAR) để dự báo các biến kinh tế vĩ mô và tài chính như GDP, lạm phát thông qua dữ liệu giá tài sản. Nghiên cứu cho thấy hiệu quả của các chỉ số tài chính như yield curve, stock prices trong việc cải thiện độ chính xác của mô hình hồi quy. Nghiên cứu có độ bao quát rộng và phân tích trên nhiều quốc gia, giúp tăng tính khái quát. Tuy nhiên, hạn chế của nghiên cứu là hiệu quả dự báo của các biến tài chính không ổn định theo thời gian và giữa các quốc gia, làm giảm tính nhất quán và độ tin cậy khi áp dụng trong thực tiễn chính sách.

Tóm lại, mặc dù các nghiên cứu dựa trên mô hình kinh tế lượng truyền thống tiếp tục đóng vai trò trọng tâm trong nghiên cứu và thực tiễn dự báo KTVM, nhưng những giới hạn về tính thích ứng, năng lực xử lý dữ liệu phức tạp và khả năng học từ dữ liệu đã thúc đẩy nhu cầu tiếp cận các phương pháp học máy và học sâu có tính linh hoạt và hiệu suất cao hơn trong bối cảnh kinh tế số hiện nay [72], [99].

1.5.1.2. Các nghiên cứu sử dụng mô hình ML/DL

Sự phát triển nhanh chóng của khoa học dữ liệu và học máy đã mở ra các phương pháp mới trong dự báo KTVM. Các mô hình ML/DL đã được sử dụng trong

hiều nghiên cứu quốc tế. Trong nghiên cứu của Masini và cộng sự (2023) [99] đã cung cấp một tổng quan sâu rộng về các phương pháp ML được ứng dụng trong dự báo chuỗi thời gian, đặc biệt là trong lĩnh vực kinh tế và tài chính. Tác giả phân tích ưu thế của các mô hình như hồi quy điều chuẩn, mô hình tập hợp (ensemble learning) và mạng nơ-ron sâu. Đánh giá cho thấy các phương pháp học máy hiện đại giúp cải thiện rõ rệt độ chính xác nhờ khả năng xử lý các mối quan hệ phi tuyến và dữ liệu có chiều cao. Tuy nhiên, nghiên cứu này mang tính tổng quan, không đi sâu vào phân tích thực nghiệm từng mô hình cụ thể.

Goulet Coulombe và cộng sự (2020) [58] thực hiện phân tích so sánh các thuật toán học máy như Kernel Ridge Regression, Random Forest, LASSO và SVR trong dự báo 5 chỉ tiêu kinh tế vĩ mô trong 5 khoảng thời gian khác nhau với bộ dữ liệu 38 năm. Các mô hình được phân loại theo đặc trưng học máy như tính phi tuyến, kiểm định chéo hàm mất mát,... áp dụng trong cả môi trường dữ liệu phong phú và dữ liệu hạn chế. Kết quả cho thấy tính phi tuyến đóng vai trò then chốt trong việc cải thiện độ chính xác dự báo, đặc biệt với dữ liệu phong phú, điều đó gợi ý rằng mô hình nhân tố phi tuyến một phần là lựa chọn phù hợp cho hầu hết biến số và giai đoạn dự báo.

Ibrahim và cộng sự (2024) [74] đề xuất một mô hình lai kết hợp giữa Random Forest, Vector Error Correction Model (VECM) và hồi quy tuyến tính nhằm tăng cường khả năng dự báo GDP và tỷ giá. Bài báo cho thấy việc tích hợp kỹ thuật ML với mô hình kinh tế lượng truyền thống có thể tận dụng được tính chính xác và khả năng diễn giải. Tuy nhiên, bài viết chủ yếu tập trung vào Nigeria, do đó cần kiểm định thêm với dữ liệu các quốc gia khác.

Kashif Beg và cộng sự (2025) [25] tiến hành so sánh nhiều mô hình học máy trong dự báo biến động ngành tài chính ở Ấn Độ. Random Forest cho kết quả vượt trội với R^2 đạt 0.998, cao hơn đáng kể so với các mô hình như Gradient Boosting, SVR hay MLP. Nghiên cứu này góp phần chứng minh tính hiệu quả của học máy phi tuyến trong môi trường biến động cao và dữ liệu tài chính ngắn hạn.

Jallow và cộng sự (2025) [81] đề xuất mô hình RNN-LSTM kết hợp học chuyển giao để dự báo GDP dựa trên dữ liệu kiểu hồi. Mô hình đạt R^2 tới 91.28%,

chứng minh hiệu quả của việc kết hợp kiến trúc học sâu và chiến lược huấn luyện cải tiến. Tuy nhiên, hiệu quả mô hình phụ thuộc lớn vào độ phù hợp của dữ liệu và việc xử lý chuỗi đầu vào.

Nghiên cứu của Mohapatra và cộng sự (2024) [53] tập trung đánh giá thực nghiệm khả năng dự báo của các mô hình học máy trong kinh tế vĩ mô, thông qua việc áp dụng một loạt thuật toán như Linear Regression, Decision Tree, RBF Neural Networks, RNN và GANs để dự báo các chỉ tiêu như GDP, CPI và chỉ số S&P 500. Bài báo so sánh hiệu quả của các mô hình ML với các phương pháp truyền thống như hồi quy đa biến, đồng thời trình bày quy trình xử lý dữ liệu, huấn luyện và đánh giá mô hình trên bộ dữ liệu thực tế. Kết quả thực nghiệm cho thấy các mô hình học máy, đặc biệt là mạng nơ-ron và cây quyết định, vượt trội hơn về độ chính xác trong dự báo ngắn hạn so với các phương pháp kinh tế lượng cổ điển. Mô hình RBF Neural Network cho sai số thấp nhất nhờ khả năng xử lý phi tuyến mạnh mẽ. Nghiên cứu cũng đề xuất kết hợp các mô hình học sâu (như GANs) với các mô hình thống kê truyền thống để cải thiện độ chính xác và khả năng khái quát hóa trong điều kiện dữ liệu hạn chế.

Oancea (2025)[104] áp dụng LSTM để dự báo GDP của Romania trong giai đoạn 1995–2023 và cho thấy mô hình này vượt trội về độ chính xác và tính linh hoạt so với các mô hình truyền thống. Ngoài LSTM, các kiến trúc lai như CNN-LSTM và Bi-LSTM đã được nghiên cứu nhằm tận dụng lợi thế của cả mạng tích chập (khai thác đặc trưng cục bộ) và mạng hồi tiếp (khai thác động lực học thời gian) [51], [84], [117].

Cơ chế attention, lần đầu tiên được phát triển bởi Bahdanau và cộng sự (2014) [20] trong lĩnh vực dịch máy, đã được điều chỉnh để tăng cường hiệu năng trong dự báo chuỗi thời gian. Qin và cộng sự (2017) [112] đề xuất cơ chế attention hai tầng (dual-stage attention) trong mạng LSTM, và chứng minh rằng nó giúp mô hình tập trung vào các thời điểm thông tin quan trọng trong lịch sử, từ đó cải thiện kết quả dự báo tài chính.

Một bước tiến nữa trong lĩnh vực mô hình hóa chuỗi thời gian là sự ra đời của kiến trúc Transformer, được đề xuất bởi Vaswani và cộng sự năm 2017 [133], Transformer sử dụng cơ chế attention để học biểu diễn thông tin toàn cục mà không phụ thuộc vào vị trí thời gian, qua đó nâng cao khả năng mô hình hóa các quan hệ dài hạn trong dữ liệu chuỗi. Gần đây, một số nghiên cứu đã mở rộng ứng dụng Transformer sang lĩnh vực dự báo kinh tế. Nghiên cứu của Li và cộng sự năm 2025 [90] cho thấy Transformer có thể được áp dụng thành công trong dự báo giá cổ phiếu. Tuy nhiên, Transformer đòi hỏi dữ liệu lớn và tài nguyên tính toán cao [90], điều này làm hạn chế khả năng ứng dụng trong các bài toán kinh tế vĩ mô - nơi dữ liệu thường có tần suất thấp, thiếu liên tục và kích thước mẫu nhỏ.

Để giảm thiểu hạn chế này, các mô hình attention thích ứng đã được nghiên cứu nhằm cho phép điều chỉnh trọng số chú ý theo ngữ cảnh dữ liệu, thay vì sử dụng cố định. Qiu, Wang và Zhou (2020) [113] đề xuất mô hình sử dụng cơ chế attention động để nâng cao chất lượng dự báo chuỗi thời gian, nhưng chưa xét đến yếu tố pha chu kỳ kinh tế trong việc xác định trọng số attention .

Một số hướng tiếp cận kết hợp mô hình LSTM với Hidden Markov Model đã được triển khai nhằm mô hình hóa hành vi chuyển trạng thái (regime switching) trong dữ liệu kinh tế [39].

Bên cạnh các công trình đã công bố trong các tạp chí khoa học, một số luận án tiến sĩ gần đây đã tiếp cận bài toán dự báo kinh tế vĩ mô từ nhiều góc độ khác nhau. Luận án của Evripidis Bantis (2024) [23] khai thác vai trò của dữ liệu bổ sung từ GoogleTrend trong cải thiện hiệu năng mô hình dự báo. Luận án của Federico Baldo (2024) [21] tập trung vào việc dự báo chuỗi thời gian với phương pháp tích hợp kiến thức chuyên môn (domain knowledge) vào ML để cải thiện độ chính xác khi dữ liệu hạn chế hoặc nhiễu. Đặc biệt, luận án của Cosimo Izzo (2022) [79] nghiên cứu ứng dụng học sâu có thể giải thích (explainable deep learning) trong kinh tế học và tài chính, mở ra hướng tiếp cận quan trọng cho tính minh bạch mô hình. Các luận án của Georgios Kontogeorgos (2023) [86] và Yu Guo (2023) [61] lần lượt sử dụng kỹ thuật Bayes và giảm chiều dữ liệu để xây dựng mô hình dự báo kinh tế với độ chính xác và

hiệu quả cao hơn. Những công trình này thể hiện rõ xu thế hiện đại hóa phương pháp dự báo trong bối cảnh dữ liệu lớn và yêu cầu cao về tính khả diễn giải.

Tóm lại, 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. Tình hình 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 công trình tiêu biểu có thể kể đến các nghiên cứu sau.

Nghiên cứu của Phạm Thành Công, Viện kinh tế Việt Nam năm 2024¹ đã so sánh mô hình ARIMA và mô hình VECM để dự báo tăng trưởng GDP quý của Việt Nam trong các quý cuối năm 2024. Nghiên cứu này chỉ ra rằng ARIMA phù hợp hơn cho dự báo tăng trưởng ngắn hạn với độ biến động thấp, trong khi VECM tốt hơn để nắm bắt xu hướng dài hạn và các biến động lớn, nhờ khả năng nắm bắt mối quan hệ dài hạn giữa các biến số kinh tế. Nghiên cứu sử dụng dữ liệu kinh tế vĩ mô hàng quý

¹ <https://kinhtevadubao.vn/du-bao-tang-truong-kinh-te-viet-nam-3-quy-cuoi-nam-2024-trong-so-sanh-giua-mo-hinh-arima-va-vecm-29070.html>

từ quý 1/2010 đến quý 1/2024 từ cục thống kê, IMF, ADB, Bloomberg và cơ sở dữ liệu CEIC, thực hiện phân tích bằng phần mềm Eview 11.

Nghiên cứu của Phạm Thành Công và Lý Đại Hùng, viện kinh tế Việt Nam năm 2025 tập trung vào dự báo tăng trưởng kinh tế Việt Nam năm 2025 thông qua việc so sánh mô hình tự hồi quy có điều kiện phương sai thay đổi tổng quát (GARCH) và lấy mẫu dữ liệu hỗn hợp (MIDAS) [1]. Kết quả chỉ ra rằng mô hình GARCH (1,0) vượt trội hơn MIDAS, đặc biệt trong khả năng xử lý các sai số lớn và dự báo tăng trưởng trung bình là 6.77%. GARCH được đánh giá cao hơn do khả năng nắm bắt phương sai thay đổi của sai số tăng trưởng kinh tế, phản ánh sự biến động của tăng trưởng theo thời gian. Liên quan đến GARCH và MIDAS cũng một số các nghiên cứu nổi bật. Đối với mô hình GARCH, một số nghiên cứu điển hình, như: Nguyễn Thanh Hương và Bùi Quang Trung (2015) [4] dự báo chỉ số VN-Index trong giai đoạn 04/08/2014 đến ngày 21/11/2014; Lê Văn Tuấn và Phùng Duy Quang (2020) [12] dự báo ảnh hưởng của dịch Covid-19 đến thị trường chứng khoán Việt Nam; Nguyễn Thị Hiền và cộng sự (2022) [3] phân tích độ biến động của hợp đồng tương lai VN30F1M trên thị trường chứng khoán phái sinh Việt Nam; Phùng Duy Quang (2024) [6] ứng dụng mô hình GARCH dự báo giá gạo Thái Lan. Đối với mô hình MIDAS, một số nghiên cứu điển hình như Trần Minh Quân (2021) [5] về dự báo lạm phát ở Việt Nam, cho thấy rằng mô hình MIDAS cung cấp dự báo chính xác hơn so với các mô hình truyền thống, với hai chỉ số về MAE và RMSE thấp hơn đáng kể. Nguyễn Thanh Tùng và Phạm Thị Lan Anh (2022) [13] sử dụng mô hình MIDAS để dự báo tăng trưởng của Việt Nam, kết hợp dữ liệu hàng ngày về tỷ giá hối đoái và lãi suất với dữ liệu GDP hàng quý từ năm 2010 đến năm 2021. Kết quả ghi nhận mô hình MIDAS không chỉ cải thiện độ chính xác của dự báo mà còn cung cấp các thông tin hữu ích cho việc ra quyết định chính sách; Lê Mai Trang và cộng sự (2022) [11] dự báo tăng trưởng GDP của Việt Nam dựa trên bộ số liệu thu thập từ năm 2006-2020 thông qua 2 mô hình là MIDAS và MF-VAR. Kết quả nghiên cứu thực nghiệm cho thấy, mô hình MIDAS cho kết quả dự báo tốt hơn so với mô hình MF-VAR.

Một nghiên cứu tiên phong của Lê Hà Thu và Roberto Leon-Gonzalez năm 2021 đã áp dụng Dynamic Model Averaging (DMA) để dự báo lạm phát và tăng trưởng ở Việt Nam, so sánh với nhiều phương pháp chuỗi thời gian khác [127]. DMA cho thấy khả năng dự báo chính xác hơn một cách nhất quán cho cả lạm phát và tăng trưởng kinh tế, đặc biệt là do khả năng thích ứng với sự thay đổi của các mối quan hệ kinh tế theo thời gian. Phương pháp này cho phép mô hình dự báo thay đổi linh hoạt theo thời gian, sử dụng các biến dự báo có độ dài và tần suất khác nhau, điều này đặc biệt hữu ích trong một nền kinh tế đang phát triển nhanh chóng như Việt Nam, nơi các cú sốc bên trong và bên ngoài có thể làm thay đổi mối quan hệ giữa các biến số vĩ mô.

Nghiên cứu của Nguyễn Thị Hiền và cộng sự, đã ứng dụng mô hình ARIMA để dự báo lạm phát ở Việt Nam trong nửa đầu năm 2023, sử dụng dữ liệu từ tháng 7/2009 đến tháng 1/2023. Mô hình ARIMA(1,1,12) được xác định là phù hợp nhất dựa trên tiêu chí AIC và SIC, với sai số dự báo trong mẫu khá nhỏ, khoảng 1%. Nghiên cứu này đã dự báo tỷ lệ lạm phát của Việt Nam trong giai đoạn từ tháng 2/2023 đến tháng 5/2023 xoay quanh mức 5% [2].

Nghiên cứu của Trần Thanh Hoa¹, năm 2017 ứng dụng VAR và ARMA để dự báo lạm phát hàng tháng và GDP theo quý, nhấn mạnh đến vai trò của cung tiền, tỷ giá và giá dầu. Mô hình VAR được đánh giá là tốt nhất cho dự báo lạm phát hàng tháng, trong khi mô hình tự hồi quy (AR) riêng lẻ có thể phù hợp cho dữ liệu hàng quý.

Ngoài ra, giai đoạn 2010–2020 tại các đơn vị như Trung tâm Thông tin và Dự báo Kinh tế – Xã hội Quốc gia (NCIF) dưới sự dẫn dắt của PGS.TS. Đỗ Văn Thành. Nhiều công trình quan trọng được thực hiện: Dự báo GDP tiềm năng và cân đối ngân sách cấu trúc (2013–2020) bằng hệ mô hình nhiều phương trình và phân tích định

¹ https://bccprogramme.org/wp-content/uploads/2021/04/tran_hoa_heidwp05-2017_2.pdf

tính – định lượng kết hợp. Sử dụng mô hình đa phương trình có cấu trúc và mô phỏng chính sách để xây dựng kịch bản tăng trưởng, lạm phát, thu chi ngân sách trong trung – dài hạn. Tuy nhiên, cách tiếp cận vẫn chủ yếu dựa trên giả định tuyến tính, tính ổn định cấu trúc, và thiếu khả năng học tự động từ dữ liệu.

Nhóm nghiên cứu của Đỗ Văn Thành và Nguyễn Minh Hải đã công bố một loạt công trình có giá trị nhằm xây dựng mô hình dự báo các chỉ số kinh tế – tài chính, điển hình là VN-Index và CPI, dựa trên sự kết hợp giữa các phương pháp thống kê truyền thống và kỹ thuật giảm chiều dữ liệu. Trong nghiên cứu năm 2016A [8] các tác giả sử dụng kiểm định nhân quả Granger để lựa chọn các chỉ số kinh tế – tài chính có ảnh hưởng đáng kể đến VN-Index từ tập dữ liệu gồm 277 biến theo ngày. Sau đó, mô hình ARDL được ước lượng bằng phương pháp OLS được triển khai để dự báo chỉ số VN-Index theo ngày. Kết quả cho thấy mô hình có độ chính xác ngoài mẫu tương đối cao, phản ánh khả năng khai thác cả quan hệ ngắn hạn và dài hạn giữa các biến đầu vào và biến mục tiêu.

Tiếp theo, trong nghiên cứu năm 2016B [9], nhóm tác giả tiếp tục phát triển hướng tiếp cận bằng cách kết hợp phương pháp Granger với phân tích thành phần chính (PCA). Việc này nhằm giảm thiểu hiện tượng đa cộng tuyến và loại bỏ các biến dư thừa trong tập dữ liệu đầu vào. Sau hai bước giảm chiều, mô hình ARDL-OLS tiếp tục được sử dụng để dự báo VN-Index. Kết quả nghiên cứu cho thấy phương pháp này có độ chính xác vượt trội so với phương pháp ở nghiên cứu trước.

Trong một hướng ứng dụng khác, nghiên cứu của Cù Thu Thủy và cộng sự (2017) [10] đã xây dựng mô hình dự báo chỉ số giá tiêu dùng (CPI) theo tháng. Phương pháp của nhóm tác giả gồm ba bước: (i) sử dụng hệ số tương quan để lựa chọn ban đầu 52 biến kinh tế có liên quan đến CPI; (ii) áp dụng PCA để giảm chiều dữ liệu; và (iii) ước lượng mô hình hồi quy tuyến tính đa biến bằng phương pháp OLS. Mô hình này cho kết quả khá khả quan, phù hợp với dữ liệu chuỗi thời gian kinh tế vĩ mô ở Việt Nam.

Bên cạnh các phương pháp hồi quy tuyến tính, Đỗ Văn Thành (2017) [7] đã đề xuất áp dụng mô hình GARCH để kiểm định hiệu quả của thị trường chứng khoán,

cũng như xử lý phần dư trong mô hình dự báo giá cổ phiếu. Việc ứng dụng GARCH giúp đảm bảo tính ổn định và không thiên lệch trong quá trình dự báo, đặc biệt là khi xảy ra các cú sốc kinh tế.

Tổng kết lại, các nghiên cứu ở Việt Nam gần đây – dù chủ yếu dựa trên nền tảng mô hình truyền thống – đã thể hiện sự tiến bộ đáng kể trong ứng dụng các kỹ thuật tiên xử lý và giảm chiều dữ liệu vào mô hình dự báo chỉ tiêu kinh tế vĩ mô. Đồng thời, các công trình này đặt nền móng cho sự chuyển dịch từ dự báo truyền thống sang hướng tiếp cận hiện đại hơn như học máy, đặc biệt là khi nhu cầu dự báo chính sách kinh tế ngày càng đòi hỏi độ chính xác, khả năng thích ứng và mức độ giải thích cao hơn.

Trong những năm gần đây, đặc biệt sau đại dịch COVID-19, khi nền kinh tế biến động mạnh và dữ liệu trở nên phong phú hơn, một số nghiên cứu bắt đầu thử nghiệm các phương pháp học máy (ML) và học sâu (DL) cho dự báo kinh tế vĩ mô:

Nghiên cứu của Trần Phước & Hoàng Anh Lê (2024) [128] đã áp dụng MLP và KNN để dự báo GDP và CPI Việt Nam trong giai đoạn 1996–2021, so sánh với mô hình VAR và LASSO, cho thấy hiệu quả dự báo của MLP vượt trội.

Nghiên cứu của Le và cộng sự tháng 2/2024 đã đánh giá khả năng dự báo của mười thuật toán học máy được chọn cho tăng trưởng kinh tế ở Việt Nam [88]. Kết quả cho thấy mô hình Random Forest mang lại dự báo tốt hơn các mô hình khác. Đối với năm 2019, dự báo của mô hình RF tương đương với các dự báo của các tổ chức quốc tế. Tuy nhiên, dự báo của RF có sự khác biệt nhỏ so với các tổ chức này trong những năm sau đó, đặc biệt là vào năm 2023, khi kết quả của RF thấp hơn các dự báo của họ, điều này phản ánh điều kiện thực tế của nền kinh tế Việt Nam.

Nghiên cứu “Predicting subnational GDP in Vietnam with remote sensing data: A machine learning approach” của Suleiman, Nguyen và Mendez (2025) [125] trình bày một phương pháp tiên tiến để ước tính GDP cấp tỉnh tại Việt Nam bằng dữ liệu viễn thám kết hợp với học máy. Bằng cách sử dụng ảnh đêm vệ tinh và các đặc trưng không gian khác, nhóm tác giả xây dựng mô hình dự báo chính xác GDP ở cấp

địa phương – nơi dữ liệu thống kê còn hạn chế. Nghiên cứu áp dụng các mô hình như Ridge Regression, Random Forest và ANN. Kết quả cho thấy mô hình học máy với dữ liệu viễn thám có thể bổ sung hiệu quả cho các nguồn dữ liệu truyền thống trong giám sát kinh tế vùng

Tuy nhiên, 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.

Để hệ thống hóa các công trình trong nước, các nghiên cứu được phân loại theo hướng tiếp cận và phương pháp dự báo. Bảng dưới đây tổng hợp các nghiên cứu tiêu biểu, phản ánh sự chuyển dịch từ mô hình kinh tế lượng truyền thống sang học máy và học sâu, phù hợp với xu hướng hiện đại hóa dự báo kinh tế vĩ mô trong bối cảnh dữ liệu lớn và trí tuệ nhân tạo.

Bảng 1. 1 Bảng phân loại nghiên cứu trong nước theo hướng tiếp cận và phương pháp

Nhóm tiếp cận	Phương pháp / Mô hình tiêu biểu	Tác giả & Năm	Đối tượng / Mục tiêu dự báo	Đặc điểm nổi bật
1. Mô hình kinh tế lượng	ARIMA, VAR, VECM, ARDL, OLS	Phạm Thành Công (2024); Trần Thanh Hoa (2017); Nguyễn Thị Hiên (2023)	GDP, lạm phát	Tiếp cận tuyến tính; phù hợp dữ liệu chuỗi thời gian truyền thống; dễ diễn giải và áp dụng
	GARCH, EGARCH	Phạm Thành Công & Lý Đại Hùng (2025); Lê Văn Tuấn & Phùng Duy Quang (2020); Nguyễn Thị Hiên (2022)	Biến động tăng trưởng, VN-Index, giá hàng hóa	Mô hình hóa phương sai thay đổi; phản ánh biến động ngắn hạn
	MIDAS, MF-VAR	Trần Minh Quân (2021); Nguyễn Thanh Tùng (2022); Lê Mai Trang (2022)	Lạm phát, tăng trưởng GDP	Kết hợp dữ liệu tần suất cao và thấp; cải thiện độ chính xác ngắn hạn

Nhóm tiếp cận	Phương pháp / Mô hình tiêu biểu	Tác giả & Năm	Đối tượng / Mục tiêu dự báo	Đặc điểm nổi bật
	DMA (Dynamic Model Averaging)	Lê Hà Thu & Roberto Leon-Gonzalez (2021)	Lạm phát, tăng trưởng GDP	Mô hình Bayes động; thích ứng theo thời gian và thay đổi cấu trúc dữ liệu
2. Mô hình kết hợp thống kê & giảm chiều dữ liệu	PCA, Granger Causality + ARDL, Multi-equation models	Đỗ Văn Thành & Nguyễn Minh Hải (2016A, 2016B); Cù Thu Thủy (2017)	VN-Index, CPI, GDP tiềm năng	Giảm đa cộng tuyến, chọn lọc biến đầu vào hợp lý, tăng độ ổn định mô hình
3. Học máy	Random Forest, LASSO, KNN, Ridge Regression	Le và cộng sự (2024); Trần Phước & Hoàng Anh Lê (2024); Suleiman, Nguyễn và Mendez (2025)	GDP, CPI, GDP cấp tỉnh	Học tự động từ dữ liệu; mô tả quan hệ phi tuyến; có khả năng mở rộng dữ liệu phi truyền thống (vệ tinh, ảnh đêm)
4. Mô hình học sâu	MLP, ANN	Trần Phước & Hoàng Anh Lê (2024)	GDP, CPI	Mô hình phi tuyến nhiều tầng; khả năng học đặc trưng ẩn; độ chính xác cao hơn các mô hình tuyến tính.

Nhìn chung, trong nghiên cứu dự báo kinh tế vĩ mô tại Việt Nam, mô hình kinh tế lượng vẫn chiếm ưu thế trong giai đoạn 2010–2020. Các mô hình lai đánh dấu giai đoạn chuyển tiếp, khi các kỹ thuật thống kê được kết hợp với phương pháp chọn biến và giảm chiều dữ liệu. Từ năm 2023, học máy và học sâu bắt đầu được ứng dụng, phản ánh xu hướng thích ứng với dữ liệu biến động và mở ra hướng tiếp cận mới dựa trên dữ liệu lớn và phi truyền thống.

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 hướng đến giải quyết:

Thứ nhất, mặc dù đã có nhiều nghiên cứu ứng dụng các mô hình kinh tế lượng, học máy và học sâu vào dự báo các chỉ tiêu kinh tế vĩ mô, phần lớn các công trình

hiện nay chỉ tập trung vào một quốc gia hoặc một khu vực cụ thể. Do đó, dữ liệu sử dụng trong các nghiên cứu này thường mang tính đơn lẻ, phản ánh đặc điểm và bối cảnh riêng của từng nền kinh tế. Tuy nhiên, mỗi quốc gia lại có cách tiếp cận và công bố số liệu thống kê kinh tế khác nhau, tùy thuộc vào mục tiêu quản lý, trình độ phát triển, và hệ thống thể chế. Sự khác biệt này không chỉ thể hiện ở hệ thống chỉ tiêu công bố, đơn vị đo lường, tần suất cập nhật, mà còn ở phương pháp thu thập và tổng hợp số liệu. Điều này đặt ra thách thức lớn trong việc mở rộng mô hình dự báo sang môi trường đa quốc gia, đòi hỏi cần có một hướng tiếp cận trong xây dựng bộ dữ liệu đa nguồn và thiết kế mô hình có khả năng khái quát hóa tốt hơn trong điều kiện dữ liệu không đồng nhất giữa các quốc gia.

Thứ hai, thiếu các mô hình dự báo có khả năng học linh hoạt và thích ứng theo chu kỳ kinh tế. Các mô hình tuyến tính như ARIMA, VAR hoặc các mô hình học máy truyền thống thường giả định cấu trúc quan hệ giữa các biến là bất biến theo thời gian. Điều này không phù hợp trong bối cảnh dữ liệu kinh tế có chu kỳ, chịu ảnh hưởng bởi các yếu tố ngoại sinh như khủng hoảng tài chính, đại dịch, hoặc thay đổi chính sách. Mặc dù các mô hình học sâu như LSTM, GRU, Transformer có tiềm năng thích ứng cao, nhưng vẫn chưa được nghiên cứu một cách hệ thống và chuyên sâu trong dự báo kinh tế vĩ mô đặc biệt theo khảo sát tài liệu chưa có mô hình học có khả năng học linh hoạt và thích ứng theo chu kỳ kinh tế.

Thứ ba, còn ít đánh giá khả năng tổng quát hóa của mô hình trong môi trường liên quốc gia. Phần lớn các nghiên cứu chỉ dừng lại ở việc thử nghiệm mô hình trên một hoặc vài nền kinh tế riêng lẻ, chưa đánh giá tính ổn định và khả năng khái quát 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, xây dựng và chuẩn hóa một cơ sở dữ liệu kinh tế vĩ mô đa quốc gia tích hợp từ nhiều nguồn đáng tin cậy, với quy trình xử lý dữ liệu bao gồm thu thập,

hợp nhất, làm sạch, chuẩn hóa, kiểm định tính dừng, phân tích chuỗi thời gian và biến đổi đặc trưng có ý nghĩa kinh tế học.

Thứ hai, đánh giá hiệu quả của các mô hình dự báo đối với GDP và một số CTKTVM khác từ đó định hướng đề xuất xây dựng mô hình dự báo hiệu quả.

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. Kiểm chứng tính khái quát và độ tin cậy của mô hình trên tập dữ liệu kinh tế vĩ mô của các quốc gia đại diện cho các nền kinh tế có đặc điểm phát triển khác biệt nhau về thể chế, quy mô và chu kỳ.

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 GDP và một số CTKTVM khác. 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

Dữ liệu kinh tế vĩ mô đóng vai trò nền tảng trong các nghiên cứu thực nghiệm và mô hình dự báo hiện đại. Tuy nhiên, sự không đồng nhất giữa các nguồn về định nghĩa, đơn vị, và tần suất thu thập đã đặt ra những thách thức đáng kể trong việc xây dựng dữ liệu đầu vào chuẩn hóa, đặc biệt với các mô hình học máy đòi hỏi tính nhất quán cao [49], [131]. Những khác biệt này không chỉ gây khó khăn trong quá trình tích hợp, mà còn ảnh hưởng trực tiếp đến độ tin cậy và khả năng học của các mô hình dự báo, đặc biệt là trong các phương pháp dựa trên học máy và phân tích chuỗi thời gian, vốn đòi hỏi dữ liệu đầu vào phải ổn định, đầy đủ và nhất quán để đảm bảo hiệu suất học và dự báo chính xác [27], [100].

Trong bối cảnh đó, chương này nhằm xây dựng bộ dữ liệu kinh tế vĩ mô đa nguồn, phục vụ cho mục tiêu mô hình hóa và dự báo liên quốc gia. Bộ dữ liệu được thiết kế để đáp ứng đồng thời yêu cầu của các mô hình thống kê truyền thống và các kiến trúc học sâu hiện đại. Nội dung chương bao gồm: (i) Quy trình xây dựng bộ dữ liệu KTVM đa nguồn, (ii) Xác định các CTKTVM cần thu thập, (iii) Xác định các nguồn dữ liệu KTVM và kỹ thuật thu thập, (iv) Phân tích đặc điểm, thống kê mô tả, (v) phát hiện xử lý ngoại lệ và gán giá trị bị thiếu, (vi) kiểm định chuỗi thời gian, và (vii) chuẩn hóa và chuyển đổi dữ liệu theo hướng tốt nhất cho việc xây dựng bộ dữ liệu đầu vào chất lượng cao cho mô hình học máy.

2.1. Giới thiệu

Trong nghiên cứu dự báo kinh tế vĩ mô, chất lượng và tính nhất quán của dữ liệu đầu vào là yếu tố then chốt quyết định đến hiệu quả và độ tin cậy của kết quả phân tích. Các nguồn dữ liệu kinh tế vĩ mô hiện nay, dù phong phú và có uy tín như WB, IMF, PWT, OECD, FRED, GSO hay FiinGroup, vẫn tồn tại những hạn chế đáng kể khi sử dụng trực tiếp trong các nghiên cứu dự báo hoặc phân tích liên quốc gia:

Sự không đồng bộ giữa các nguồn dữ liệu:

- Sự khác biệt về khái niệm chỉ tiêu, phương pháp tính toán, đơn vị đo lường và cách mã hóa dữ liệu gây khó khăn trong việc tích hợp;
- Chu kỳ công bố và tần suất dữ liệu không thống nhất (năm, quý hoặc bất quy tắc), làm giảm tính liên tục và so sánh theo thời gian.

Phân mảnh và thiếu bao phủ

- 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]. Ví dụ, World Bank¹ chủ yếu cung cấp các chỉ tiêu liên quan đến tăng trưởng, giáo dục, y tế và phát triển bền vững với dữ liệu có độ phủ rộng nhưng tần suất chủ yếu theo năm; trong khi đó, Quỹ Tiền tệ Quốc tế² lại tập trung vào các chỉ tiêu tài khóa, tiền tệ và cán cân vãng lai, thường được cập nhật theo quý hoặc tháng. Penn World Table³ cung cấp các chỉ số hiệu suất kinh tế được điều chỉnh sức mua; OECD⁴ và FRED⁵ lại có hệ thống thống kê tài chính vĩ mô chi tiết riêng biệt. Ngoài ra, các cơ quan thống kê quốc gia như GSO⁶ của Việt Nam hay BEA⁷ của Hoa Kỳ sử dụng hệ thống mã chỉ tiêu nội bộ và lịch công bố khác nhau.
- Dữ liệu bị thiếu hoặc không đồng đều giữa các giai đoạn, ảnh hưởng tới khả năng phân tích xu hướng dài hạn.

Khó khăn trong tích hợp và chuẩn hóa

¹ <https://data.worldbank.org/>

² <https://data.imf.org>

³ <https://www.rug.nl/ggdc/productivity/pwt/>

⁴ <https://www.oecd.org/en/data.html>

⁵ <https://fred.stlouisfed.org/>

⁶ <https://www.gso.gov.vn>

⁷ <https://www.bea.gov/data>

- Quá trình đồng bộ dữ liệu từ nhiều nguồn đòi hỏi các bước tiền xử lý phức tạp để loại bỏ sai lệch và đảm bảo tính so sánh;
- Các dữ liệu gốc thường chưa ở định dạng tối ưu cho phân tích định lượng, yêu cầu bổ sung bước chuyển đổi và cấu trúc hóa.

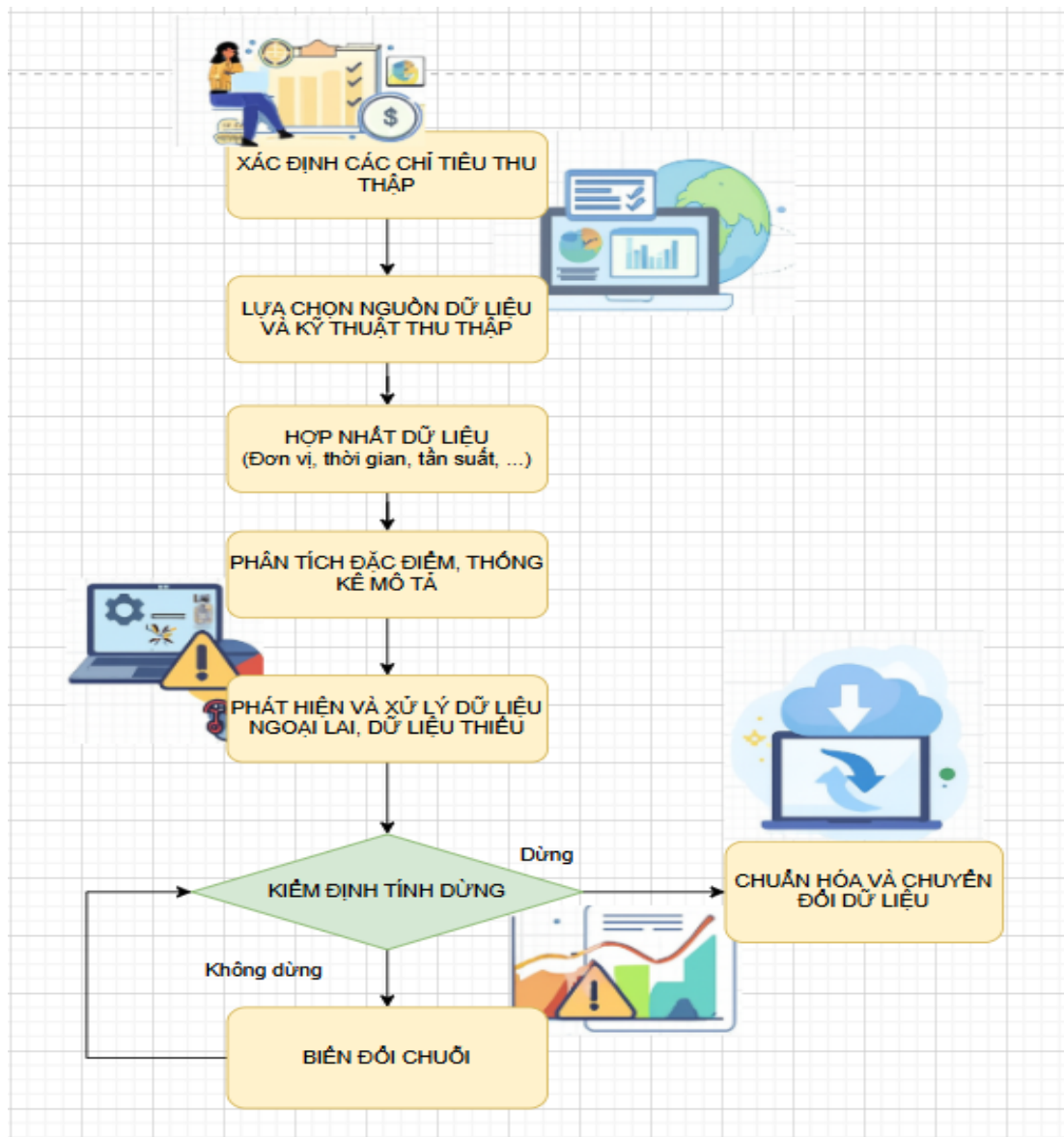
Trong bối cảnh đó, một bộ dữ liệu đa nguồn, chuẩn hóa, có độ tin cậy cao và bao phủ rộng về thời gian – không gian là điều kiện cần thiết để thực hiện các nghiên cứu dự báo và phân tích kinh tế vĩ mô một cách hiệu quả. Bộ dữ liệu này cần đáp ứng các tiêu chí:

- Độ tin cậy cao: lấy từ các nguồn thống kê chính thức và có uy tín quốc tế;
- Tính nhất quán: thống nhất về khái niệm, đơn vị đo và cấu trúc thời gian;
- Độ bao phủ rộng: nhiều quốc gia, nhiều giai đoạn, phục vụ cả nghiên cứu đơn quốc gia và liên quốc gia;
- Khả năng tái sử dụng: có thể áp dụng cho nhiều mô hình và phương pháp phân tích khác nhau, từ thống kê truyền thống đến học máy và học sâu.

Do đó, việc xây dựng bộ dữ liệu đáp ứng các tiêu chí này 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 kinh tế vĩ mô đa nguồn được thực hiện theo các bước chính trong hình 2.1.



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

(1) Xác định các chỉ tiêu thu thập

Danh sách các chỉ tiêu như GDP thực, lạm phát, thất nghiệp, đầu tư, chi tiêu chính phủ, cung tiền (M2), xuất nhập khẩu, v.v. được chọn dựa trên lý thuyết kinh tế vĩ mô về mối quan hệ đồng hành, dẫn dắt hoặc trễ pha trong chu kỳ kinh tế, đồng thời tương thích với khả năng học của mô hình.

(2) Lựa chọn nguồn dữ liệu và kỹ thuật thu thập

Các nguồn dữ liệu chính bao gồm World Bank, OECD, Penn World Table, cục thống kê các quốc gia,... Việc lựa chọn dựa trên tính cập nhật, độ tin cậy và khả

năng cung cấp chuỗi thời gian dài hạn và chi tiết về kinh tế vĩ mô. Các tài liệu như *Guidelines on Integrated Economic Statistics* của UN nhấn mạnh sự cần thiết phải tích hợp đa nguồn để cung cấp dữ liệu đồng bộ và chất lượng[131].

(3) Hợp nhất dữ liệu

Quá trình này cần chuẩn hóa đơn vị (giá cố định/hiện hành, tỷ lệ phần trăm), tần suất thời gian (năm, quý), và định dạng ngày tháng. Bước hợp nhất đòi hỏi xử lý khác biệt về định nghĩa và phân loại giữa các nguồn dữ liệu. Kinh nghiệm trong các hướng dẫn tích hợp dữ liệu chính thức cho thấy việc sử dụng khung tích hợp là cần thiết để tránh trùng lặp và thiếu bao phủ [131].

(4) Phân tích đặc điểm và thống kê mô tả

Một số chỉ tiêu quan trọng được phân tích biểu đồ, thống kê mô tả (xu hướng trung bình, độ lệch chuẩn, tương quan giữa biến, v.v.) để đánh giá sơ bộ các đặc tính như tính ổn định, mùa vụ, và phát hiện bất thường. Đây là bước tiền xử lý quan trọng để xác định biến nào cần chuyển đổi hoặc loại trừ trước khi đưa vào mô hình.

(5) Phát hiện và xử lý dữ liệu ngoại lai và thiếu hụt

Thông qua các phương pháp như boxplot, z-score, clustering, phát hiện điểm ngoại lai, sau đó áp dụng nội suy (interpolation), multiple imputation hoặc xử lý ngoại lai như missing-as-missing hoặc Winsorization. Các nghiên cứu về *multiple imputation* và *robust outlier handling* trong chuỗi thời gian gợi ý các kỹ thuật xử lý phù hợp trong trường hợp ML/DL [27]

(6) Kiểm định tính dừng chuỗi thời gian

Sử dụng các kiểm định thống kê phổ biến như ADF, KPSS, PP để xác định tính dừng của chuỗi. Nếu chuỗi không dừng, thực hiện biến đổi (lấy sai phân). Phương pháp này được áp dụng rộng rãi trong phân tích chuỗi thời gian kinh tế học [87].

(7) Chuẩn hóa và chuyển đổi dữ liệu

Biến được chuẩn hóa theo z-score hoặc min-max, các biến có phân phối lệch như CPI hoặc tỷ lệ thất nghiệp có thể được log-transform hoặc sqrt-transform để giảm ảnh hưởng của ngoại lệ và dữ liệu phi tuyến.

2.3. Các CTKTVM sử dụng trong nghiên cứu

Việc lựa chọn các chỉ tiêu kinh tế vĩ mô cho mô hình dự báo GDP và các chỉ tiêu liên quan được thực hiện dựa trên ba nguyên tắc:

Thứ nhất, cơ sở lý thuyết kinh tế học vĩ mô: Chỉ tiêu được chọn phải đại diện cho các khía cạnh cơ bản của nền kinh tế như tăng trưởng (GDP), giá cả (lạm phát), lao động (tỷ lệ thất nghiệp), thương mại (xuất xuất nhập khẩu), và tài chính – tiền tệ (cung tiền, đầu tư công, chi tiêu chính phủ). Đây là cách tiếp cận tương đồng với phương pháp sử dụng *diffusion index* trong nghiên cứu dự báo chỉ tiêu kinh tế vĩ mô của Stock và Watson [124].

Theo các mô hình tăng trưởng kinh tế cổ điển và hiện đại (Solow-Swan, Mankiw-Romer-Weil, và các biến thể mới), tốc độ tăng trưởng GDP chịu ảnh hưởng của nhiều yếu tố như vốn vật chất, vốn con người, lao động và cơ cấu chi tiêu [97]. Ngoài ra, các chỉ tiêu như lạm phát, thất nghiệp và tỷ giá có quan hệ chặt chẽ với chu kỳ kinh tế và chính sách điều tiết vĩ mô.

Thứ hai, cơ sở thực nghiệm từ các nghiên cứu trước: Nhiều công trình thực nghiệm chỉ ra rằng bộ biến gồm các yếu tố sản xuất, các biến chính sách, và các biến thị trường có khả năng nâng cao độ chính xác dự báo GDP và các chỉ tiêu liên quan [122].

Ưu tiên các chỉ tiêu được công bố định kỳ bởi các tổ chức quốc tế uy tín như World Bank, IMF, PWT hoặc cơ quan thống kê quốc gia.

Thứ ba, tính sẵn có và khả năng chuẩn hóa dữ liệu đa quốc gia: Các chỉ tiêu được chọn đều có dữ liệu công bố từ các tổ chức quốc tế uy tín như World Bank, IMF, PWT hoặc cơ quan thống kê quốc gia. Đảm bảo khả năng đồng bộ hóa, tính so sánh và độ phủ thời gian – không gian cần thiết cho nghiên cứu liên quốc gia.

Trên cơ sở các nguyên tắc đó, luận án xác định 7 nhóm chỉ tiêu cần thu thập, mỗi nhóm tương ứng với một trục động lực chính của quá trình hình thành và biến động GDP [97] với mục tiêu trung tâm là dự báo GDP quốc gia. Các CTKTVM được chọn đáp ứng các mục đích trong bảng 2.1.

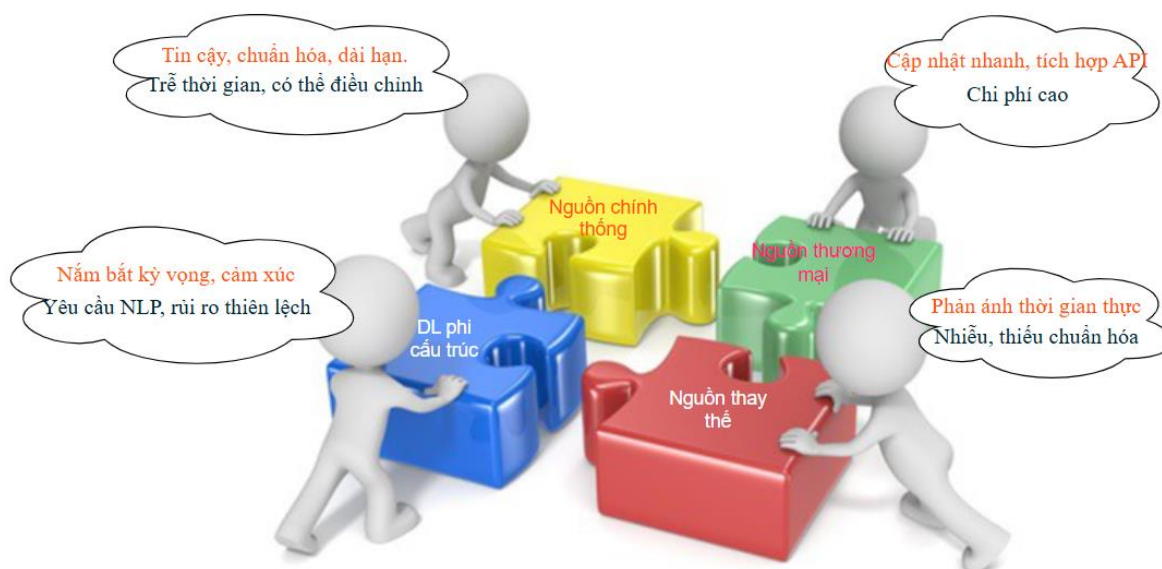
Bảng 2. 1 Các chỉ tiêu kinh tế vĩ mô cần thu thập

STT	Mục đích	Nhóm chỉ tiêu lựa chọn	Các chỉ tiêu lựa chọn
1	Phản ánh trực tiếp sản lượng thực tế và các yếu tố đầu vào trong quá trình sản xuất như lao động, số giờ làm việc, hoặc năng suất lao động. Nhóm chỉ tiêu này cung cấp dữ liệu đầu vào cho việc mô hình hóa động thái tăng trưởng, là thành phần cốt lõi trong phân tích chuỗi thời gian GDP và các chỉ tiêu liên quan.	Đo lường quy mô nền kinh tế, nguồn lực và năng suất	- GDP thực tế theo phía chi tiêu, - GDP thực tế theo phía sản xuất, - Dân số, - Số giờ làm việc trung bình mỗi năm của người lao động, - Chi số vốn con người, - Số người có việc làm, - Tỷ lệ thất nghiệp
2	Đo lường các thành phần chính của GDP từ góc độ cầu (tiêu dùng, đầu tư, tiết kiệm, vốn vật chất), giúp phản ánh chính xác quy mô kinh tế thực và khả năng mở rộng sản xuất.	Tổng sản lượng và tích lũy vốn	- Tiêu dùng thực tế, - Hấp thụ nội địa thực tế, - GDP thực tế theo phía chi tiêu tại mức PPP hiện hành, - GDP thực tế phía sản xuất tại mức PPP hiện hành, - Tổng vốn vật chất của nền kinh tế
3	Cung cấp góc nhìn đa chiều từ sản xuất, thu nhập và tiêu dùng. Các chỉ tiêu thuộc ba cấu phần này giúp mô hình học được mối quan hệ cấu thành GDP từ nhiều chiều cạnh, tăng khả năng dự báo chính xác trong các điều kiện khác nhau.	Chỉ tiêu từ Hệ thống Tài khoản quốc gia (SNA)	- GDP thực tế theo giá quốc gia, - Tiêu dùng thực tế theo giá quốc gia, - Hấp thụ trong nước thực tế theo giá cố định, - Tổng vốn vật chất thực tế theo giá quốc gia, - Dịch vụ vốn theo giá quốc gia, - Năng suất nhân tố tổng hợp (TFP), - TFP có ý nghĩa phúc lợi, - Tỷ trọng thu nhập lao động trong GDP, - Tỷ suất sinh lợi nội tại thực tế của vốn đầu tư, - Tỷ lệ khấu hao trung bình của vốn vật chất hàng năm.
4	Đảm bảo rằng mô hình học được biến động thực của nền kinh tế. Đảm bảo tính đồng nhất của dữ liệu và loại bỏ ảnh hưởng của biến động giá cả, sức mua hoặc tỷ giá hối đoái. Các chỉ tiêu này giúp so sánh và phân tích GDP trên cùng một mặt bằng giá, hỗ trợ nâng cao tính ổn định và khả năng dự báo..	Mức giá / tỷ giá / sức mua	- Tỷ giá hối đoái, - Mức giá tiêu dùng thực tế, - Mức giá của hấp thụ nội địa thực tế, - Mức giá GDP theo phía sản xuất.

STT	Mục đích	Nhóm chỉ tiêu lựa chọn	Các chỉ tiêu lựa chọn
5	Đảm bảo sự nhất quán giữa các thành phần cấu thành GDP và các chỉ tiêu bổ trợ. Việc chuẩn hóa này giúp giảm thiểu sai lệch do khác biệt trong đơn vị đo hoặc phạm vi thống kê, từ đó cải thiện độ tin cậy của mô hình.	Mức giá PPP theo cấu phần chỉ tiêu và vốn	- Mức giá tiêu dùng, - Mức giá đầu tư, - Mức giá chi tiêu chính phủ, - Mức giá xuất khẩu, - Mức giá nhập khẩu, - Mức giá vốn vật chất, - Mức giá dịch vụ.
6	Mô hình học sâu thường hoạt động tốt hơn khi các biến được chuẩn hóa. Các tỷ trọng như tiêu dùng hộ gia đình, đầu tư, xuất khẩu, v.v. trong GDP giúp biểu diễn tương quan kinh tế dưới dạng tương đối, từ đó tăng tính ổn định và khả năng so sánh giữa các quốc gia. Giúp mô hình nhận diện được động lực chính đóng góp vào tăng trưởng GDP.	Tỷ trọng trong GDP theo PPP hiện hành	- Tỷ trọng tiêu dùng hộ gia đình, - Tỷ trọng đầu tư, - Tỷ trọng chi tiêu chính phủ, - Tỷ trọng xuất khẩu hàng hóa, - Tỷ trọng nhập khẩu hàng hóa, - Tỷ trọng sai số thống kê và thương mại ròng
7	Phản ánh tác động của biến động giá cả (CPI) và lạm phát đến sức mua, chi phí sản xuất và tiêu dùng. Đây là yếu tố then chốt để hiệu chỉnh mô hình dự báo, đặc biệt trong bối cảnh các biến động giá cả tác động trực tiếp đến động thái GDP và hành vi kinh tế.	Tiêu dùng và lạm phát	- Chỉ số giá tiêu dùng, - Tỷ lệ lạm phát.

2.4. Nguồn dữ liệu

Việc lựa chọn và sử dụng các nguồn dữ liệu phù hợp là điều tối quan trọng để đảm bảo tính chính xác, độ tin cậy và sự phù hợp của các kết quả nghiên cứu [131]. Nguồn dữ liệu được sử dụng trong các nghiên cứu kinh tế vĩ mô gồm bốn nhóm chính: (i) nguồn dữ liệu chính thống, (ii) nguồn dữ liệu thương mại, (iii) dữ liệu thay thế, và (iv) dữ liệu phi cấu trúc. Mỗi nhóm có những đặc điểm, ưu điểm và hạn chế riêng, ảnh hưởng trực tiếp đến chất lượng mô hình dự báo được thể hiện trong hình 2.2.



Hình 2. 2 Các nguồn dữ liệu

Căn cứ trên bảng các CTKTVM cần thu thập, Penn World Table¹ và World Bank² được luận án lựa chọn làm hai nguồn dữ liệu chính, hai nguồn này được bổ sung bằng dữ liệu từ GSO³, FiinPro⁴... khi cần, nhưng vẫn đảm bảo PWT và WB là lõi dữ liệu để duy trì tính đồng bộ dựa trên các lí do sau:

Độ tin cậy và tính chính thống

- PWT được xây dựng và duy trì bởi các nhóm nghiên cứu uy tín (Groningen Growth and Development Centre, University of California, Davis) và được cộng đồng học thuật quốc tế sử dụng rộng rãi⁵;
- WB (WDI) là nguồn dữ liệu chính thức của Ngân hàng Thế giới, bao quát hàng trăm chỉ tiêu kinh tế – xã hội, thu thập và chuẩn hóa từ các cơ quan thống kê quốc gia và tổ chức quốc tế.

¹ <https://www.rug.nl/ggdc/productivity/pwt/>

² <https://data.worldbank.org/>

³ <https://www.gso.gov.vn>

⁴ <https://www.fiinpro.com>

⁵ <https://doi.org/10.1257/aer.20130954>

Độ bao phủ thời gian và không gian

- PWT cung cấp dữ liệu chuỗi thời gian dài (từ năm 1950) cho hơn 180 quốc gia, đặc biệt mạnh về các chỉ tiêu kinh tế lượng hóa như GDP, vốn vật chất, lao động, vốn con người, và tỷ trọng chi tiêu. PWT cung cấp các chỉ tiêu kinh tế so sánh giữa các quốc gia, được điều chỉnh về sức mua tương đương.
- WB (WDI) bao phủ hơn 200 quốc gia và vùng lãnh thổ, với dữ liệu trải dài nhiều thập kỷ, đáp ứng tốt yêu cầu phân tích liên quốc gia và dài hạn. WB cung cấp hơn 3000 chỉ tiêu phát triển, bao gồm chỉ số phát triển con người, vốn con người, năng suất lao động, cấu trúc kinh tế, và các chỉ tiêu liên quan đến nghèo đói, bất bình đẳng, v.v.

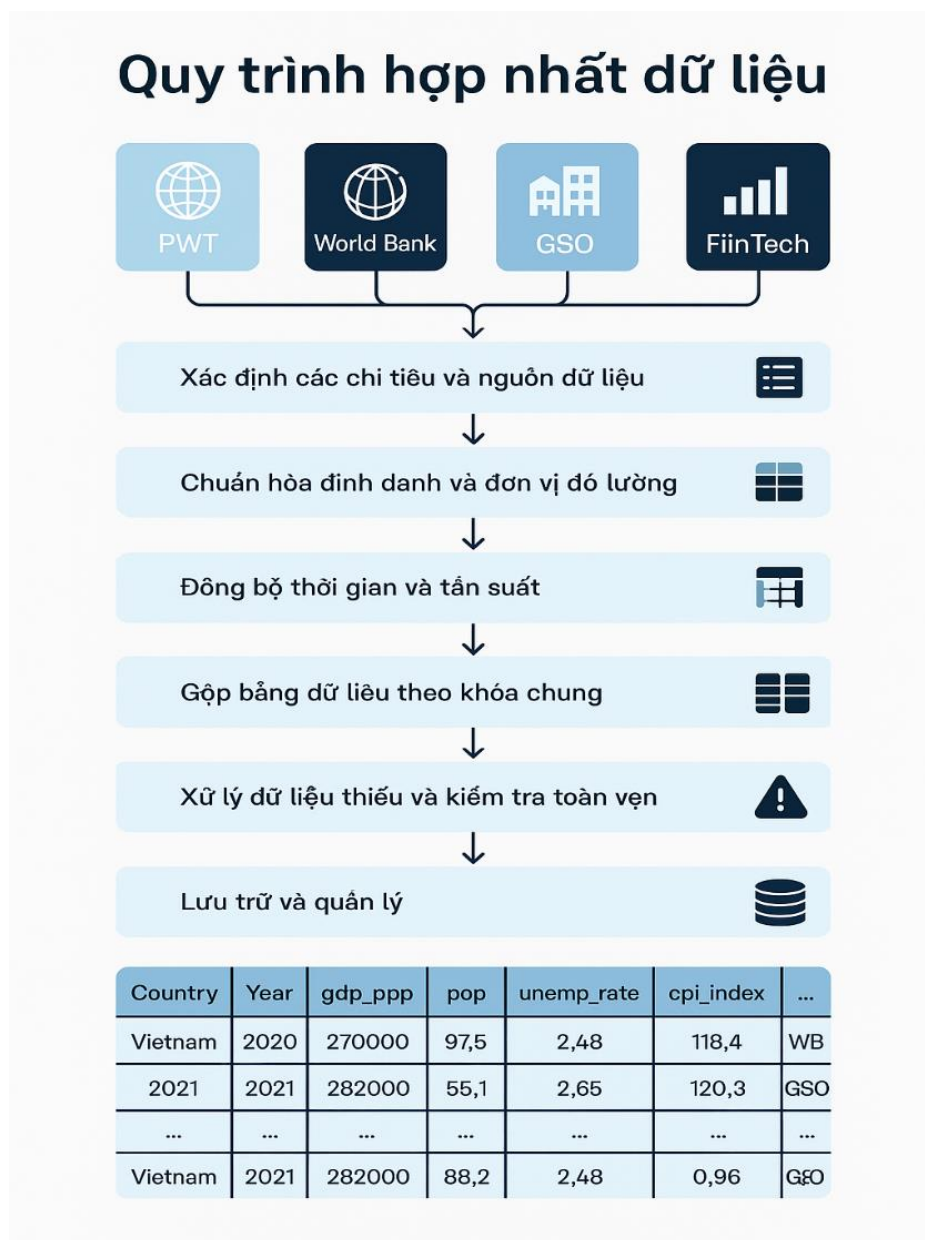
Tính nhất quán và khả năng chuẩn hóa

- PWT áp dụng phương pháp tính thống nhất cho tất cả các quốc gia, đảm bảo khả năng so sánh trực tiếp;
- WB (WDI) cung cấp dữ liệu đã được điều chỉnh về định nghĩa, đơn vị đo và phương pháp thu thập, giảm thiểu sai lệch giữa các quốc gia.

Bên cạnh hai nguồn dữ liệu quốc tế đóng vai trò lõi PWT và WB, nghiên cứu sử dụng thêm các nguồn dữ liệu trong nước là GSO và Fiintech để tăng độ chi tiết, tính cập nhật và khả năng phản ánh đặc thù kinh tế Việt Nam. GSO là cơ quan chủ lực công bố dữ liệu GDP theo quý và năm, chỉ số giá tiêu dùng (CPI), chỉ số sản xuất công nghiệp (IPI), dân số, tỷ lệ thất nghiệp, năng suất lao động, và cơ cấu kinh tế theo ngành của Việt Nam. GSO còn công bố niên giám thống kê và các ấn phẩm chuyên đề. FiinGroup là đơn vị dữ liệu tài chính – kinh tế tư nhân hàng đầu tại Việt Nam, nổi bật với khả năng cập nhật nhanh và tổ chức dữ liệu theo cấu trúc hỗ trợ phân tích chuyên sâu. Dữ liệu của FiinTech bao gồm thông tin chi tiết về hoạt động thương mại, doanh nghiệp niêm yết, chỉ số ngành, thị trường xuất nhập khẩu và các biến vi mô liên quan đến hoạt động sản xuất – kinh doanh.

2.5. Hợp nhất và chuẩn hóa dữ liệu

Mục 2.3 và 2.4 đã xác định các CTKTVM cần thu thập và nguồn thu thập dữ liệu. Quy trình hợp nhất dữ liệu được thực hiện theo các bước trong hình 2.3.



Hình 2. 3 Quy trình hợp nhất dữ liệu

Bước 1: Xác định các chỉ tiêu và nguồn dữ liệu tương ứng (Bảng 2.2)

Bước 2: Chuẩn hóa định danh và đơn vị đo lường

Nội dung thực hiện	Cách thực hiện
Quy đổi đơn vị	Đơn vị tiền tệ: triệu USD Năm cơ sở: 2017 Tỷ lệ: %
Đặt tên biến thống nhất	Theo tên biến trong các bảng phụ lục
Ghi nhận nguồn gốc biến	Thêm cột phụ ghi nguồn gốc

Bước 3: Chuẩn hóa thời gian và tần suất

Nội dung thực hiện	Cách thực hiện	
	Chỉ tiêu	Phương pháp
Chuyển đổi tần suất (về năm)	GDP thực tế theo phía chi tiêu, GDP thực tế theo phía sản xuất, tiêu dùng thực tế, hấp thụ nội địa thực tế, GDP theo PPP hiện hành phía chi tiêu, GDP theo PPP hiện hành phía sản xuất, GDP thực tế theo giá quốc gia, tiêu dùng thực tế theo giá quốc gia, hấp thụ nội địa thực tế theo giá quốc gia, tổng vốn vật chất thực tế theo giá quốc gia.	Lấy tổng
	Dân số, chỉ số vốn con người, tỷ trọng thu nhập lao động trong GDP, tỷ suất sinh lợi nội tại thực tế của vốn đầu tư, năng suất nhân tố tổng hợp, năng suất nhân tố tổng hợp có ý nghĩa phúc lợi.	Lấy giá trị cuối
	Các biến còn lại	Lấy giá trị trung bình

Bước 4: Gộp bảng dữ liệu theo khóa chung

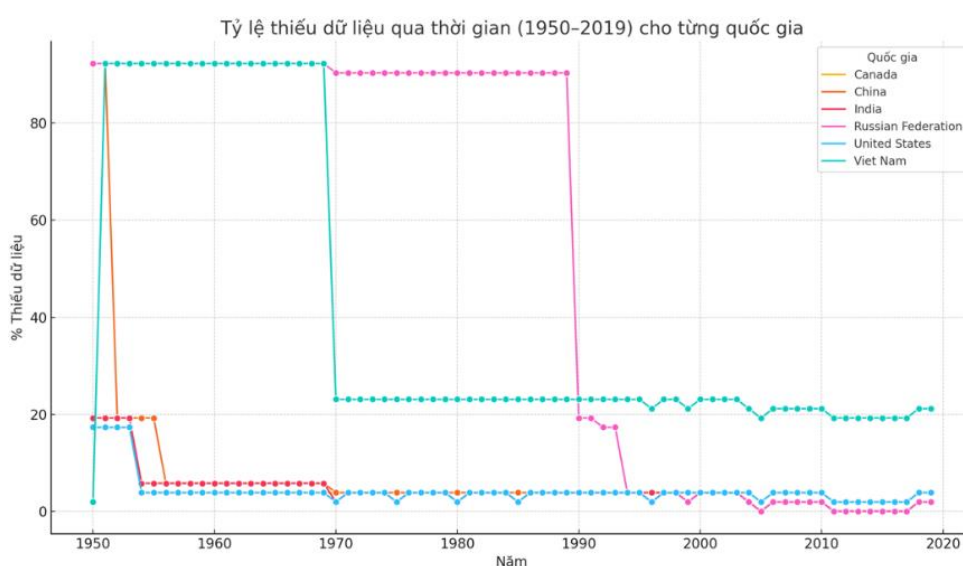
Các bảng dữ liệu được gộp theo khóa chung bao gồm: country (quốc gia), year (năm), và các cột chỉ tiêu thống nhất.

Phương pháp gộp sử dụng là left join tuần tự, trong đó PWT được chọn làm bảng chính do độ ổn định và độ phủ rộng.

Bước 5: Xử lý tính thiếu và kiểm tra toàn vẹn

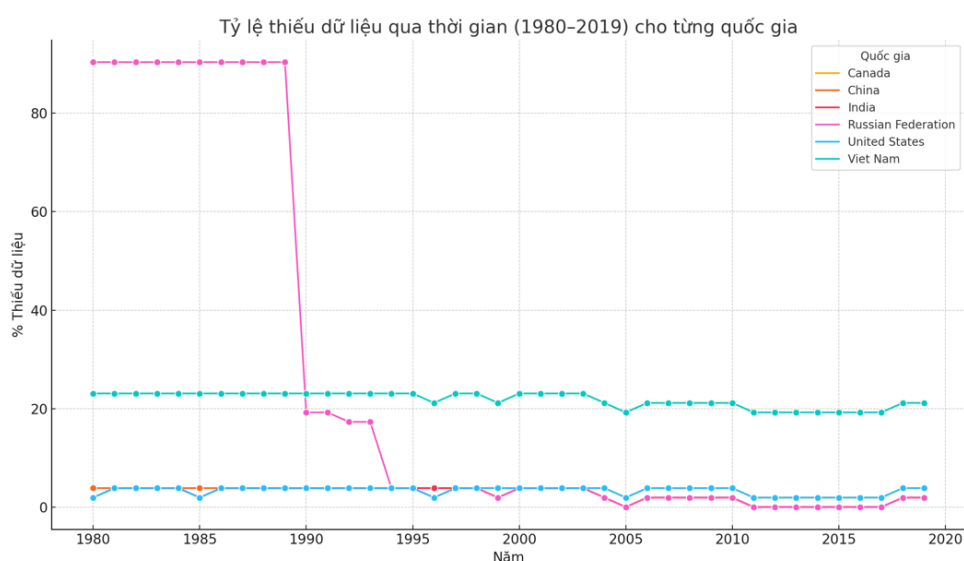
Rà soát trên nhóm 6 quốc gia nghiên cứu (Việt Nam, Nga, Mỹ, Canada, Trung Quốc và Ấn Độ) cho thấy:

- Giai đoạn toàn bộ (1950–2019): Tỷ lệ thiếu dữ liệu trung bình cho 6 quốc gia ở mức 19,12%, nhưng phân bố không đồng đều, được thể hiện trong hình 2.4. Trong đó, Nga là quốc gia có tỷ lệ thiếu dữ liệu cao nhất, lên tới khoảng 54% tổng số giá trị vì lý do Nga không có dữ liệu trước năm 1990, tiếp đến là Việt Nam trên 30% vì Việt Nam không có dữ liệu trước năm 1970, trong khi các quốc gia còn lại dao động trong khoảng 4–7%. Dữ liệu của Việt Nam trước giai đoạn 1980 chủ yếu là dữ liệu ước tính.



Hình 2. 4 Tỷ lệ thiếu dữ liệu giai đoạn 1950-2019

- Giai đoạn 1980–2019: Mức độ thiếu dữ liệu giảm đáng kể, với tỷ lệ trung bình chỉ còn 9,89% cho toàn bộ nhóm quốc gia. Điều này cho thấy từ thập niên 1980 trở đi, dữ liệu trở nên đầy đủ và ổn định hơn cho hầu hết các quốc gia, ngoại trừ Nga – vẫn duy trì tỷ lệ thiếu cao do không có dữ liệu giai đoạn trước 1980. Xu hướng này được minh họa rõ trong Hình 2.5.



Hình 2. 5 Tỷ lệ thiếu dữ liệu giai đoạn 1980-2019

Từ kết quả phân tích trên, nghiên cứu này sẽ tập trung vào giai đoạn 1980–2019 nhằm đảm bảo dữ liệu đầu vào có tính đầy đủ và ổn định cao, đồng thời hạn chế sai lệch do thiếu dữ liệu nghiêm trọng ở các giai đoạn trước. Việc rà soát cho thấy phần lớn dữ liệu thiếu tập trung ở một số biến nhất định, trong khi đa số biến khác có tỷ lệ thiếu rất thấp. Cụ thể, trong giai đoạn này, tỷ lệ thiếu dữ liệu của Canada, Trung Quốc, Ấn Độ và Hoa Kỳ chỉ dao động từ 2,55% đến 3,32%, trong khi Việt Nam có tỷ lệ thiếu cao hơn đáng kể (21,83%) do một số ít biến không được thu thập đầy đủ. Do đó, nghiên cứu sẽ áp dụng tiêu chí loại bỏ một số này nhằm giảm thiểu sai số và nâng cao độ tin cậy của các mô hình dự báo.

Sau khi hợp nhất dữ liệu được bộ dữ liệu gồm các biến cùng với mã và đơn vị trong bảng 2.2 (Ý nghĩa của biến được trình bày trong bảng phụ lục).

Bảng 2. 2 Tổng hợp các chỉ tiêu KTVM sau hợp nhất

<i>STT</i>	<i>Mã chỉ tiêu</i>	<i>Tên chỉ tiêu</i>	<i>Vai trò trong dự báo</i>	<i>Nguồn dữ liệu chính</i>	<i>Nguồn dữ liệu bổ sung</i>	<i>Đơn vị</i>	<i>Ghi chú</i>
1	rgdpe	GDP thực tế theo phía chi tiêu	Biến phụ thuộc chính trong dự báo GDP; biến giải thích cho CPI, thất nghiệp và tỷ giá thông qua mối quan hệ tăng trưởng–ổn định kinh tế vĩ mô	PWT	WB, GSO	Triệu USD (2017)	Đã hiệu chỉnh PPP
2	rgdpo	GDP thực tế theo phía sản xuất	Biến tham chiếu và kiểm chứng kết quả dự báo GDP	PWT	WB, GSO	Triệu USD (2017)	Đã hiệu chỉnh PPP
3	pop	Dân số	Ảnh hưởng đến tăng trưởng GDP tiềm năng, nhu cầu tiêu dùng (CPI) và quy mô thị trường lao động (thất nghiệp)	WB	GSO	Triệu người	
4	Emp	Số người có việc làm	Cơ sở tính tăng trưởng lao động; biến giải thích cho thất nghiệp, GDP và năng suất lao động	WB	GSO	Triệu người	
5	ahv	Số giờ làm việc trung bình mỗi năm của người lao động	tác động đến năng suất, tăng trưởng GDP và thu nhập bình quân	PWT	GSO	Giờ/năm/người	
6	hc	Chỉ số vốn con người	Chỉ tiêu chất lượng lao động; biến giải thích cho năng suất, tăng trưởng GDP và sức cạnh tranh (tác động tới tỷ giá dài hạn)	PWT	WB	Chỉ số	
7	une_r	Tỷ lệ thất nghiệp	Biến phụ thuộc chính trong dự báo thất nghiệp; biến giải thích cho dự báo GDP	WB	GSO	%	
8	ccon	Tiêu dùng thực tế	Thành phần của GDP theo chi tiêu; biến giải thích cho CPI	WB	PWT	Triệu USD (2017)	
9	cda	Hấp thụ nội địa thực tế	Thước đo tổng cầu; biến giải thích GDP và CPI	WB	PWT	Triệu USD (2017)	

<i>STT</i>	<i>Mã chỉ tiêu</i>	<i>Tên chỉ tiêu</i>	<i>Vai trò trong dự báo</i>	<i>Nguồn dữ liệu chính</i>	<i>Nguồn dữ liệu bổ sung</i>	<i>Đơn vị</i>	<i>Ghi chú</i>
10	cgdpe	GDP thực tế theo phía chỉ tiêu tại mức PPP hiện hành	Biến tham chiếu hỗ trợ phân tích sức mua và so sánh quốc tế; dùng trong các mô hình GDP và phân tích sức mạnh tiền tệ (tỷ giá)	PWT	WB	Triệu USD (2017)	PPP hiện hành
11	cgdpo	GDP thực tế phía sản xuất tại mức PPP hiện hành	Biến tham chiếu bổ sung cho GDP; hỗ trợ đánh giá cơ cấu sản xuất và ảnh hưởng tới cán cân thương mại, tỷ giá	PWT	WB	Triệu USD (2017)	PPP hiện hành
12	cn	Tổng vốn vật chất của nền kinh tế,	Biến giải thích cho GDP, năng suất lao động, và gián tiếp tác động tới CPI, tỷ giá qua năng lực sản xuất	PWT	WB	Triệu USD (2017)	
13	rgdpna	GDP thực tế theo giá quốc gia	Biến phụ thuộc/giải thích khi phân tích tác động của lạm phát nội địa (CPI) và tỷ giá danh nghĩa	PWT	WB	Triệu USD (2017)	
14	rconna	Tiêu dùng thực tế theo giá quốc gia	Biến giải thích chính cho CPI và tổng cầu; đồng thời ảnh hưởng tới GDP qua kênh tiêu dùng	PWT	WB	Triệu USD (2017)	
15	rdana	Hấp thụ trong nước thực tế theo giá cố định	Đo lường tổng cầu nội địa đã loại bỏ yếu tố giá; giải thích GDP và tác động tới tỷ giá qua nhập khẩu ròng	PWT	WB	Triệu USD (2017)	
16	rnna	Tổng vốn vật chất thực tế theo giá quốc gia	Biến phản ánh năng lực sản xuất và đầu tư; ảnh hưởng dài hạn tới GDP, CPI và tỷ giá	PWT	WB	Triệu USD (2017)	
17	rkna	Dịch vụ vốn theo giá quốc gia	Phản ánh chi phí vốn, tác động tới đầu tư, tăng trưởng GDP và kỳ vọng tỷ giá	PWT	WB	Triệu USD (2017)	
18	rtfpna	Năng suất nhân tố tổng hợp(TFP)	Biến giải thích cốt lõi cho tăng trưởng GDP dài hạn; ảnh hưởng tới cạnh tranh quốc tế và tỷ giá thực	PWT	WB	Chỉ số	
19	rwtfpna	TFP có ý nghĩa phúc lợi	Biến giải thích cho tăng trưởng GDP gắn với phúc lợi xã hội;	PWT	WB	Chỉ số	
20	labsh	Tỷ trọng thu nhập lao động trong GDP (giá hiện hành)	Biến giải thích cho phân phối thu nhập, sức mua và tác động tới tiêu dùng, thất nghiệp	PWT	WB	%	

<i>STT</i>	<i>Mã chỉ tiêu</i>	<i>Tên chỉ tiêu</i>	<i>Vai trò trong dự báo</i>	<i>Nguồn dữ liệu chính</i>	<i>Nguồn dữ liệu bổ sung</i>	<i>Đơn vị</i>	<i>Ghi chú</i>
21	irr	Tỷ suất sinh lợi nội tại thực tế của vốn đầu tư	Biến đánh giá hiệu quả đầu tư, ảnh hưởng tới kỳ vọng tăng trưởng GDP và tỷ giá	PWT	WB	%	
22	delta	Tỷ lệ khấu hao trung bình của vốn vật chất hàng năm.	Ảnh hưởng tới năng lực sản xuất, chi phí vốn, và gián tiếp tác động tới GDP, CPI	PWT	WB	%	
23	xr	Tỷ giá hối đoái	Biến phụ thuộc hoặc giải thích trong dự báo tỷ giá; ảnh hưởng trực tiếp tới CPI và gián tiếp tới GDP qua thương mại	WB		Nội tệ/USD	
24	pl_con	Mức giá tiêu dùng thực tế (CCON)	Biến phụ thuộc khi dự báo CPI; biến giải thích cho tiêu dùng và tổng cầu trong dự báo GDP	PWT	WB	%	
25	pl_da	Mức giá của hấp thụ nội địa thực tế (CDA)	Đo lường mức giá tổng cầu nội địa; ảnh hưởng tới CPI và cán cân thương mại	PWT	WB	%	
26	pl_gdp_o	Mức giá GDP theo phía sản xuất (CGDPo)	Chỉ báo lạm phát từ phía cung; ảnh hưởng tới GDP thực và tỷ giá thực	PWT	WB	%	
27	pl_c	Mức giá tiêu dùng	Biến chính để dự báo CPI; tác động trực tiếp tới sức mua, tiêu dùng và tăng trưởng GDP	PWT	WB	%	
28	pl_i	Mức giá đầu tư	Ảnh hưởng tới quyết định đầu tư, năng lực sản xuất, từ đó tác động tới GDP, CPI và tỷ giá	PWT	WB	%	
29	pl_g	Mức giá chi tiêu chính phủ	Ảnh hưởng tới chi tiêu công thực, từ đó tác động tới GDP, CPI và cân đối ngân sách	PWT	WB	%	
30	pl_x	Mức giá xuất khẩu	Ảnh hưởng trực tiếp tới cán cân thương mại, tỷ giá và GDP; gián tiếp tới CPI qua giá nhập khẩu	PWT	WB	%	
31	pl_m	Mức giá nhập khẩu	Biến giải thích quan trọng cho CPI; tác động tới cán cân thương mại và GDP	PWT	WB	%	

<i>STT</i>	<i>Mã chỉ tiêu</i>	<i>Tên chỉ tiêu</i>	<i>Vai trò trong dự báo</i>	<i>Nguồn dữ liệu chính</i>	<i>Nguồn dữ liệu bổ sung</i>	<i>Đơn vị</i>	<i>Ghi chú</i>
32	pl_n	Mức giá vốn vật chất	Ảnh hưởng tới chi phí đầu tư và năng lực sản xuất; tác động dài hạn tới GDP và lạm phát	PWT	WB	%	
33	pl_k	Mức giá dịch vụ	Chỉ báo lạm phát trong khu vực dịch vụ; ảnh hưởng tới CPI và GDP	PWT	WB	%	
34	csn_c	Tỷ trọng tiêu dùng hộ gia đình	Phản ánh cầu tiêu dùng; biến giải thích GDP, CPI	PWT	WB	%	
35	csn_i	Tỷ trọng đầu tư	Động lực tăng trưởng dài hạn, biến giải thích GDP	PWT	WB	%	
36	csn_g	Tỷ trọng chi tiêu chính phủ	đánh giá tác động của chi tiêu công tới GDP và CPI	PWT	WB	%	
37	csn_x	Tỷ trọng xuất khẩu hàng hóa	Ảnh hưởng trực tiếp tới GDP qua thành phần xuất khẩu ròng và tỷ giá	PWT	WB	%	
38	csn_m	Tỷ trọng nhập khẩu hàng hóa	Biến giải thích GDP, CPI và cán cân thương mại	PWT	WB	%	
39	csn_r	Tỷ trọng sai số thống kê và thương mại ròng	Ảnh hưởng gián tiếp tới GDP và CPI	PWT	WB	%	
40	cpi	Chỉ số giá tiêu dùng	Biến phụ thuộc khi dự báo CPI; biến giải thích trong dự báo GDP và lạm phát	WB	GSO, FiinGroup	%	
41	inf_r	Tỷ lệ lạm phát	Biến phụ thuộc hoặc giải thích trong dự báo lạm phát; tác động tới GDP, CPI, tỷ giá	WB	GSO, FiinGroup	%/năm	

2.6. Kiểm tra chất lượng và chuẩn bị dữ liệu

2.6.1. Phân tích đặc điểm bộ dữ liệu

1. Thông tin chung:

- Số chỉ tiêu: 41 trong bảng 2.2 (Ý nghĩa của các chỉ tiêu này được mô tả trong bảng phụ lục PL1 – PL7)
- Loại dữ liệu: Tất cả các biến đều thuộc loại số thực
- Phạm vi: Mỹ, Canada, Nga, Ấn Độ, Trung Quốc và Việt Nam.
- Khoảng thời gian: Năm 1950 đến 2019. Ngoại trừ một số quốc gia như Nga từ 1990-2019, Việt Nam từ 1970 đến 2019.

2. Thiếu dữ liệu

Sau khi đã kiểm tra thiếu và tính toàn vẹn của bộ dữ liệu trong bước 2.5 Hợp nhất dữ liệu thì:

- Tỷ lệ thiếu trên toàn bộ tập dữ liệu thấp: chỉ khoảng 1.9%–3.8% ở một số biến, và không đồng đều ở các quốc gia.
- Các biến có dữ liệu thiếu gồm: pl_k, ck, rwtfpna, rtfpna, rkna, labsh, irr tùy thuộc vào từng quốc gia.
- Các dữ liệu thiếu chủ yếu nằm ở giai đoạn trước năm 1990

3. Đặc điểm cấu trúc

Tất cả các biến có dạng chuỗi thời gian đơn biến, được sắp xếp theo thứ tự thời gian đều đặn, thuận lợi cho việc chuẩn hóa và huấn luyện mô hình học sâu tuần tự.

Dữ liệu phản ánh đa chiều về sản lượng (GDP thực), lao động (dân số, số người có việc làm, giờ làm), giá cả (CPI, deflator), chi tiêu (tiêu dùng, đầu tư, chính phủ), thương mại (xuất, nhập khẩu), và các tỷ trọng kinh tế (tỷ trọng đầu tư, chi tiêu hộ,...).

2.6.2. Phân tích thống kê mô tả

Bộ dữ liệu sử dụng trong nghiên cứu bao gồm các quan sát hàng năm với các biến phản ánh các khía cạnh kinh tế vĩ mô như sản lượng, tiêu dùng, lao động, giá cả và cơ cấu chi tiêu, v.v. Mục kiểm tra chất lượng và chuẩn hóa bộ dữ liệu này được luận án thực hiện cho từng quốc gia nhằm đánh giá phân phối dữ liệu, tính biến động và mức độ tuân theo phân phối chuẩn. Từ phần này, luận án thực hiện với bộ dữ liệu của Việt Nam, dữ liệu của các quốc gia còn lại thực hiện tương tự.

Bảng 2.2 ở dưới là thống kê mô tả của bộ dữ liệu KTVM Việt Nam với ý nghĩa của các chỉ tiêu được mô tả ở phụ lục. Kết quả thống kê cho thấy GDP thực theo hướng tiếp cận tiêu dùng (rgdpe) và sản lượng (rgdpo) lần lượt đạt trung bình khoảng 211.457 triệu USD và 218.746 triệu USD (giá năm 2017), với độ lệch chuẩn rất cao (~198.000–199.000 triệu USD), phản ánh xu hướng tăng trưởng rõ rệt và biến động đáng kể theo thời gian. Cả hai biến này đều có hệ số lệch (skewness) dương lớn (trên 1.1), cho thấy phân phối lệch phải, tức là có một số năm tăng đột biến. Hệ số nhọn (kurtosis) ở mức trung bình cho thấy dữ liệu có phân phối gần chuẩn. Tuy nhiên, kiểm định Jarque-Bera (JB) khẳng định rằng hai biến không tuân theo phân phối chuẩn ở mức ý nghĩa 5% ($p < 0.01$).

Các biến nhân khẩu học như dân số (pop) và lao động (emp) có phân phối ổn định hơn. Trung bình dân số đạt khoảng 71.63 triệu người và số lao động trung bình là 33.99 triệu người. Cả hai biến có skewness gần 0 và kiểm định JB không bác bỏ giả thuyết phân phối chuẩn ($p > 0.15$), cho thấy chúng có thể phù hợp với các mô hình tuyến tính hoặc hồi quy cổ điển.

Biến vốn con người (hc) có giá trị trung bình là 1.99 với độ lệch chuẩn 0.392, cho thấy xu hướng tăng dần theo thời gian. Tuy nhiên, kiểm định JB cho thấy phân phối lệch phải nhẹ và có khác biệt so với phân phối chuẩn ở mức ý nghĩa 5% ($p = 0.04199$), gợi ý rằng nên chuẩn hóa hoặc biến đổi log biến này khi đưa vào mô hình học máy.

Các chỉ tiêu như tỷ trọng chi tiêu hộ gia đình (csh_c), đầu tư (csh_i), chính phủ (csh_g), xuất khẩu (csh_x) và nhập khẩu (csh_m) cũng được thống kê cho thấy sự phân tán cao, đặc biệt là biến csh_x (tỷ trọng xuất khẩu) với độ lệch chuẩn lớn, phản ánh mức độ hội nhập kinh tế gia tăng nhanh chóng của Việt Nam trong giai đoạn đổi mới và hội nhập quốc tế.

Nhìn chung, phần lớn các biến kinh tế vĩ mô trong bộ dữ liệu có phân phối lệch và/hoặc nhọn, do đó trong các bước xử lý tiếp theo, cần chuẩn hóa để làm mượt dữ liệu, đồng thời tăng hiệu quả huấn luyện mô hình học sâu và giảm nguy cơ overfitting. Việc nắm rõ đặc điểm thống kê của từng biến không chỉ giúp lựa chọn mô hình dự báo phù hợp mà còn hỗ trợ việc thiết kế các bước tiền xử lý dữ liệu một cách khoa học và có hệ thống.

Bảng 2. 3 Bảng tham số thống kê các chỉ số KTVM Việt Nam

Chỉ tiêu	Trung bình	Trung vị	Lớn nhất	Nhỏ nhất	Độ lệch chuẩn	Skewness	Kurtosis	Xác suất	JB
rgdpe	211456.7087	127008.1524	750726.7500	36863.5859	198345.0211	1.2534	0.4603	13.5341	0.0012
rgdpo	218745.7755	129262.8633	724123.3750	40070.1016	199569.0782	1.1125	-0.0086	10.3132	0.0058
pop	71.6255	74.2808	96.4621	43.4048	16.2652	-0.2093	-1.2635	3.6909	0.1580
emp	33.9870	32.3319	54.0691	15.5761	12.3937	0.1998	-1.2581	3.6303	0.1628
avh	2398.4008	2348.2380	2986.8475	2131.9682	207.0952	1.1537	0.7116	12.1466	0.0023
hc	1.9942	1.8036	2.8700	1.4942	0.3921	0.8176	-0.6077	6.3403	0.0420
une_r	1.92	1.97	2.87	1	0.5	-0.3019	-0.2819	0.8083	0.4255
ccon	174970.4129	122552.5001	582677.0625	38076.7734	147274.3600	1.2738	0.6128	14.3042	0.0008
cda	223259.8556	140294.9688	758821.9375	40851.3438	200564.9518	1.1541	0.2051	11.1880	0.0037
cgdpe	213479.2722	129090.1875	747853.7500	36742.9141	199318.3154	1.2139	0.3418	12.5238	0.0019
cgdpo	219944.1977	131795.2070	723142.6875	40268.1133	199385.2152	1.0914	-0.0419	9.9296	0.0070
cn	493715.3188	134752.7110	2009263.3750	34543.2461	588372.0109	1.0712	-0.2267	9.6690	0.0080
rgdpna	232252.5681	151592.4609	741653.5625	42713.8867	198575.7993	1.0117	-0.1209	8.5600	0.0138
rconna	187187.9384	130029.0508	582796.0000	42934.9727	147034.8719	1.1201	0.2361	10.5715	0.0051

Chỉ tiêu	Trung bình	Trung vị	Lớn nhất	Nhỏ nhất	Độ lệch chuẩn	Skewness	Kurtosis	Xác suất	JB
rdana	225370.4533	141703.9375	762540.8750	40080.1875	201415.7687	1.1136	0.1355	10.3719	0.0056
rnna	489921.7185	199546.1172	2011976.0000	52886.7422	549285.7857	1.2591	0.4220	13.5828	0.0011
delta	0.0455	0.0394	0.0632	0.0354	0.0095	0.6536	-1.2732	6.9365	0.0312
xr	9776.2904	11035.4167	23050.2417	2.9932	8505.7141	0.0421	-1.5269	4.8720	0.0875
pl_con	0.1444	0.1239	0.3358	0.0215	0.1118	0.5736	-1.1628	5.5585	0.0621
pl_da	0.1544	0.1429	0.3526	0.0230	0.1164	0.5013	-1.2088	5.1385	0.0766
pl_gdp o	0.1493	0.1400	0.3622	0.0210	0.1143	0.5842	-1.0259	5.0368	0.0806
cs_h_c	0.7287	0.7192	0.9307	0.5578	0.0990	0.1185	-1.0035	2.2149	0.3304
cs_h_i	0.1573	0.1331	0.3217	0.0645	0.0883	0.4664	-1.3009	5.3380	0.0693
cs_h_g	0.1442	0.1515	0.1915	0.0888	0.0329	-0.4126	-1.1801	4.3199	0.1153
cs_h_x	0.1523	0.0976	0.5260	0.0002	0.1657	0.8738	-0.4359	6.7579	0.0341
cs_h_m	-0.1617	-0.0930	-0.0012	-0.5191	0.1657	-0.7880	-0.7230	6.2634	0.0436
cs_h_r	-0.0208	-0.0128	0.1155	-0.2374	0.0725	-0.7025	0.5432	4.7269	0.0941
pl_c	0.1563	0.1397	0.3519	0.0238	0.1167	0.4902	-1.2638	5.3306	0.0696
pl_i	0.2091	0.2282	0.4104	0.0440	0.1236	0.1445	-1.3532	3.9889	0.1361
pl_g	0.0819	0.0614	0.2451	0.0093	0.0790	0.9846	-0.4846	8.5672	0.0138
pl_x	0.4038	0.4201	0.6956	0.1201	0.1804	-0.1139	-1.0991	2.6247	0.2692
pl_m	0.4319	0.4815	0.7126	0.1373	0.1823	-0.3833	-1.1532	3.9951	0.1357
pl_n	0.1500	0.1552	0.2562	0.0395	0.0726	-0.0736	-1.3817	4.0223	0.1338
cpi	89.39	82.3	145.3	46.1	31.66	0.426	-1.1265	0.3845	1.9118
inf_r	6.93	6.6	23.1	-1.8	5.91	1.1902	1.2348	0.0319	6.8914

Nguồn: Tác giả

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

Phát hiện ngoại lệ

Dữ liệu kinh tế vĩ mô thường chứa các điểm bất thường do các nguyên nhân như sai số đo lường, lỗi nhập liệu, thay đổi chính sách đột ngột, hoặc các cú sốc kinh tế như khủng hoảng tài chính và đại dịch. Việc phát hiện và xử lý các giá trị ngoại lệ này là điều kiện bắt buộc để tránh làm sai lệch các mô hình học máy hoặc các phân tích thống kê truyền thống.

Có nhiều phương pháp được sử dụng để phát hiện ngoại lệ. Với dữ liệu có phân phối chuẩn, kỹ thuật Z-score được sử dụng để xác định các điểm cách xa trung bình nhiều độ lệch chuẩn. Đối với dữ liệu có phân phối lệch, phương pháp dựa trên khoảng tứ phân vị (IQR) là phù hợp hơn: các điểm nằm ngoài khoảng $[Q1 - 1.5 \times IQR, Q3 + 1.5 \times IQR]$ được coi là ngoại lệ. Đối với dữ liệu đa biến và phi tuyến, các thuật toán học máy như Isolation Forest [94] hoặc One-Class SVM, cũng như phương pháp mật độ như Local Outlier Factor (LOF) [32], được sử dụng để phát hiện các điểm bất thường mà không cần giả định phân phối.

Trong bảng thống kê mô 2.2 phần lớn các chỉ tiêu có phân phối lệch nên luận án đã tiến hành phát hiện và xử lý ngoại lệ thông qua phương pháp khoảng tứ phân vị (Interquartile Range – IQR), kết quả phân tích cho thấy một số biến kinh tế vĩ mô có xuất hiện các giá trị ngoại lệ nằm ngoài ngưỡng ± 1.5 lần độ rộng của khoảng tứ phân vị. Cụ thể, bảng 2.3 Rgdpe có ba quan sát vượt trội trong giai đoạn 2017–2019, phản ánh mức tăng trưởng kinh tế đột biến của Việt Nam trong những năm gần đây.

Bảng 2. 4 Các ngoại lệ

Năm	Giá trị rgdpe (USD, giá cố định)
2017	653,410.81
2018	705,553.38
2019	750,726.75

Ngoài ra, Biến csh_r (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. Trong một số trường hợp, giá trị ngoại lệ có thể được mã hóa lại thành biến chỉ thị nhằm làm rõ bối cảnh phân tích. Thay vì loại bỏ các quan sát ngoại lệ, luận án lựa chọn tiếp cận theo hướng bảo toàn thông tin kinh tế, đồng thời tiến hành chuẩn hóa toàn bộ các biến định lượng bằng phương pháp Z-score.

Xử lý giá trị bị thiếu

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.

2.6.4. Kiểm định chuỗi thời gian

Việc đánh giá tính dừng của các chuỗi thời gian kinh tế vĩ mô là bước quan trọng nhằm đảm bảo độ chính xác và hiệu quả trong quá trình dự báo. Một chuỗi không dừng thường dẫn đến sai lệch trong ước lượng, kết luận thống kê thiếu tin cậy, và làm giảm hiệu quả của các mô hình như ARIMA, VAR hoặc các kiến trúc học sâu. Do đó, trong nghiên cứu này, kiểm định tính dừng được tiến hành bằng cách kết hợp đồng thời hai phương pháp phổ biến là ADF và KPSS [87].

Sử dụng kết hợp hai kiểm định là cần thiết vì chúng có giả thuyết gốc trái ngược nhau: ADF kiểm tra giả thuyết H_0 là chuỗi không dừng (có gốc đơn vị), trong khi KPSS kiểm tra giả thuyết H_0 là chuỗi dừng. Cách tiếp cận này giúp kiểm tra chéo và tăng cường độ tin cậy của kết luận. Ngoài ra, KPSS thường nhạy hơn trong việc phát hiện các xu hướng trôi dạt hoặc phương sai thay đổi – những đặc điểm phổ biến

trong dữ liệu kinh tế vĩ mô. Khi cả hai kiểm định đều cho kết luận rõ ràng (ví dụ ADF bác bỏ H_0 và KPSS không bác bỏ H_0), ta có thể xác nhận chuỗi là dừng ($I(0)$); ngược lại, nếu ADF không bác bỏ H_0 và KPSS bác bỏ H_0 , chuỗi được xem là không dừng và cần sai phân ($I(1)$). Trường hợp cả hai kiểm định đều không rõ ràng có thể chỉ ra rằng chuỗi là $I(2)$ – tức là cần sai phân hai lần để đạt được trạng thái dừng.

Bảng 2.4 là kết quả kiểm tra tính dừng của bộ chỉ số KTVM theo phương pháp kiểm định Dickey – Fuller được tăng cường (ADF) và kiểm định KPSS. Kết quả được trình bày dưới dạng bậc tích phân của chuỗi, ký hiệu là $I(d)$, trong đó d là số lần cần sai phân để đạt được chuỗi dừng.

Bảng 2. 5 Bảng kiểm định tính dừng của một số biến KTVM

Biến	ADF Statistic	ADF p-value	KPSS Statistic	KPSS p-value	Kiểu dừng	Mức ý nghĩa
rgdpe	2.9900	1.0000	0.9538	0.0100	$I(1)$	5%
rgdpo	10.4627	1.0000	0.9681	0.0100	$I(1)$	5%
pop	-3.4832	0.0084	1.0990	0.0100	$I(2)$	5%
emp	-3.7102	0.0040	1.0841	0.0100	$I(2)$	5%
avh	-3.3128	0.0143	0.9087	0.0100	$I(2)$	5%
hc	1.1168	0.9954	1.0068	0.0100	$I(1)$	5%
une_r	1.8597	0.3513	0.4313	0.0637	$I(2)$	5%
ccon	17.2921	1.0000	0.9588	0.0100	$I(1)$	5%
cda	4.2802	1.0000	0.9747	0.0100	$I(1)$	5%
cgdpe	2.2329	0.9989	0.9605	0.0100	$I(1)$	5%
cgdpo	11.7950	1.0000	0.9743	0.0100	$I(1)$	5%
cn	4.4656	1.0000	0.9411	0.0100	$I(1)$	5%
rgdpna	3.2355	1.0000	1.0039	0.0100	$I(1)$	5%
rconna	18.0425	1.0000	0.9939	0.0100	$I(1)$	5%
rdana	4.2811	1.0000	0.9868	0.0100	$I(1)$	5%
rnna	0.8356	0.9922	0.9373	0.0100	$I(1)$	5%
delta	-1.3178	0.6210	0.7951	0.0100	$I(1)$	5%
xr	0.1177	0.9673	1.0591	0.0100	$I(1)$	5%

Biến	ADF Statistic	ADF p-value	KPSS Statistic	KPSS p-value	Kiểu dùng	Mức ý nghĩa
pl_con	0.4229	0.9823	1.0340	0.0100	I(1)	5%
pl_da	0.2180	0.9733	1.0445	0.0100	I(1)	5%
pl_gdpo	0.7951	0.9916	1.0424	0.0100	I(1)	5%
cs_h_c	-1.4175	0.5739	0.6354	0.0194	I(1)	5%
cs_h_i	-1.2245	0.6630	0.9210	0.0100	I(1)	5%
cs_h_g	-0.1602	0.9431	0.7256	0.0112	I(1)	5%
cs_h_x	3.0471	1.0000	1.0093	0.0100	I(1)	5%
cs_h_m	2.1980	0.9989	0.9923	0.0100	I(1)	5%
cs_h_r	-2.2203	0.1990	0.1929	0.1000	I(1)	5%
pl_c	0.5041	0.9850	1.0453	0.0100	I(1)	5%
pl_i	-0.5948	0.8722	1.0272	0.0100	I(1)	5%
pl_g	2.4076	0.9990	0.9546	0.0100	I(1)	5%
pl_x	-0.5581	0.8802	1.0094	0.0100	I(1)	5%
pl_m	-0.9421	0.7738	0.9823	0.0100	I(1)	5%
pl_n	-0.9269	0.7790	0.9812	0.0100	I(1)	5%
cpi	3.70624	1	0.6692	0.0163	I(1)	5%
inf_r	-3.9462	0.0017	0.19760	0.1	I(0)	5%

Nguồn: Tác giả

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 (cs_h_c, cs_h_i), và mức giá tương đối (p pl_c, pl_g), hc, v.v... đều không dừng tại cấp độ gốc (p-value ADF > 0.05 và p-value KPSS < 0.05). Kiểu dùng I(1), tức là cần sai phân bậc nhất để trở nên dừng. Điều này là phổ biến trong các chuỗi kinh tế vĩ mô dài hạn, vốn chịu ảnh hưởng của xu hướng tăng trưởng hoặc các yếu tố cấu trúc không ổn định theo thời gian.

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.

Việc các chuỗi không dừng là điều thường thấy trong dữ liệu kinh tế vĩ mô, phản ánh các xu hướng dài hạn như tăng trưởng dân số, tích lũy vốn hoặc cải thiện năng suất lao động. Tuy nhiên, để đảm bảo mô hình hóa chính xác, các chuỗi không dừng cần được xử lý thông qua các kỹ thuật như sai phân bậc nhất, biến đổi log hoặc lọc xu hướng. Những kỹ thuật này giúp ổn định phương sai và kỳ vọng của chuỗi, từ đó cải thiện tính hiệu quả và độ tin cậy của các 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.

2.6.5. Chuyển đổi dữ liệu

Chuẩn hóa dữ liệu không chỉ cải thiện hiệu suất của các thuật toán lựa chọn đặc trưng mà còn tăng cường khả năng khái quát hóa của mô hình học máy, đặc biệt trong các bài toán có số chiều lớn hoặc dữ liệu không đồng nhất [64]. Xuất phát từ đặc điểm của dữ liệu kinh tế vĩ mô và yêu cầu kỹ thuật của các mô hình học sâu được triển khai, 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 với hai lý do chính:

Thứ nhất, dữ liệu kinh tế vĩ mô thường có sự khác biệt đáng kể về đơn vị đo lường và độ lớn giữa các biến. Ví dụ: GDP tính theo triệu USD, tỷ lệ thất nghiệp theo phần trăm, chỉ số giá tiêu dùng là chỉ số tổng hợp không có đơn vị đo cụ thể. Sự không đồng nhất này có thể gây sai lệch trong quá trình tối ưu hóa trọng số, đặc biệt là trong các mô hình nhạy cảm với quy mô như mạng nơ-ron hồi tiếp (RNN, LSTM) hoặc các thuật toán có sử dụng gradient descent. Nếu không được chuẩn hóa, các biến có giá trị lớn hơn có xu hướng chi phối quá trình học, làm giảm hiệu quả của mô hình và tăng nguy cơ hội tụ chậm hoặc rơi vào cực trị địa phương [116].

Thứ hai, chuẩn hóa Z-score – tức đưa mỗi biến về dạng có trung bình bằng 0 và độ lệch chuẩn bằng 1 – có ưu điểm nổi bật là duy trì hình dạng phân phối gốc của

dữ liệu, trong khi vẫn đảm bảo loại bỏ ảnh hưởng của quy mô tuyệt đối. Phương pháp này không bị ảnh hưởng bởi giá trị biên như chuẩn hóa Min-Max, và do đó thích hợp hơn trong các trường hợp dữ liệu có phân phối gần chuẩn hoặc có chứa giá trị ngoại lệ ở mức vừa phải.

Trong luận án, quá trình chuẩn hóa Z-score được thực hiện theo công thức:

$$z_i = \frac{x_i - \mu}{\sigma} \quad (2.1)$$

Trong đó:

- x_i là giá trị gốc của biến;
- μ là trung bình mẫu của biến;
- σ là độ lệch chuẩn mẫu của biến;
- và z_i là giá trị đã được chuẩn hóa.

Phép biến đổi này được áp dụng riêng biệt cho từng biến, sử dụng thông tin thống kê (trung bình và độ lệch chuẩn) tính từ mục 2.5 (Bảng thống kê mô tả) nhằm đảm bảo tính nhất quán và tránh rò rỉ dữ liệu. 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.

2.7. 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 quốc gia 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. CÁC MÔ HÌNH DỰ BÁO CHỈ TIÊU KINH TẾ VĨ MÔ

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.

Mục tiêu chính của chương bao gồm:

- Đánh giá ưu điểm, hạn chế của từng nhóm mô hình: mô hình kinh tế lượng (hồi quy tuyến tính, arima, v.v.), học máy và học sâu trong dự báo GDP và một số chỉ tiêu kinh tế vĩ mô khác. Từ đó đưa ra điều kiện áp dụng cho từng nhóm mô hình,
- Thực nghiệm so sánh hiệu suất của các mô hình trong dự báo GDP và một số CTKTVM khác,
- Định hướng cho việc đề xuất xây dựng mô hình dự báo GDP ở chương 4.

3.1. Mô hình kinh tế lượng

Các mô hình kinh tế lượng truyền thống như hồi quy tuyến tính, ARIMA [29], mô hình vector tự hồi quy [118], mô hình hiệu chỉnh sai số vector [82], mô hình cân bằng tổng thể động [119], v.v. là những công cụ nền tảng trong phân tích và dự báo kinh tế vĩ mô. Các mô hình này dựa trên các lý thuyết kinh tế được xác lập, cho phép mô tả mối quan hệ nhân quả giữa các biến kinh tế thông qua các cấu trúc phương trình có thể kiểm định và giải thích rõ ràng. Mặc dù có những hạn chế nhất định khi đối mặt với dữ liệu có tính phi tuyến hoặc không ổn định nhưng nhờ đặc tính minh bạch và khả năng lý giải, các mô hình kinh tế lượng vẫn là công cụ đáng tin cậy trong nhiều ứng dụng thực tiễn.

***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 một hoặc nhiều biến độc lập. Các hệ số hồi quy được ước lượng thông qua phương pháp bình phương tối thiểu, nhằm tối thiểu hóa tổng sai số bình

phương giữa giá trị thực tế và giá trị dự báo. Mô hình này thường được sử dụng để xác định và lượng hóa tác động của các yếu tố kinh tế đến chỉ tiêu mục tiêu.

Mô hình ARIMA [29]: 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 một chuỗi thời gian đơn biến. Phương pháp luận Box–Jenkins được sử dụng để nhận dạng, ước lượng và kiểm định mô hình. ARIMA phù hợp với các chuỗi có xu hướng và không cần biến giải thích bên ngoài.

Mô hình VAR [118]: mở rộng mô hình AR sang nhiều biến phụ thuộc, mô hình hóa động lực giữa các biến trong hệ thống. Mô hình không yêu cầu xác định quan hệ nhân quả trước, do đó cho phép phân tích mối liên kết động giữa các chỉ tiêu kinh tế mà không cần cấu trúc lý thuyết cụ thể. 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 [82] : thích hợp cho các chuỗi có quan hệ đồng liên kết. VECM là dạng mở rộng của VAR được sử dụng khi các chuỗi thời gian có mối quan hệ đồng liên kết dài hạn. Mô hình này vừa mô tả động học ngắn hạn, vừa đảm bảo rằng hệ thống tiến tới trạng thái cân bằng dài hạn thông qua cơ chế hiệu chỉnh sai số.

Mô hình DSGE [119]: là mô hình vĩ mô kết cấu được xây dựng dựa trên các giả định hành vi tối ưu của tác nhân kinh tế dưới tác động của các cú sốc ngẫu nhiên, và giải quyết bằng phương pháp xấp xỉ tuyến tính quanh trạng thái cân bằng.

***Ưu điểm của các mô hình kinh tế lượng**

Các phương pháp kinh tế lượng truyền thống được đánh giá cao nhờ nền tảng lý thuyết vững chắc, xuất phát từ các nguyên lý cốt lõi trong kinh tế học vĩ mô. Các mô hình này cho phép mô hình hóa mối quan hệ nhân quả giữa các biến kinh tế theo cách có cấu trúc và có cơ sở học thuật rõ ràng. Tiêu biểu như mô hình DSGE, với khả năng mô phỏng phản ứng động của nền kinh tế đối với các cú sốc chính sách tiền tệ hoặc tài khóa trong môi trường có kỳ vọng hợp lý và điều chỉnh chậm, đã được Smets và Wouters (2003) [119] ứng dụng thành công cho Khu vực đồng euro.

Bên cạnh nền tảng lý thuyết, một điểm mạnh khác của các mô hình truyền thống là khả năng diễn giải kinh tế rõ ràng. Các tham số trong mô hình thường mang ý nghĩa cụ thể, phản ánh mức độ ảnh hưởng của các biến giải thích đối với biến phụ thuộc. Điều này giúp các nhà nghiên cứu dễ dàng giải thích kết quả theo cách phù hợp với tư duy kinh tế học, đồng thời cung cấp cơ sở định lượng để hỗ trợ hoạch định và điều chỉnh chính sách công.

Không chỉ được công nhận trong nghiên cứu hàn lâm, các mô hình này còn được áp dụng rộng rãi trong thực tiễn điều hành chính sách. Nhiều ngân hàng trung ương lớn như Cục Dự trữ Liên bang Mỹ (FED), Ngân hàng Trung ương châu Âu (ECB) và Ngân hàng Trung ương Nhật Bản (BOJ) đã sử dụng các mô hình như SVAR, DSGE và BVAR trong công tác phân tích và dự báo kinh tế vĩ mô. Chẳng hạn, mô hình BVAR với các phân phối tiên nghiệm hợp lý được Litterman (1986) [93] phát triển đã giúp nâng cao hiệu quả dự báo tại các thể chế tài chính lớn.

Một lợi thế quan trọng nữa của các mô hình truyền thống là hệ thống công cụ phân tích định lượng đi kèm. Các kỹ thuật như kiểm định đồng liên kết (Johansen, 1991) [82], kiểm định tính dừng chuỗi thời gian, phân rã phương sai, và đặc biệt là hàm phản ứng xung động (impulse response functions) cho phép nghiên cứu mối quan hệ động và nguyên nhân giữa các biến kinh tế. Nhờ đó, độ tin cậy và chiều sâu của phân tích định lượng trong mô hình hóa kinh tế vĩ mô được nâng cao đáng kể.

Tổng thể, sự kết hợp giữa nền tảng lý thuyết, khả năng diễn giải rõ ràng, tính ứng dụng cao và hệ thống công cụ thống kê phong phú đã giúp các mô hình kinh tế lượng truyền thống duy trì được vị thế trung tâm trong nghiên cứu và thực tiễn kinh tế học suốt nhiều thập kỷ.

Dù sở hữu nhiều ưu điểm, các mô hình kinh tế lượng truyền thống vẫn gặp phải không ít hạn chế và thách thức, nhất là khi ứng dụng vào dự báo trong bối cảnh kinh tế biến động phức tạp.

***Hạn chế của các mô hình kinh tế lượng**

Một trong những hạn chế rõ nét nhất của các nghiên cứu dựa trên các mô hình kinh tế lượng truyền thống như ARIMA, VAR hay VECM là giả định tuyến tính cứng nhắc giữa các biến kinh tế. Các nghiên cứu này thường xây dựng trên cơ sở cho rằng mỗi quan hệ giữa các biến đầu vào và biến đầu ra là tuyến tính và cố định theo thời gian. Tuy nhiên, trong thực tế, các quan hệ kinh tế thường mang tính phi tuyến, phụ thuộc vào điều kiện thể chế, kỳ vọng, và có xu hướng thay đổi theo chu kỳ kinh tế hoặc do các cú sốc bên ngoài. Điều này khiến các mô hình tuyến tính truyền thống khó lòng phản ánh đầy đủ tính phức tạp và năng động của hệ thống kinh tế thực tế.

Bên cạnh đó, các mô hình này cũng gặp nhiều khó khăn trong bối cảnh dữ liệu lớn và có chiều cao. Khi số lượng biến đưa vào mô hình tăng nhanh, đặc biệt vượt quá số quan sát, khả năng ước lượng ổn định của mô hình bị suy giảm đáng kể. Hiện tượng quá khớp dễ xảy ra, làm giảm tính tổng quát hóa và độ tin cậy của dự báo. Điều này là một bất lợi đáng kể trong thời đại dữ liệu phong phú như hiện nay, nơi mà hàng trăm chỉ báo kinh tế có thể được thu thập theo thời gian thực từ nhiều nguồn khác nhau.

Không chỉ vậy, các mô hình kinh tế lượng truyền thống còn nhạy cảm với quá trình lựa chọn mô hình và thông số. Việc xác định đúng độ trễ, phân loại chính xác giữa biến nội sinh và ngoại sinh, hay lựa chọn phân phối tiên nghiệm trong BVAR đều có ảnh hưởng lớn đến kết quả phân tích và dự báo. Như đã được Sims và Litterman [118], [93] chỉ ra, một sai lệch nhỏ trong bước thiết lập có thể dẫn đến kết quả dự báo sai lệch nghiêm trọng. Tính phụ thuộc cao vào kỹ năng và kinh nghiệm của người nghiên cứu làm giảm tính tự động hóa và khả năng tái lập mô hình trong thực tiễn.

Cuối cùng, một điểm yếu đáng lưu ý của các mô hình truyền thống là hiệu suất dự báo thường kém cạnh tranh hơn so với các phương pháp hiện đại như học máy và học sâu trong bối cảnh dữ liệu kinh tế hiện đại. Các nghiên cứu gần đây đã cho thấy các mô hình học máy như random forest, gradient boosting hay mạng nơ-ron có khả năng nắm bắt tốt hơn tính phi tuyến và động thái phức tạp trong dữ liệu kinh tế, từ đó cải thiện đáng kể độ chính xác dự báo [18]. Điều này đặc biệt đúng trong các bối cảnh

có độ bất định cao hoặc giai đoạn chuyển đổi kinh tế, khi các mô hình tuyến tính truyền thống dễ bị lỗi hệ thống do không thích ứng được với sự thay đổi cấu trúc.

Tổng hợp lại, các mô hình kinh tế lượng truyền thống có những đóng góp quan trọng trong nghiên cứu kinh tế học thực nghiệm, nhưng các hạn chế về tính tuyến tính, khả năng mở rộng, độ nhạy mô hình và hiệu suất dự báo đang trở thành rào cản lớn trong việc ứng dụng chúng vào phân tích kinh tế hiện đại. Đây cũng là lý do ngày càng nhiều nghiên cứu hướng đến tích hợp các phương pháp truyền thống với kỹ thuật học máy hiện đại nhằm khắc phục các điểm yếu nói trên.

***Điều kiện áp dụng các mô hình kinh tế lượng**

Việc áp dụng các mô hình kinh tế lượng trong dự báo kinh tế vĩ mô đòi hỏi một số điều kiện nhất định để đảm bảo tính hiệu quả và độ tin cậy của kết quả dự báo. Cụ thể:

Thứ nhất, dữ liệu chuỗi thời gian sử dụng cần có độ dài lịch sử đầy đủ và độ phân giải ổn định, chẳng hạn như dữ liệu theo quý hoặc theo năm. Ngoài ra, chuỗi dữ liệu phải đảm bảo tính dừng, đây là một giả định quan trọng trong hầu hết các mô hình truyền thống. Trường hợp chuỗi không dừng, cần áp dụng các phương pháp biến đổi như sai phân hoặc kiểm định đồng liên kết để đưa chuỗi về dạng dừng.

Thứ hai, số lượng biến đầu vào nên được giới hạn ở mức hợp lý, thông thường dưới 20 biến, nhằm hạn chế tình trạng đa cộng tuyến và giảm thiểu rủi ro quá khớp trong quá trình ước lượng tham số mô hình.

Thứ ba, các mô hình truyền thống đặc biệt phù hợp với mục tiêu dự báo biến số thường xuyên được sử dụng trong phân tích chính sách. Việc lựa chọn mô hình trong các trường hợp này cần tính đến yêu cầu cao về khả năng diễn giải để hỗ trợ cho hoạt động hoạch định và điều hành vĩ mô.

Thứ tư, một giả định quan trọng là tính ổn định của cấu trúc mô hình theo thời gian, tức là mối quan hệ giữa các biến không thay đổi đáng kể trong giai đoạn phân tích. Các cú sốc mang tính cơ cấu hoặc thay đổi chính sách có thể làm suy giảm độ chính xác của mô hình nếu không được xử lý thích hợp.

Thứ năm, dữ liệu đầu vào phải mang tính định lượng và có cấu trúc rõ ràng. Các mô hình truyền thống không phù hợp để xử lý dữ liệu phi cấu trúc như văn bản,

hình ảnh, hoặc dữ liệu thời gian thực từ Internet, vốn đòi hỏi các phương pháp học máy tiên tiến hơn để khai thác hiệu quả.

3.2. Mô hình học máy

Mặc dù các mô hình truyền thống đã đóng vai trò quan trọng trong lĩnh vực dự báo kinh tế vĩ mô nhờ tính đơn giản, dễ diễn giải và nền tảng lý thuyết vững chắc, nhưng những hạn chế cố hữu như giả định tuyến tính, yêu cầu dữ liệu dừng, và thiếu khả năng mô hình hóa các quan hệ phi tuyến đã dần bộc lộ khi hệ thống kinh tế ngày càng trở nên phức tạp và biến động nhanh hơn [95], [121].

Trong bối cảnh đó, sự phát triển vượt bậc của công nghệ tính toán và khả năng thu thập dữ liệu quy mô lớn đã mở ra hướng tiếp cận mới dựa trên học máy. Khác với các mô hình truyền thống vốn được xây dựng dựa trên cấu trúc mô hình cố định và giả định kinh tế học chặt chẽ, các phương pháp ML mang tính dữ liệu dẫn dắt, cho phép mô hình học trực tiếp từ dữ liệu và phát hiện các quy luật ẩn mà không cần xác định trước cấu trúc hàm.

Các mô hình ML không chỉ có khả năng xử lý dữ liệu phi tuyến, không dừng và có tương tác phức tạp, mà còn thích ứng tốt với sự đa dạng của dữ liệu kinh tế hiện đại bao gồm cả dữ liệu truyền thống và phi truyền thống. Nhờ đó, các phương pháp này ngày càng được nghiên cứu và ứng dụng rộng rãi trong dự báo các chỉ tiêu kinh tế vĩ mô như GDP, CPI, tỷ lệ thất nghiệp, cung tiền và tăng trưởng đầu tư.

*Nguyên lý hoạt động

Các mô hình học máy được xây dựng dựa trên nguyên tắc học từ dữ liệu nhằm phát hiện các mối quan hệ tiềm ẩn và đưa ra dự báo mà không cần giả định trước cấu trúc hàm giữa biến đầu vào và biến đầu ra.

Random Forest (RF) [30]: là một phương pháp học máy theo hướng tập hợp (ensemble learning), trong đó nhiều cây quyết định [114] được huấn luyện song song trên các tập dữ liệu bootstrap với lựa chọn ngẫu nhiên biến đầu vào tại mỗi nút. Nguyên lý hoạt động của RF dựa trên cơ chế biểu quyết hoặc trung bình hóa kết quả của các cây đơn lẻ nhằm giảm phương sai và nâng cao độ ổn định của mô hình. Bản

chất tập hợp của Random Forest, kết hợp nhiều cây quyết định, là nguyên nhân dẫn đến hiệu suất vượt trội của nó trong việc xử lý tính phi tuyến tính và giảm thiểu overfitting.

Extreme Gradient Boosting (XGBoost) [40]: hoạt động theo nguyên lý tăng cường tuần tự, trong đó mỗi cây mới được xây dựng nhằm khắc phục phần sai số còn lại từ các cây trước đó. Dựa trên nguyên lý huấn luyện tuần tự các mô hình yếu (thường là các cây quyết định) nhằm giảm thiểu sai số tích lũy theo từng vòng lặp học. Mỗi vòng lặp kế tiếp sẽ học từ phần dư của vòng trước để cải thiện chất lượng dự đoán tổng thể.

Điểm nổi bật của XGBoost so với các kỹ thuật boosting cổ điển như AdaBoost hay Gradient Boosting truyền thống nằm ở khả năng tích hợp các thành phần tối ưu hóa hiện đại như regularization (L1/L2) để kiểm soát độ phức tạp của mô hình và tránh hiện tượng quá khớp. Ngoài ra, XGBoost áp dụng chiến lược tree pruning sớm nhằm cắt tỉa các nhánh cây không cần thiết ngay từ giai đoạn xây dựng mô hình, giúp tiết kiệm tài nguyên tính toán và tăng tính tổng quát hóa cho mô hình [24].

Bên cạnh đó, XGBoost còn được thiết kế với kiến trúc tính toán song song, cho phép xử lý dữ liệu theo khối và tối ưu hóa tốc độ huấn luyện. Cơ chế tối ưu hóa bộ nhớ hiệu quả, kết hợp với khả năng xử lý dữ liệu thiếu một cách tự động, khiến XGBoost trở thành một công cụ đặc biệt mạnh mẽ trong các bài toán có dữ liệu lớn, đa chiều, và chứa nhiễu [15].

Máy học vector tựa (Support vector regression - SVM) : là một kỹ thuật học máy có giám sát được phát triển bởi Cortes và Vapnik (1995) [43], dựa trên nguyên lý tối ưu biên độ và giảm thiểu rủi ro cấu trúc nhằm tìm siêu phẳng tối ưu trong không gian đặc trưng chiều cao để phân tách dữ liệu hoặc phù hợp với dữ liệu huấn luyện[62]. SVM thực hiện nguyên tắc giảm thiểu rủi ro cấu trúc để tối đa hóa khả năng tổng quát hóa. Hàm Kernel cho phép SVM xử lý các mối quan hệ phi tuyến tính bằng cách ánh xạ dữ liệu đầu vào ban đầu vào một không gian đặc trưng chiều cao hơn. Các hàm kernel phổ biến bao gồm Linear, Polynomial, Radial Basis Function (RBF) và Sigmoid. SVR là một kỹ thuật SVM để dự đoán các giá trị liên tục, tìm một siêu phẳng phù hợp nhất với các điểm dữ liệu và giảm thiểu lỗi dự đoán. Sức mạnh

của SVM trong việc xử lý phi tuyến tính bắt nguồn từ việc sử dụng các hàm kernel, vốn là nguyên nhân cho việc ánh xạ dữ liệu vào một không gian chiều cao hơn, nơi phân tách/hồi quy tuyến tính trở nên khả thi.

Perceptron đa tầng (Multilayer Perceptron - MLP): Gồm một lớp đầu vào, một lớp đầu ra và một hoặc nhiều lớp ẩn. Các kết nối giữa các nơ-ron luôn truyền thẳng về phía trước. Thường sử dụng các thuật ngữ trễ của chuỗi làm đầu vào và dự báo làm đầu ra [101].

Ngoài các mô hình nêu trên, còn có nhiều thuật toán học máy khác đang được khám phá và thử nghiệm trong các ứng dụng dự báo kinh tế, thể hiện tiềm năng đáng kể trong việc cải thiện độ chính xác và khả năng thích ứng với môi trường dữ liệu biến động.

***Ưu điểm của các mô hình học máy**

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 Random Forest 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 phân phối chuẩn của chuỗi bằng cách sử dụng hàm mất mát Gaussian [137].

***Hạn chế của các mô hình học máy**

Mặc dù các mô hình học máy truyền thống (như Random Forest, XGBoost, SVM, MLP) đã chứng minh tiềm năng trong việc xử lý dữ liệu phi tuyến [58] và có tính tương tác phức tạp, nhưng chúng vẫn tồn tại một số hạn chế nhất định khi so sánh với mô hình kinh tế lượng và học sâu trong dự báo các chỉ tiêu kinh tế vĩ mô.

Thứ nhất, so với mô hình kinh tế lượng, các phương pháp học máy thường thiếu khả năng giải thích rõ ràng về mặt kinh tế. Trong khi các mô hình như ARIMA,

VAR, VECM hay DSGE được xây dựng trên nền tảng lý thuyết kinh tế vững chắc và cho phép phân tích mối quan hệ nhân quả giữa các biến, thì học máy lại hoạt động chủ yếu theo hướng “hộp đen”, khó diễn giải về mặt cơ chế tác động hoặc giải thích chính sách.

Thứ hai, phần lớn các thuật toán học máy truyền thống không có cơ chế nội tại để khai thác động học chuỗi. Điều này khiến các mô hình học máy phải phụ thuộc nhiều vào kỹ thuật tạo đặc trưng thủ công, làm giảm khả năng thích ứng linh hoạt với dữ liệu kinh tế có cấu trúc biến đổi theo thời gian.

Ngoài ra, một số mô hình như SVM hay MLP không mở rộng tốt với dữ liệu cực lớn hoặc có tính không đồng nhất cao – đặc trưng phổ biến trong dữ liệu kinh tế vĩ mô hiện đại [126].

***Đánh giá hiệu suất mô hình trong các nghiên cứu cụ thể**

Nghiên cứu của K. Beg [25] cho thấy Random Forest vượt trội hơn các mô hình khác trong việc dự đoán biến động ngành, đạt R-squared cao nhất là 0.998 và MAE thấp nhất là 0.765. XGBoost và Extra Tree đạt điểm ROC-AUC cao nhất là 0.86 trong dự đoán các chương trình được IMF hỗ trợ, vượt trội so với các phương pháp kinh tế lượng truyền thống [24]. Trong một nghiên cứu về dự báo tăng trưởng GDP, Random Forest đạt R^2 là 0.9998, trong khi Gradient Boosting đạt R^2 là 0.9979, cả hai đều vượt trội so với hồi quy tuyến tính [53]. Hiệu suất mạnh mẽ của các mô hình dựa trên cây, đặc biệt là khả năng nắm bắt các phi tuyến tính và tính mạnh mẽ, cho thấy chúng có giá trị cao đối với các nhà thực hành, đặc biệt trong các tình huống mà các mối quan hệ kinh tế phức tạp và không được hiểu rõ bằng các mô hình tuyến tính truyền thống [58]. Tuy nhiên, chi phí tính toán của chúng đối với các tập dữ liệu rất lớn vẫn là một cân nhắc thực tế.

SVM cung cấp kỹ thuật dự báo chính xác và hiệu quả hơn cho dữ liệu tài chính so với ANN và ARIMA, với RMSE thấp hơn đáng kể [106]. Trong dự báo suy thoái, một mô hình SVM sử dụng kernel RBF đạt độ chính xác kiểm định chéo là 74.2% và độ chính xác ngoài mẫu là 66.7%, đạt độ chính xác 100% trong việc dự báo các trường

hợp "Suy thoái" thực tế. Tuy nhiên, nó tạo ra 6 cảnh báo sai [57]. SVR cho thấy kết quả đáng tin cậy trong một số kịch bản nhưng gặp khó khăn trong việc nắm bắt toàn bộ mẫu cơ bản của dữ liệu trong các trường hợp đột biến. Mặc dù SVM nhìn chung tốt với "kích thước mẫu nhỏ", nghiên cứu được cung cấp về ứng dụng của nó trong dự báo kinh tế vĩ mô sử dụng một tập dữ liệu tương đối lớn (1967-2011 cho lãi suất và GDP của Hoa Kỳ) [72]. Điều này cho thấy rằng mặc dù về mặt lý thuyết có khả năng, lợi thế thực tế được chứng minh của nó đối với các tập dữ liệu kinh tế vĩ mô rất nhỏ có thể bị hạn chế trong các tài liệu hiện có.

Các nghiên cứu thực nghiệm cho thấy ANN/MLP thường đạt độ chính xác dự báo cao hơn các mô hình truyền thống trong các bài toán với dữ liệu phi tuyến hoặc có tính biến động mạnh [41], [67], [101], [107]. MLP thể hiện khả năng học tốt các mẫu phức tạp trong chuỗi thời gian kinh tế như GDP, lạm phát và chỉ số tài chính. Tuy nhiên, hiệu quả của mô hình phụ thuộc lớn vào cấu trúc mạng và kỹ thuật huấn luyện.

Để có cái nhìn rõ hơn về vai trò, đặc điểm và hiệu quả của các mô hình học máy trong dự báo kinh tế vĩ mô, luận án tiến hành tổng hợp và phân tích một số nghiên cứu thực nghiệm tiêu biểu đã công bố trong lĩnh vực này. Bảng 3.1 trình bày tổng quan về nguyên lý hoạt động, ưu – nhược điểm kỹ thuật, hiệu quả dự báo, chỉ tiêu kinh tế đã áp dụng và thời gian dự báo tương ứng của từng mô hình. Sự tổng hợp này cung cấp cơ sở thực tiễn quan trọng để đánh giá mức độ phù hợp của từng phương pháp khi áp dụng cho các bài toán dự báo kinh tế cụ thể.

Bảng 3. 1 So sánh các mô hình học máy trong dự báo kinh tế vĩ mô

Loại mô hình ML	Nguyên lý hoạt động chính	Ưu điểm chính trong dự báo KTVM	Nhược điểm chính trong dự báo KTVM	Hiệu suất điển hình (Chỉ số)	CTKTV Mthường được dự báo	Khoảng thời gian dự báo
Random Forest	Tập hợp cây quyết định, lấy mẫu bootstrap, chọn đặc trưng ngẫu nhiên cho mỗi nút, trung bình hóa dự đoán [30].	Xử lý phi tuyến tính, giảm overfitting, độ chính xác cao, tính diễn giải tốt về tầm quan trọng của đặc trưng [25].	Tính toán chuyên sâu, kém chính xác với sự kiện kinh tế hoàn toàn mới.[53]	R^2 : 0.998 (biến động ngành), 0.9998 (tăng trưởng GDP). [31]	GDP, biến động ngành, tín hiệu giao dịch vĩ mô. [25]	Ngắn đến trung hạn. [53]
XGBoost	Tăng cường gradient, huấn luyện tuần tự các mô hình yếu, sửa lỗi. [24]	Độ chính xác cao, tốc độ nhanh, xử lý dữ liệu lớn, mạnh mẽ với dữ liệu thiếu, cung cấp tầm quan trọng đặc trưng. [24]	Cần tinh chỉnh siêu tham số tỉ mỉ, hiệu suất có thể giảm với dữ liệu hoàn toàn mới. [15]	ROC-AUC: 0.86 (dự đoán chương trình IMF). R^2 : 0.9979 (tăng trưởng GDP). [24], [53]	GDP, chương trình IMF, biến động giá. [24]	Ngắn đến trung hạn. [24]
SVM/SVR	Tìm siêu phẳng tối ưu, sử dụng hàm kernel cho phi tuyến tính. [62]	Hiệu quả với mẫu nhỏ, xử lý phi tuyến tính và nhiễu, khả năng tổng quát hóa tốt, không có cực tiểu cục bộ.[62]	Nhạy cảm với tham số, khó chọn tham số tối ưu, tốn thời gian với kernel. [40]	RMSE: 0.00703 (dữ liệu tài chính) [106]. Độ chính xác ngoài mẫu: 66.7% (dự báo suy thoái). [57]	Chỉ số tài chính, dự báo suy thoái.	Ngắn đến trung hạn. [62]
ANN/MLP	Mạng nơ-ron truyền thẳng, nhiều lớp, xấp xỉ hàm. [101]	Xấp xỉ hàm phi tuyến tính, bắt mẫu phức tạp.	Yêu cầu nhiều dữ liệu, vấn đề hộp đen, nhạy cảm với kiến trúc khi mẫu nhỏ. [42]	Hiệu suất đa dạng, thường vượt trội ARIMA cho dữ liệu phi tuyến. [107]	GDP, lạm phát, chỉ số tài chính. [67]	Ngắn đến trung hạn.[41]

Từ tổng hợp trong bảng 3.1 có thể nhận thấy rằng các mô hình học máy như Random Forest, XGBoost, SVM, v.v. đều thể hiện khả năng dự báo khá tốt trong nhiều ngữ cảnh khác nhau của kinh tế vĩ mô, đặc biệt là đối với các chỉ tiêu có tính phi tuyến, chuỗi dữ liệu ngắn hoặc có mức độ nhiễu cao. Tuy nhiên, mỗi mô hình đều có giới hạn riêng và mức độ phù hợp khác nhau tùy thuộc vào cấu trúc dữ liệu, độ dài chuỗi thời gian và mục tiêu ứng dụng (dự báo ngắn hạn hay dài hạn).

Điều kiện áp dụng các mô hình ML trong dự báo CTKTVM

Các mô hình học máy truyền thống, đặc biệt là những mô hình dựa trên cây quyết định như Random Forest và XGBoost, cũng như mô hình Hỗ trợ Vector (SVM), được xem là lựa chọn phù hợp trong các bối cảnh dữ liệu có cấu trúc phức tạp nhưng không quá lớn về quy mô. Cụ thể:

Thứ nhất, các mô hình này thích hợp với các tập dữ liệu có quy mô trung bình nghĩa là không quá nhỏ đến mức thiếu thông tin để huấn luyện, nhưng cũng không quá lớn đến mức gây quá tải tính toán hoặc yêu cầu hạ tầng phần cứng cao. Trong môi trường kinh tế vĩ mô thực tế, việc thu thập dữ liệu đầy đủ theo thời gian có thể gặp hạn chế, do đó mô hình học máy truyền thống có thể khai thác hiệu quả những tập dữ liệu sẵn có với độ dài chuỗi thời gian vừa phải.

Thứ hai, trong những trường hợp mà mối quan hệ giữa các biến đầu vào và biến mục tiêu có tính phi tuyến cao và tương tác phức tạp, các mô hình tuyến tính như hồi quy bội hoặc ARIMA không còn đáp ứng tốt. Các thuật toán như Random Forest và XGBoost sử dụng cơ chế phân tách dựa trên thông tin entropy hoặc độ lợi Gini, cho phép phát hiện những mẫu ẩn phi tuyến mà các mô hình tuyến tính bỏ sót.

Thứ ba, các mô hình này đặc biệt phù hợp với dữ liệu đa chiều, khi số lượng biến độc lập lớn hơn đáng kể so với độ dài chuỗi thời gian. Trong bối cảnh như vậy, tính năng chọn lọc đặc trưng nội tại của Random Forest và XGBoost đóng vai trò quan trọng, giúp giảm thiểu hiện tượng overfitting và cải thiện khả năng tổng quát hóa của mô hình.

Thứ tư, các mô hình học máy truyền thống thường yêu cầu thời gian huấn luyện và tối ưu ít hơn so với các mô hình học sâu, đồng thời dễ triển khai hơn trong các hệ thống thực tế. Mặc dù không đạt độ chính xác tuyệt đối cao nhất, nhưng chúng vẫn cung cấp hiệu suất dự báo đáng tin cậy trong nhiều trường hợp, đặc biệt là khi yêu cầu giải thích mô hình ở mức trung bình thay vì hoàn toàn "hộp đen".

Cuối cùng, một ưu điểm nổi bật khác là khả năng xếp hạng mức độ ảnh hưởng của các biến đầu vào thông qua chỉ số quan trọng biến. Điều này đặc biệt hữu ích trong các nghiên cứu ứng dụng kinh tế, nơi mà việc hiểu rõ tầm quan trọng của từng yếu tố kinh tế là điều kiện tiên quyết cho hoạch định chính sách. Random Forest và XGBoost đều cung cấp công cụ trực quan và hiệu quả cho mục tiêu này.

Tóm lại, trong những tình huống mà dữ liệu vừa đủ lớn, cấu trúc phi tuyến, yêu cầu hiệu suất cao nhưng không quá khắt khe về diễn giải và hạ tầng tính toán, thì các mô hình học máy truyền thống vẫn giữ vai trò quan trọng và là bước đệm hữu ích trước khi chuyển sang các kiến trúc học sâu phức tạp hơn.

3.3. Mô hình học sâu

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. Theo Hull (2021) [72] thì các mô hình học sâu không những vượt trội trong việc khai thác động học phi tuyến và mối quan hệ phức tạp giữa các biến kinh tế, mà còn đặc biệt phù hợp với bối cảnh dự báo kinh tế hiện đại với dữ liệu lớn và biến động nhanh.

***Nguyên lý hoạt động:**

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 (Recurrent Neural Networks - 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.

Bộ nhớ ngắn dài hạn (Long Short-Term Memory - LSTM) và GRU (Gated Recurrent Unit) [69]: 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.

Các cơ chế cổng tinh vi trong LSTMs và GRUs là nguyên nhân cho khả năng của chúng trong việc giảm thiểu vấn đề gradient biến mất, vốn là một hạn chế quan trọng đối với các RNN đơn giản hơn trong việc xử lý các phụ thuộc dài hạn. Điều này trực tiếp dẫn đến hiệu suất được cải thiện cho các chuỗi dài hơn. GRU có thiết kế đơn giản hơn với hai cổng (cổng đặt lại và cổng cập nhật), giúp đào tạo nhanh hơn và hội tụ nhanh hơn.

Mạng tích chập (Temporal Convolutional Network - TCN): Sử dụng các phép toán tích chập để nắm bắt cả phụ thuộc ngắn hạn và dài hạn [138].

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 [42]:**

Mạng nơ-ron có thể yêu cầu nhiều sức mạnh tính toán và mất nhiều thời gian để hội tụ nếu chúng rất phức tạp. LSTM có nhiều tham số, đòi hỏi sức mạnh tính toán đáng kể và làm chậm quá trình xử lý dữ liệu. RNN đơn giản dễ bị vấn đề gradient biến mất và bùng nổ. Mặc dù LSTM và GRU giải quyết điều này, vấn đề gradient biến mất tiềm ẩn vẫn tồn tại, làm phức tạp việc đào tạo các mô hình sâu trên các chuỗi dài. Nhiều mô hình ANN, đặc biệt là các thuật toán học sâu, hoạt động như "hộp đen", khiến quá trình ra quyết định của chúng phức tạp và khó hiểu. Việc sử dụng ANN trong các nghiên cứu kinh tế vĩ mô tương đối hạn chế do kích thước mẫu nhỏ và tính chất tần suất thấp của dữ liệu kinh tế vĩ mô. Hiệu suất của mạng sẽ nhạy cảm hơn nhiều với kiến trúc/thiết kế khi dữ liệu huấn luyện khan hiếm. Có một sự đánh đổi rõ ràng giữa độ phức tạp/hiệu suất của mô hình và chi phí tính toán/tính diễn giải. Mặc dù các mô hình Transformer mang lại độ chính xác vượt trội và khả năng song song hóa, chúng đi kèm với nhu cầu tính toán cao hơn và thách thức về tính diễn giải.

***Đánh giá hiệu suất của các mô hình học sâu trong các nghiên cứu cụ thể:**

Hầu hết các thuật toán học sâu (GRU, LSTM, TCN, Transformer) đều vượt trội so với các mô hình khác trong dự báo chỉ số điều kiện tài chính vĩ mô [91]. GRU thường cho kết quả tốt nhất, đặc biệt là trong dự báo ngoài mẫu đơn biến cho tỷ giá hối đoái và dự báo ngoài mẫu đa biến cho các chỉ số thị trường chứng khoán. GRU cũng có thời gian huấn luyện ngắn hơn và tốc độ hội tụ nhanh hơn LSTM. Transformer đạt R-squared cao nhất (0.7688) và các giá trị lỗi thấp nhất (MSE, RMSE, MSLE, MAPE) trong dự báo GDP bình quân đầu người khi bao gồm CPI và

tỷ lệ thất nghiệp. LSTM thể hiện sự vượt trội về độ chính xác dự đoán, MSE, RMSE và R^2 trong dự báo các biến kinh tế [115]. Hiệu suất vượt trội của các mô hình học sâu so với các mô hình chuỗi thời gian truyền thống cho thấy rằng các "mối quan hệ phi tuyến tính phức tạp" vốn có trong dữ liệu kinh tế thực sự đáng kể và không được nắm bắt đầy đủ bởi các phương pháp cũ hơn. Cấu trúc đa lớp của học sâu đặc biệt phù hợp để khám phá những mối quan hệ này.

Bảng 3.2 dưới đây tổng hợp một số mô hình học sâu điển hình đã được sử dụng trong các nghiên cứu dự báo các chỉ tiêu kinh tế vĩ mô.

Bảng 3. 2 So sánh các mô hình học sâu trong dự báo kinh tế vĩ mô

Loại Mô hình ML	Nguyên lý Hoạt động Chính	Ưu điểm chính trong dự báo KTVM	Nhược điểm chính trong dự báo KTVM	Hiệu suất điển hình (Chỉ số)	Chỉ tiêu kinh tế vĩ mô thường được dự báo	Khoảng thời gian dự báo
ANN/MLP	Mạng nơ-ron truyền thẳng, nhiều lớp, xấp xỉ hàm. [101]	Xấp xỉ hàm phi tuyến, bắt mẫu phức tạp.	Yêu cầu nhiều dữ liệu, vấn đề hộp đen, nhạy cảm với kiến trúc khi mẫu nhỏ. [42]	Hiệu suất đa dạng, thường vượt trội ARIMA cho dữ liệu phi tuyến. [107]	GDP, lạm phát, chỉ số tài chính. [67]	Ngắn đến trung hạn.[41]
RNN/LSTM/GRU	Xử lý dữ liệu tuần tự, có bộ nhớ (qua các công cho LSTM/GRU). [67]	Bất phụ thuộc thời gian dài, giải quyết vấn đề gradient biến mất (LSTM/GRU), độ chính xác cao [67]. GRU nhanh hơn LSTM. [19]	Chi phí tính toán cao (LSTM), vấn đề gradient (RNN đơn giản), hộp đen.	GRU: MAPE 1.69% (giá Bitcoin). R^2 : 0.7078 (GDP bình quân đầu người + UR). [67], [17]	GDP, chỉ số tài chính, tỷ giá hối đoái, lạm phát.[67]	Ngắn đến dài hạn. [58]

Loại Mô hình ML	Nguyên lý Hoạt động Chính	Ưu điểm chính trong dự báo KTVM	Nhược điểm chính trong dự báo KTVM	Hiệu suất điển hình (Chỉ số)	Chỉ tiêu kinh tế vĩ mô thường được dự báo	Khoảng thời gian dự báo
TCN	Mạng tích chập thời gian, bất phụ thuộc ngắn và dài hạn. [138]	Hiệu suất dự báo xuất sắc, đặc biệt trong các kịch bản đầu vào đơn và đa. [91]	Tính toán phức tạp. [138]	Vượt trội ARIMA/LS TM (RMSE 0.26, MAPE 5.3% cho biến động giá nông sản). [138]	Chỉ số điều kiện tài chính vĩ mô, biến động giá. [91]	Ngắn đến trung hạn. [138]
Transformer	Cơ chế tự chú ý, xử lý song song, bất phụ thuộc tầm xa. [110]	Độ chính xác vượt trội, xử lý phi tuyến tính, song song hóa, mở rộng tốt với dữ liệu lớn. [38]	Chi phí tính toán cao, tính diễn giải thấp. [110]	R^2 : 0.7688 (GDP bình quân đầu người + CPI + UR). [38]	GDP, chỉ số tài chính. [38]	Ngắn đến dài hạn. [110]

Từ bảng 3.2 có thể nhận thấy rằng các mô hình học sâu, đặc biệt là LSTM và các biến thể như Bi-LSTM và GRU, thể hiện hiệu suất vượt trội trong các bài toán dự báo chuỗi thời gian kinh tế nhờ khả năng ghi nhớ dài hạn, xử lý dữ liệu phi tuyến, và thích nghi với cấu trúc động của chuỗi.

***Điều kiện áp dụng học sâu trong dự báo CTKTVM**

Các mô hình học sâu, đặc biệt là LSTM, thường được ưu tiên áp dụng trong các trường hợp sau:

Thứ nhất, chuỗi thời gian chứa các quan hệ phụ thuộc dài hạn và cấu trúc động phức tạp, chẳng hạn như dữ liệu GDP, CPI, tỷ lệ thất nghiệp, v.v. trong giai đoạn dài.

Thứ hai, mối quan hệ giữa các biến độc lập và biến mục tiêu có tính phi tuyến cao, đa chiều, khó đặc tả bằng mô hình tuyến tính hoặc mô hình cây quyết định truyền thống.

Thứ ba, dữ liệu có tính chu kỳ rõ rệt, chịu ảnh hưởng mạnh từ các yếu tố ngoại sinh, đòi hỏi mô hình có khả năng ghi nhớ, thích ứng và cập nhật linh hoạt theo thời gian.

Thứ bốn, khi khối lượng dữ liệu khai thác từ nhiều nguồn khác nhau, các mô hình học sâu có lợi thế nhờ khả năng học biểu diễn và tích hợp dữ liệu hiệu quả.

Tuy nhiên, điều kiện tiên quyết để áp dụng hiệu quả các mô hình học sâu là hạ tầng tính toán mạnh (GPU, TPU), và quy trình tối ưu hóa siêu tham số chặt chẽ. Bên cạnh đó, việc giải thích kết quả của các mô hình học sâu vẫn là một thách thức, đặc biệt trong lĩnh vực kinh tế học nơi minh bạch và giải thích mô hình là yêu cầu quan trọng để hỗ trợ hoạch định chính sách.

Như vậy, xét một cách tổng thể, các mô hình học máy và học sâu không nhằm thay thế hoàn toàn các phương pháp kinh tế lượng truyền thống, nhưng là một công cụ bổ sung mạnh mẽ và cần thiết trong bối cảnh kinh tế hiện đại. Việc lựa chọn mô hình phù hợp tùy thuộc vào nhiều yếu tố như chỉ tiêu dự báo, mục tiêu dự báo dài, trung hay ngắn hạn, chất lượng dữ liệu ,v.v

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%, kiểm tra: 20%.

3.4.2. 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),
- Random Forest: 50 cây, độ sâu 3,
- XGBoost: 50 cây, max_depth=3, learning_rate=0.05,
- SVR chọn $C=1$, $\epsilon=0.1$ với kernel RBF,
- LSTM:

Tham số	Giá trị
Số đơn vị ẩn (Số nơ-ron)	10
Hàm kích hoạt (activation)	tanh
Số vòng lặp (epochs)	200
Số lô (batch_size)	32
Bộ tối ưu (optimizer)	Adam
Hàm mất mát (loss)	mse

3.4.3. Chỉ số đánh giá hiệu suất của mô hình

Để đánh giá sai số của dự báo thực nghiệm sử dụng chỉ số MAE có công thức (1.8), RMSE có công thức (1.9) và chỉ số R^2 có công thức (1.11) trong mục 1.4 “Chỉ số đánh giá hiệu suất của mô hình”.

3.4.4. Kết quả và so sánh

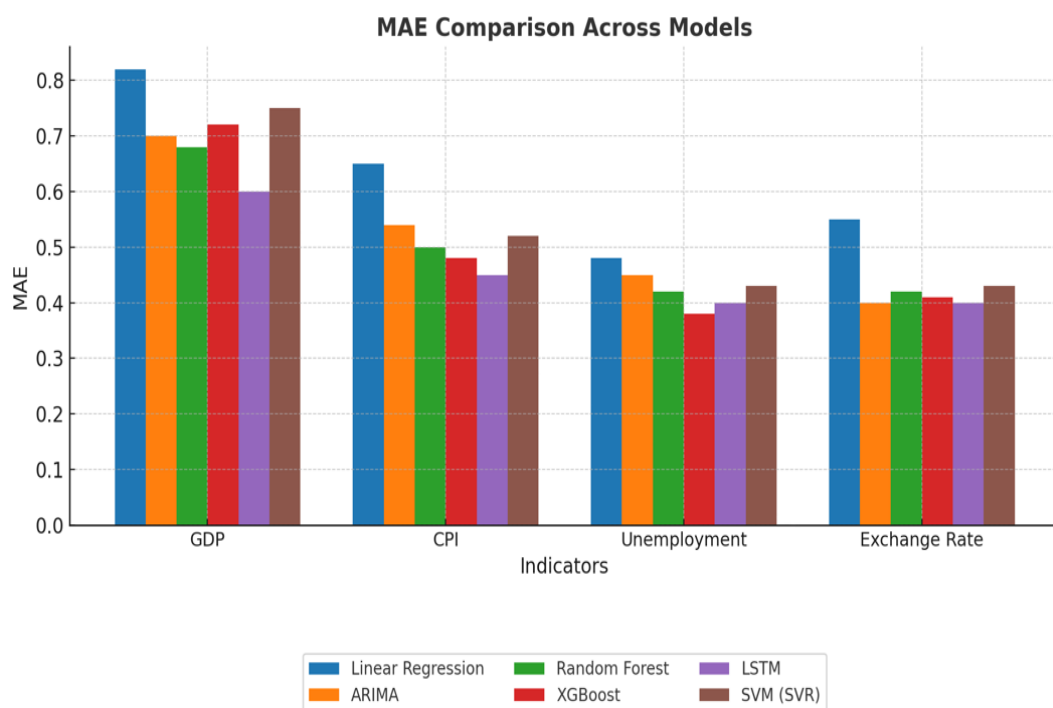
Trong thực nghiệm này, sáu mô hình gồm các mô hình kinh tế lượng và các mô hình học máy/học sâu đã được đánh giá đồng thời trên bốn chỉ tiêu kinh tế vĩ mô: GDP, tỉ lệ lạm phát, tỉ lệ thất nghiệp và tỉ giá hối đoái. Để phân tích rõ mức độ khác biệt giữa các nhóm phương pháp, chúng tôi tổng hợp các kết quả MAE, RMSE và R^2 trong Bảng 3.3

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

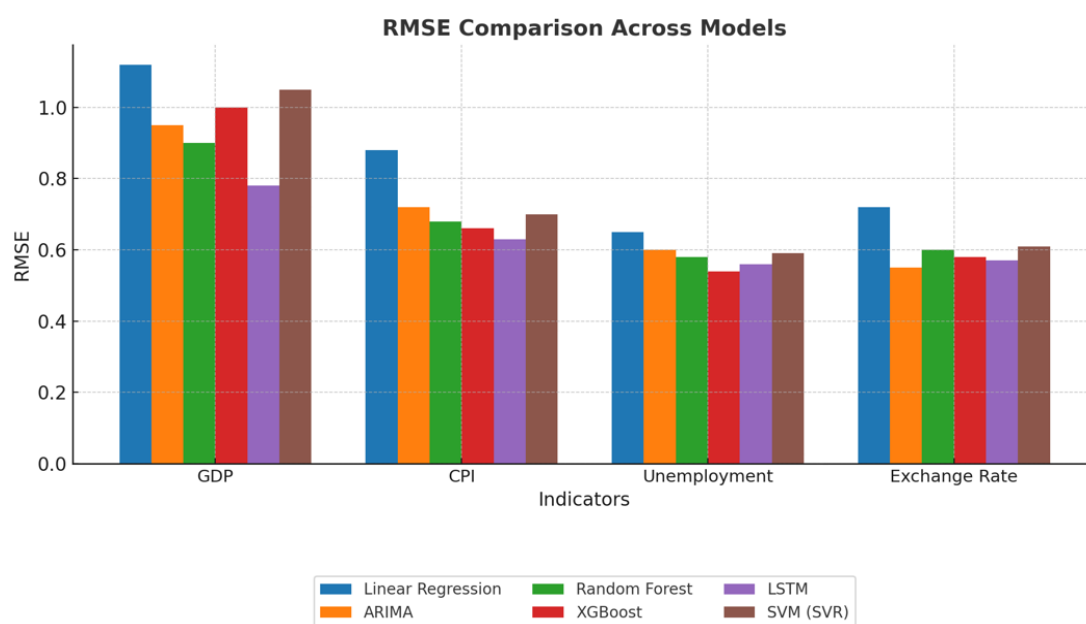
Chỉ tiêu	Mô hình	MAE	RMSE	R^2
GDP	Linear Regression	0.82	1.12	0.89
	ARIMA	0.70	0.95	0.92
	Random Forest	0.68	0.90	0.94
	XGBoost	0.72	1.00	0.93
	SVM	1.75	1.05	0.91
	LSTM	0.60	0.77	0.96
Tỷ lệ lạm phát	Linear Regression	0.65	0.88	0.91
	ARIMA	0.54	0.72	0.94
	Random Forest	0.50	0.68	0.95
	XGBoost	0.48	0.66	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
	SVM	0.43	0.59	0.91
	LSTM	0.40	0.56	0.92

Chỉ tiêu	Mô hình	MAE	RMSE	R ²
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

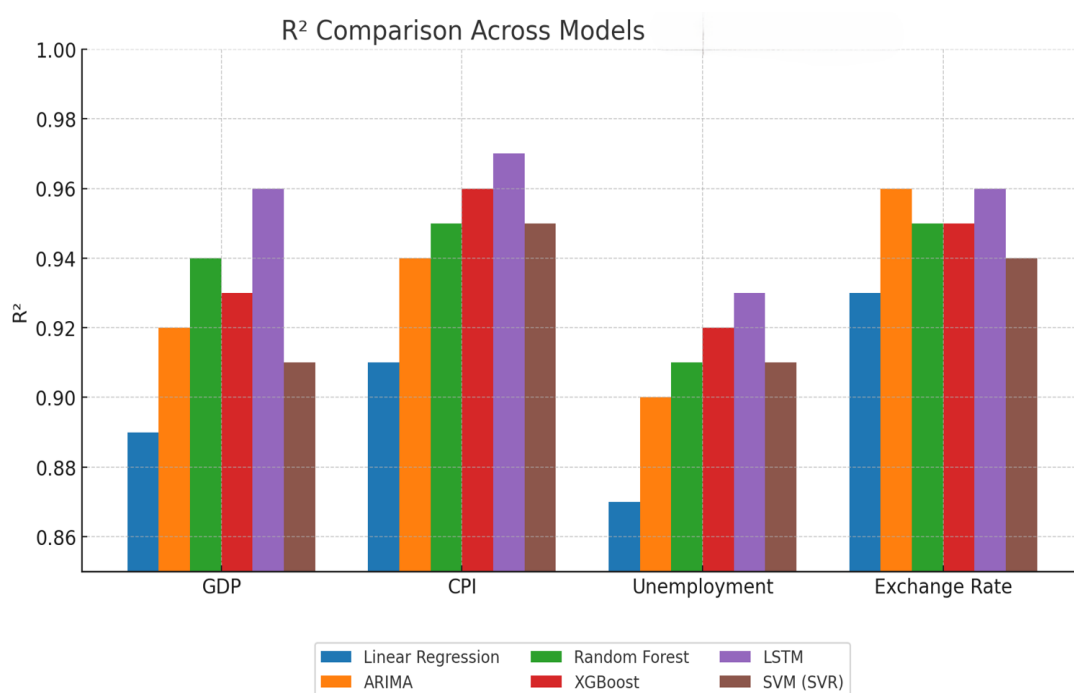
Kết quả bảng 3.7 được trực quan hóa qua hình 3.1 (MAE) , hình 3.2 (RMSE) và hình 3.3 (R²).



Hình 3. 1 MAE của các mô hình



Hình 3. 2 RMSE của các mô hình



Hình 3. 3 R^2 của các mô hình

Kết quả thực nghiệm trên bốn chỉ tiêu kinh tế vĩ mô của Việt Nam giai đoạn 1980–2019 (GDP, lạm phát – CPI, tỷ lệ thất nghiệp, và tỷ giá hối đoái) cho thấy sự

khác biệt rõ rệt về hiệu suất giữa các nhóm mô hình kinh tế lượng truyền thống và các mô hình học máy/học sâu.

Đối với GDP, mô hình LSTM đạt kết quả vượt trội với $RMSE = 0.77$, $MAE = 0.60$ và $R^2 = 0.96$, phản ánh khả năng học được các quan hệ phi tuyến và ghi nhớ phụ thuộc dài hạn đặc trưng của chuỗi thời gian GDP. So với ARIMA ($RMSE = 0.95$, $R^2 = 0.92$) và XGBoost ($RMSE = 1.00$, $R^2 = 0.93$), LSTM thể hiện mức giảm RMSE lần lượt khoảng 18,9% và 23,0%, qua đó khẳng định ưu thế của học sâu đối với chỉ tiêu có cấu trúc động thái phức tạp.

Với lạm phát (CPI), các mô hình học máy (Random Forest, XGBoost) và LSTM có R^2 tương đương 0.96–0.97, vượt trội so với Linear Regression (0.91) và ARIMA (0.94). Đặc biệt, XGBoost đạt $RMSE = 0.66$ và $MAE = 0.48$, cho thấy khả năng mô hình hóa hiệu quả biến động CPI vốn chịu ảnh hưởng của nhiều yếu tố phi tuyến và nhiễu ngẫu nhiên.

Ở tỷ lệ thất nghiệp, XGBoost và LSTM là hai mô hình có kết quả tốt nhất với RMSE lần lượt 0.54 và 0.54, R^2 đạt 0.92–0.93. So với các mô hình truyền thống ($RMSE > 0.58$, $R^2 \leq 0.91$), nhóm mô hình ML/DL giảm sai số khoảng 7–15%, cho thấy lợi thế trong việc xử lý dữ liệu nhiễu và chuỗi có biên độ dao động nhỏ.

Đối với tỷ giá hối đoái, ARIMA đạt kết quả $RMSE = 0.55$, $R^2 = 0.96$, chỉ nhỉnh hơn LSTM và XGBoost ở mức tối thiểu. Điều này gợi ý rằng với các chuỗi có cấu trúc tuyến tính rõ rệt và tính chu kỳ cao như tỷ giá, mô hình kinh tế lượng truyền thống vẫn có thể cạnh tranh với các kỹ thuật học máy/học sâu.

SVM, mặc dù được đánh giá cao trong các bài toán phi tuyến, chỉ đạt mức độ cạnh tranh trung bình (GDP $RMSE=1.05$) chủ yếu vì sự nhạy cảm với việc lựa chọn kernel và tham số C , ϵ .

Tổng hợp kết quả cho thấy, nhóm mô hình học máy và học sâu thể hiện ưu thế rõ rệt ở đa số chỉ tiêu, đặc biệt với các biến có tính phi tuyến mạnh như GDP, CPI, và tỷ lệ thất nghiệp. Tuy nhiên, đối với chỉ tiêu có tính tuyến tính cao như tỷ giá, sự khác biệt giữa hai nhóm mô hình không quá lớn. Kết quả này nhấn mạnh sự cần thiết

của việc lựa chọn mô hình phù hợp theo đặc tính dữ liệu, đồng thời cho thấy tiềm năng ứng dụng của học sâu trong dự báo kinh tế vĩ mô 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 trong chương này tập trung đánh giá hiệu suất của các nhóm mô hình dự báo thuộc ba nhóm: mô hình thống kê truyền thống (hồi quy, ARIMA), mô hình học máy truyền thống (SVR, Random Forest, XGBoost) và mô hình học sâu (LSTM), khi áp dụng vào bài toán dự báo các chỉ tiêu kinh tế vĩ mô then chốt như GDP, lạm phát (CPI), tỷ lệ thất nghiệp và tỷ giá hối đoái trên tập dữ liệu đã xây dựng ở chương 2. Những chỉ tiêu này đại diện cho ba nhóm biến trọng yếu: tăng trưởng, giá cả và thị trường lao động – thường xuyên được sử dụng để phản ánh trạng thái chu kỳ và mức độ ổn định kinh tế vĩ mô trong thực tiễn.

Thứ nhất, các mô hình học máy truyền thống như SVR, Random Forest, và XGBoost cho thấy hiệu quả dự báo khá tốt trong các trường hợp chuỗi thời gian có cấu trúc ổn định, mối quan hệ phi tuyến tồn tại rõ ràng giữa các biến đầu vào và đầu ra, và khi độ dài chuỗi không quá lớn. Đặc biệt, XGBoost – với cơ chế boosting tích hợp regularization – thể hiện khả năng học mạnh mẽ từ dữ liệu nhiễu, giảm thiểu overfitting và nâng cao độ chính xác trong nhiều bài toán dự báo kinh tế ngắn hạn hoặc trung hạn. Các mô hình dựa trên cây như Random Forest, mặc dù có thể thiếu độ nhạy đối với các biến thời gian, lại có ưu thế khi cần giải thích và xếp hạng vai trò các biến đầu vào, đặc biệt hữu ích trong dự báo tỷ lệ thất nghiệp, nơi có nhiều biến giải thích mang tính rời rạc và tương tác phi tuyến phức tạp.

Thứ hai, mô hình thống kê truyền thống như hồi quy tuyến tính và ARIMA vẫn giữ vai trò nền tảng trong dự báo kinh tế vĩ mô nhờ tính đơn giản, dễ diễn giải và được hỗ trợ bởi lý thuyết kinh tế lượng vững chắc. ARIMA phát huy hiệu quả trong các chuỗi dữ liệu ổn định, ngắn hạn và có đặc điểm tuyến tính rõ ràng. Tuy nhiên, nhược điểm cố hữu của các mô hình này là không thể nắm bắt quan hệ phi tuyến, kém thích ứng với dữ liệu nhiễu nhiễu động hoặc thay đổi theo cấu trúc thời gian. Ngoài

ra, trong bối cảnh hiện nay khi các cú sốc kinh tế diễn ra thường, các giả định tĩnh tại của ARIMA trở nên hạn chế.

Thứ ba, mô hình học sâu LSTM chứng minh khả năng vượt trội trong việc học các mẫu động phức tạp và ghi nhớ phụ thuộc thời gian dài – điều mà các mô hình khác khó đạt được. Với kiến trúc cổng (gating mechanisms), mô hình này cho phép điều chỉnh luồng thông tin đầu vào, từ đó học được chuỗi phụ thuộc dài hạn một cách hiệu quả, đặc biệt trong bài toán như dự báo GDP – nơi hành vi kinh tế phản ánh thông tin tích lũy theo thời gian. LSTM tỏ ra ưu việt trong điều kiện chuỗi có tính không ổn định hoặc chịu ảnh hưởng bởi các yếu tố vĩ mô không quan sát được trực tiếp.

Cuối cùng, thảo luận tổng thể cho thấy rằng không có mô hình nào là tối ưu tuyệt đối cho mọi loại chỉ tiêu, mục tiêu hoặc cấu trúc dữ liệu. Tuy nhiên, trong bối cảnh dữ liệu kinh tế vĩ mô ngày càng phức tạp, có tính phi tuyến, chứa nhiều nhiễu, và chịu ảnh hưởng bởi chu kỳ kinh tế, các mô hình học sâu nổi lên như một công cụ hữu hiệu – bổ sung và nâng cao hiệu suất dự báo bên cạnh các mô hình truyền thống, đặc biệt là ưu thế trong dự báo GDP quốc gia.

3.6. Kết luận chương

Chương 3 đã tiến hành đánh giá hiệu quả dự báo của các nhóm mô hình đại diện cho ba hướng tiếp cận chính: mô hình kinh tế lượng (hồi quy, ARIMA), mô hình học máy (RF, XGBoost), và mô hình học sâu (LSTM) trong bối cảnh dự báo các chỉ tiêu kinh tế 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. Tuy ML/DL không thay thế các mô hình kinh tế lượng truyền thống trong dự báo CTKTVM nhưng các mô hình hiện đại này đã chứng minh ưu thế rõ rệt trong nhiều tình huống thực tế. Đặc biệt, các mô hình học sâu như LSTM thể hiện

tính vượt trội trong việc học các quan hệ phi tuyến, ghi nhớ thông tin dài hạn trong chuỗi thời gian, và thích ứng tốt với những biến động dữ liệu phức tạp – vốn ngày càng phổ biến trong bối cảnh kinh tế hiện đại. Khả năng này giúp cải thiện đáng kể hiệu suất dự báo, đặc biệt đối với các chỉ tiêu có tính chu kỳ như GDP.

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 DỰ BÁO GDP QUỐC GIA

Tổng sản phẩm quốc nội là một trong những chỉ tiêu kinh tế vĩ mô quan trọng nhất, phản ánh quy mô và hiệu suất hoạt động của nền kinh tế quốc gia [97]. Tuy nhiên, việc dự báo GDP luôn gặp nhiều khó khăn do đặc điểm phi tuyến, chịu ảnh hưởng mạnh bởi các yếu tố bất định như khủng hoảng tài chính, đại dịch toàn cầu hoặc thay đổi chính sách kinh tế. Trong bối cảnh đó, thách thức đặt ra là làm thế nào để xây dựng được các mô hình dự báo có khả năng học được quan hệ biến động phức tạp giữa các yếu tố đầu vào và GDP.

Mục tiêu của chương này là đề xuất và kiểm định hiệu quả mô hình đề xuất (PAA-LSTM) trong dự báo GDP quốc gia. Cụ thể, nội dung của chương bao gồm: (i) Đề xuất kiến trúc của mô hình LSTM nhiều tầng kết hợp attention chú ý theo pha; (ii) phương pháp xác định pha chu kỳ kinh tế dựa trên thuật toán Bry-Boschan [33] và phương án mở rộng sang các pha phụ; (iii) xây dựng cơ chế ánh xạ trọng số attention riêng biệt cho từng pha; (iv) Thực nghiệm mô hình đề xuất dự báo GDP quốc gia của các nền kinh tế khác nhau; và (v) đánh giá hiệu suất mô hình, so sánh với các mô hình ML/DL và mô hình truyền thống nhằm kiểm chứng tính ưu việt và tính khái quát của mô hình đề xuất.

4.1. Giới thiệu bài toán

4.1.1. Đặt vấn đề

Các nhà hoạch định chính sách và các nhà kinh tế học dựa vào một loạt các chỉ tiêu kinh tế vĩ mô để định hướng các quyết định và chính sách xã hội tác động tới nền kinh tế thế giới [81]. Trong số các chỉ tiêu kinh tế vĩ mô, tổng sản phẩm quốc nội được sử dụng phổ biến nhất. GDP được coi là một chỉ báo thống kê quan trọng về sự phát triển và tiến bộ quốc gia [89]. GDP cũng có mối tương quan với tỷ lệ việc làm [16], [132], qua đó có thể cung cấp thông tin về các cơ hội thương mại và đầu tư [26]. Kỳ vọng là một yếu tố cốt lõi trong các lý thuyết kinh tế vĩ mô, đề cập đến những dự

báo và quan điểm của các chuyên gia về các chỉ tiêu kinh tế như doanh thu, thu nhập, thuế và giá cả trong tương lai [35], [45]. Các dự báo kinh tế vĩ mô cũng có thể tác động đến cách mà mỗi cá nhân kỳ vọng nền kinh tế sẽ diễn biến trong tương lai. Chính vì GDP có mối liên kết chặt chẽ với các chỉ tiêu kinh tế vĩ mô khác như giá cổ phiếu và tỷ lệ thất nghiệp [56], nên GDP được coi là chỉ tiêu "cốt lõi" để đo lường sự phát triển và tăng trưởng của một quốc gia. Việc dự báo chính xác GDP rất quan trọng đối với chính sách kinh tế nhằm giảm thiểu sai sót trong việc đưa ra quyết định [96]. Quỹ tiền tệ quốc tế công bố các báo cáo hàng quý về quan điểm của họ đối với sự phát triển kinh tế thế giới [78], trong đó bao gồm quan điểm về GDP có thể ảnh hưởng đến chính sách đối nội và đối ngoại của một quốc gia [52], [134].

Dự báo GDP thường đòi hỏi sử dụng các biến ngoại sinh (đặc trưng) như giá cổ phiếu, vốn có thể được thu thập với tần suất cao hơn so với dữ liệu GDP [102]. Việc dự báo GDP yêu cầu sử dụng các kỹ thuật mô hình hóa đặc thù của từng lĩnh vực, giúp hợp nhất dữ liệu trên các khoảng thời gian khác nhau, chẳng hạn như phương trình cầu nối (bridge equations) hoặc các mô hình nhân tố động (dynamic factor models) [50], [54], [98], [108]. Các mô hình này nói chung là những phương trình tuyến tính được xây dựng trên một chuỗi các biến ngoại sinh và biến tự hồi quy đã được lấy độ trễ hoặc lấy sai phân [103]. Chúng nằm trong nhóm các mô hình chuỗi thời gian truyền thống rộng hơn, được xác định bởi các phương pháp trung bình trượt tích hợp tự hồi quy (ARIMA) [27] và mô hình véc-tơ tự hồi quy (VAR) [118]. Điểm yếu của các phương pháp này là bị giới hạn trong các mối quan hệ tuyến tính, vốn là sự đánh đổi bắt buộc khi thiếu các bộ dữ liệu lớn. Tuy nhiên, những cách tiếp cận này liên tục không xác định được những thay đổi quan trọng trong nền kinh tế [48], [108]. Việc xác định phương pháp dự báo tốt nhất cho một mục tiêu kinh tế nhất định vẫn là một câu hỏi nghiên cứu mở. Các mô hình học máy có khả năng nắm bắt tốt hơn các mối quan hệ phi tuyến và thích ứng với dữ liệu nhiễu, vì vậy chúng đã thu hút được sự quan tâm của các nhà nghiên cứu dự báo kinh tế [18].

Bên cạnh thách thức của việc xây dựng mô hình dự báo GDP đến từ nguyên nhân là do bản chất đặc thù của dữ liệu KTVM, cụ thể, là các chuỗi dữ liệu này thường

có những đặc điểm phức tạp như: phi tuyến, chứa nhiều nhiễu, có độ trễ trong phản ứng với chính sách, thì GDP là chỉ tiêu đặc biệt chịu ảnh hưởng mạnh mẽ từ chu kỳ kinh tế [34]. Chu kỳ kinh tế thường diễn ra theo các pha nối tiếp nhau – bao gồm suy thoái, phục hồi, tăng trưởng và tăng trưởng mạnh – mỗi pha lại gắn với những động lực và cơ chế điều hành khác nhau [34]. Việc các chỉ tiêu kinh tế biến động không đồng đều theo từng pha chu kỳ mở ra hướng nghiên cứu về các mô hình dự báo có khả năng thích ứng linh hoạt theo thời kỳ kinh tế.

Ngoài ra, dự báo dài hạn cũng rất quan trọng đối với việc hoạch định chính sách, vì nhiều quốc gia cần lên kế hoạch kinh tế và đầu tư trước nhiều thập kỷ. Một ví dụ điển hình là Trung Quốc, đã bắt đầu lập kế hoạch và đầu tư vào châu Phi từ những năm 1990 [85]. Quỹ tiền tệ Quốc tế và Ngân hàng Thế giới đã có các dự báo kinh tế cho thập kỷ tới, được sử dụng như đầu vào quan trọng trong việc xây dựng chính sách, xem xét các yếu tố như biến đổi khí hậu, tăng trưởng hoặc suy giảm dân số và nguồn cung năng lượng [109], [134]. Tuy nhiên, các mô hình hiện tại thường mất ổn định và thiếu khả năng tổng quát hóa khi kéo dài tầm dự báo, đặc biệt trong môi trường kinh tế có tính chu kỳ và chịu nhiều tác động bất định.

Từ những phân tích trên, có thể khái quát vấn đề nghiên cứu đặt ra là làm thế nào để xây dựng một mô hình dự báo GDP hiệu quả:

- (1) Có khả năng học được các quan hệ phi tuyến và xử lý hiệu quả dữ liệu nhiễu và bất định,
- (2) Thích ứng linh hoạt với các pha của chu kỳ kinh tế để nâng cao hiệu quả dự báo trong điều kiện nền kinh tế toàn cầu biến động mạnh.

4.1.2. Hướng tiếp cận giải quyết vấn đề

Để giải quyết các vấn đề của bài toán dự báo chỉ tiêu kinh tế vĩ mô đặt ra trong mục 4.1.1, luận án hướng đến xây dựng một mô hình theo tiếp cận sau:

Thứ nhất, để giải quyết các vấn đề (1) xây dựng mô hình có khả năng học được các mối quan hệ phi tuyến, xử lý dữ liệu nhiễu và bất định bất định, các mô hình học máy thể hiện tính ưu việt hơn so với các mô hình kinh tế lượng [18]. Vì vậy, luận án

lựa chọn các phương pháp ML/DL là hướng tiếp cận chủ đạo trong dự báo chuỗi GDP quốc gia. Minh chứng rõ nét cho xu hướng này là các nghiên cứu quốc tế như Oancea (2025) [104], hay Lim và Zohren (2021) [92] cũng đã khẳng định hiệu quả vượt trội của các kiến trúc học sâu trong các bài toán dự báo kinh tế dựa trên chuỗi thời gian phi tuyến, độ trễ dài và dữ liệu có độ phức tạp cao. Ngoài ra công trình [CT3] và [CT4] của chúng tôi cũng cho thấy tiềm năng của các mô hình ML/DL trong dự báo CTKTVM. Trong đó, công trình [CT3] khẳng định giá trị của ML/DL so với mô hình ARIMA trong dự báo các chuỗi thời gian tần suất cao. Công trình [CT4] đã so sánh và đánh giá hiệu quả dự báo của các mô hình ML/DL và mô hình truyền thống để dự báo các chỉ tiêu kinh tế quan trọng như GDP, lạm phát, tỷ lệ thất nghiệp và tỷ giá hối đoái tại Việt Nam và các quốc gia Đông Nam Á, kết quả cho thấy độ chính xác của ML/DL được cải thiện đáng kể so với các mô hình kinh tế lượng truyền thống.

Thứ hai, mặc dù các mô hình học máy đã cho thấy hiệu quả trong việc xử lý dữ liệu phi tuyến [58], nhưng đặc thù của bài toán dự báo GDP quốc gia vốn chịu ảnh hưởng bởi các quan hệ động phức tạp và độ trễ kéo dài, đòi hỏi các kiến trúc có khả năng ghi nhớ dài hạn và học sâu theo thời gian. Kết quả tổng hợp tại Bảng 3.1 và Bảng 3.2 cho thấy các mô hình học sâu thường đạt hiệu suất vượt trội trong dự báo chuỗi thời gian mở rộng. Trên cơ sở đó, luận án lựa chọn học sâu làm nền tảng chính cho việc phát triển mô hình dự báo GDP quốc gia.

Trong số các mô hình học sâu, kiến trúc LSTM - một biến thể của mạng nơ-ron hồi tiếp được đề xuất bởi Hochreiter và Schmidhuber (1997) [69] – đã thể hiện rõ năng lực ghi nhớ dài hạn và mô hình hóa phi tuyến hiệu quả. Minh chứng cụ thể có thể kể đến nghiên cứu của B. Oancea (2025) [104], trong đó LSTM được triển khai để dự báo GDP hàng quý của Romania giai đoạn 1995–2023 và đạt hiệu quả vượt trội so với các mô hình kinh tế lượng tiêu chuẩn cả về độ chính xác và tính linh hoạt.

Song song với đó, mô hình Transformer, được giới thiệu bởi Vaswani và cộng sự (2017) [133], đã mở ra một bước đột phá trong lĩnh vực xử lý chuỗi dữ liệu, đặc biệt trong xử lý ngôn ngữ tự nhiên, và sau đó được mở rộng sang các bài toán dự báo

kinh tế. Nhờ cơ chế Attention và khả năng học biểu diễn không phụ thuộc vào thứ tự thời gian, Transformer có tiềm năng khắc phục nhiều hạn chế của LSTM, đặc biệt trong việc xử lý chuỗi dữ liệu dài và có cấu trúc phức tạp. Tuy nhiên, theo nghiên cứu của Li và cộng sự (2025) [91], việc ứng dụng rộng rãi Transformer trong dự báo kinh tế vĩ mô vẫn gặp nhiều trở ngại, do yêu cầu lượng dữ liệu đầu vào rất lớn và chi phí tính toán cao, trong khi dữ liệu kinh tế vĩ mô thường có giới hạn về độ dài chuỗi và tính liên tục.

Từ các phân tích nêu trên, luận án đề xuất hướng tiếp cận kết hợp ưu điểm của LSTM và Transformer thông qua việc xây dựng một mô hình LSTM tích hợp cơ chế Attention – thành phần cốt lõi trong kiến trúc Transformer. Cơ chế Attention cho phép mô hình gán trọng số linh hoạt cho từng bước thời gian, từ đó tập trung vào các điểm dữ liệu có ảnh hưởng lớn đến biến mục tiêu, góp phần cải thiện đáng kể khả năng học và dự báo.

Cuối cùng, theo lý thuyết kinh tế học cổ điển do Burns và Mitchell đề xuất năm 1946 [34], các chỉ tiêu kinh tế vĩ mô như GDP, lạm phát, hay tỷ lệ thất nghiệp vận động theo chu kỳ gồm bốn pha cơ bản: suy thoái, phục hồi, tăng trưởng và tăng trưởng mạnh. Việc bỏ qua đặc trưng chu kỳ trong mô hình hóa có thể làm giảm đáng kể khả năng thích ứng thời gian thực của hệ thống dự báo. Mặc dù lý thuyết chu kỳ đã được nghiên cứu sâu trong kinh tế học, nhưng theo khảo sát của luận án, việc tích hợp tri thức này vào mô hình học sâu vẫn chưa được thực hiện. Xuất phát từ khoảng trống đó, luận án nhận định rằng một vấn đề còn chưa được giải quyết là: làm thế nào để tối thiết kế cơ chế Attention sao cho có thể thích ứng với từng pha cụ thể của chu kỳ kinh tế. Việc điều chỉnh trọng số chú ý theo trạng thái chu kỳ – ví dụ như chú trọng đến tiêu dùng hộ gia đình trong giai đoạn phục hồi hay đến đầu tư công trong giai đoạn suy thoái – được kỳ vọng sẽ nâng cao độ chính xác, tính ổn định và khả năng tổng quát hóa của mô hình dự báo.

Trên cơ sở đó, 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

Xuất phát từ hướng tiếp cận giải quyết vấn đề đã được nhận diện trong mục 4.1.2, mục này đề xuất một kiến trúc học sâu mới PAA-LSTM. Mô hình được thiết kế nhằm khắc phục hạn chế của các phương pháp dự báo chuỗi thời gian kinh tế truyền thống và hiện đại vốn chưa khai thác hiệu quả thông tin về trạng thái chu kỳ kinh tế và giải quyết các vấn đề đặt ra ở mục 4.1.1.

Cụ thể, PAA-LSTM là sự kết hợp giữa mạng LSTM nhiều tầng – có khả năng học các phụ thuộc dài hạn và quan hệ phi tuyến trong chuỗi dữ liệu kinh tế vĩ mô – với một cơ chế attention thích ứng theo pha, được xây dựng có chủ đích để điều chỉnh trọng số học theo từng giai đoạn cụ thể của chu kỳ kinh tế. Thay vì xử lý dữ liệu đầu vào theo trình tự tuyến tính một cách đồng nhất, mô hình này tận dụng thông tin pha (bao gồm các trạng thái: suy thoái, phục hồi, tăng trưởng, và tăng trưởng mạnh) như một biến điều tiết trong quá trình phân bổ trọng số chú ý theo thời gian.

Cách tiếp cận này cho phép mô hình tập trung mạnh hơn vào những mốc thời gian có ý nghĩa kinh tế đặc biệt – chẳng hạn như thời điểm xảy ra cú sốc vĩ mô, điểm đảo chiều của chu kỳ, hay giai đoạn phản ứng chính sách. Nhờ đó, PAA-LSTM kỳ vọng đạt được hiệu suất dự báo cao hơn, đồng thời gia tăng tính thích ứng và khả năng khái quát trong các bối cảnh kinh tế có tính chu kỳ rõ rệt.

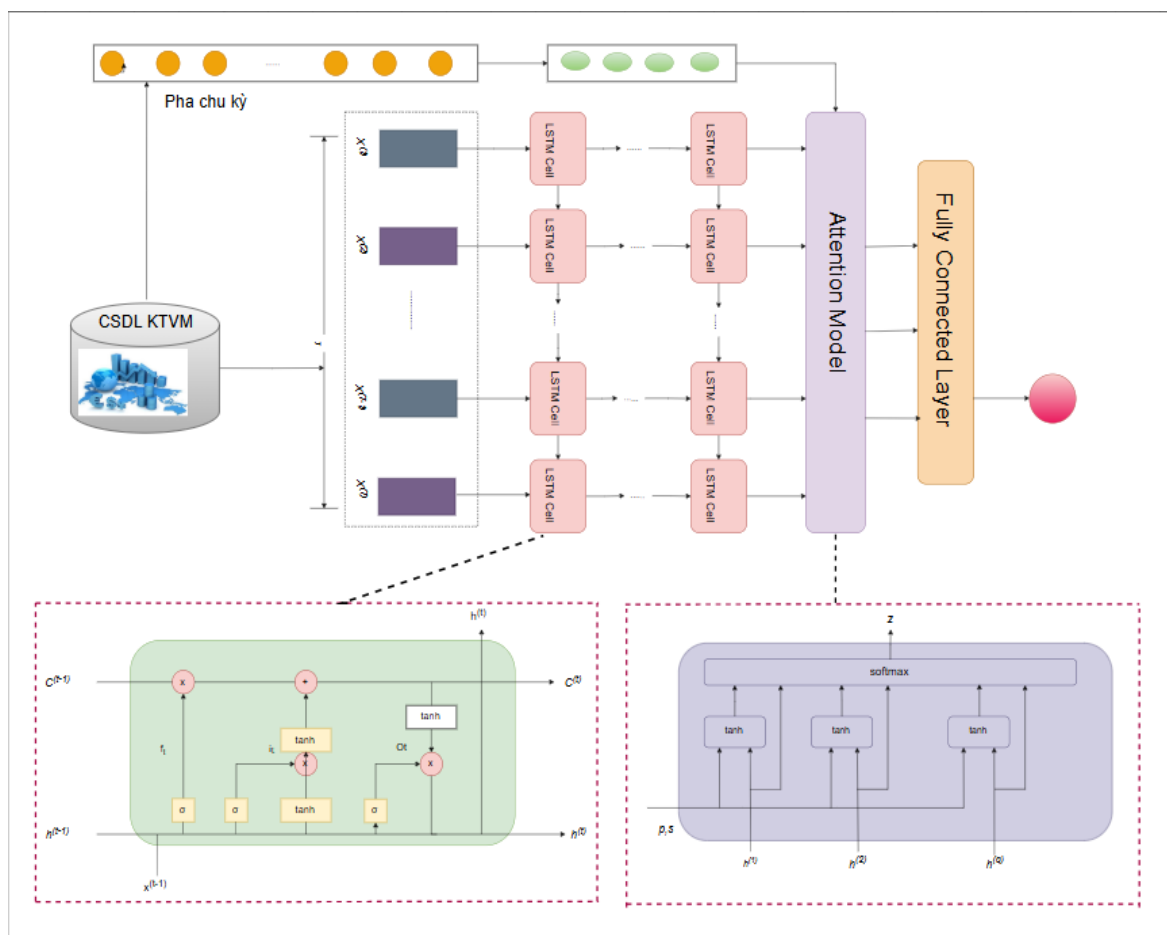
Kiến trúc tổng thể của mô hình PAA-LSTM bao gồm ba thành phần chính: mạng LSTM nhiều tầng, cơ chế attention thích ứng theo pha, và lớp đầu ra kết nối đầy đủ (fully-connected output layer) – minh họa trong Hình 4.1. Mục tiêu của kiến trúc này là tận dụng khả năng của LSTM trong việc mô hình hóa các phụ thuộc dài hạn và mối quan hệ phi tuyến trong chuỗi thời gian, đồng thời nâng cao tính linh hoạt theo thời gian và khả năng tập trung chọn lọc thông qua cơ chế attention thích ứng theo pha.

- LSTM nhiều tầng: Thành phần này chịu trách nhiệm trích xuất các đặc trưng động học sâu từ chuỗi thời gian đầu vào, bao gồm các biến như tăng trưởng GDP, tỷ lệ đầu tư, tiêu dùng hộ gia đình, và chỉ số việc làm. Việc sử dụng nhiều lớp LSTM cho phép mô hình học được các biểu diễn trừu tượng về động lực kinh tế vĩ mô;
- Cơ chế attention thích ứng theo pha: Khác với các cơ chế attention tiêu chuẩn, thành phần này được thiết kế riêng để thích ứng với từng pha của chu kỳ kinh tế (suy thoái, phục hồi, tăng trưởng, tăng trưởng mạnh). Mỗi pha tương ứng với một tập hợp tham số attention riêng biệt, cho phép mô hình ưu tiên các mốc thời gian quan trọng tùy theo ngữ cảnh kinh tế cụ thể;
- Lớp đầu ra kết nối đầy đủ: Lớp này kết hợp thông tin ngữ cảnh từ cả LSTM và attention để tạo ra kết quả dự báo cho bước thời gian tiếp theo.

Sự tích hợp của ba thành phần trên tạo nên một kiến trúc có khả năng nắm bắt đồng thời các quan hệ phụ thuộc dài hạn và mẫu biến động theo chu kỳ trong dữ liệu kinh tế vĩ mô, từ đó giúp mô hình trở nên thích nghi hơn và có độ chính xác cao hơn trong các môi trường kinh tế biến động mạnh.

Trong mô hình được đề xuất, cơ chế attention không còn được áp dụng một cách đồng đều trên toàn bộ chuỗi thời gian như trong các phương pháp truyền thống. Thay vào đó, attention được điều chỉnh một cách linh hoạt theo từng pha của chu kỳ kinh tế. Việc này được thực hiện bằng cách gán mỗi pha – bao gồm suy thoái, phục hồi, tăng trưởng và tăng trưởng mạnh – với một tập trọng số attention riêng biệt. Các trọng số này điều chỉnh mức độ tập trung của mô hình vào những thời điểm cụ thể, tùy thuộc vào đặc trưng của từng pha kinh tế.

Chẳng hạn, trong giai đoạn suy thoái, mô hình có thể ưu tiên trọng số cao hơn cho các bước thời gian liên quan đến chỉ số tiêu dùng và lao động. Ngược lại, trong giai đoạn phục hồi hoặc tăng trưởng, cơ chế attention sẽ học cách tập trung nhiều hơn vào các biến như đầu tư hoặc sản xuất công nghiệp.



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

Nhờ đó, cơ chế attention thích ứng theo pha giúp mô hình trở nên linh hoạt hơn, đồng thời vẫn giữ được khả năng chọn lọc thông tin như attention truyền thống, nhưng theo một cách phù hợp hơn với tính động và chu kỳ của hệ thống kinh tế vĩ mô.

4.3. Mô tả kiến trúc đề xuất

Trong kiến trúc mô hình được đề xuất, cơ chế attention không được áp dụng một cách tĩnh và đồng nhất trên toàn bộ chuỗi thời gian như trong các mô hình học sâu truyền thống. Thay vào đó, attention được thiết kế để thích ứng một cách linh hoạt

với từng pha cụ thể của chu kỳ kinh tế. Quá trình này bao gồm hai bước triển khai chính:

(1) *Phân đoạn chu kỳ kinh tế thành các pha có ý nghĩa kinh tế học rõ ràng,*

(2) *Xây dựng cơ chế ánh xạ trọng số attention riêng biệt cho từng pha.*

4.3.1. Phân đoạn chu kỳ kinh tế

Để xác định các pha của chu kỳ kinh tế như 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.

Về cơ bản, phương pháp BB tiến hành lọc nhiễu chuỗi GDP thực tế. Sau đó, chuỗi đã được làm mượt sẽ được phân tích để xác định các điểm cực trị cục bộ – đại diện cho đỉnh (điểm chuyển từ tăng sang giảm) và đáy (điểm chuyển từ giảm sang tăng) trong chu kỳ.

Các điểm cực trị được phát hiện thông qua quy tắc hình thức:

- Đỉnh (peak) tại thời điểm t được xác định khi:

$$y_{t-1} < y_t \text{ và } y_t > y_{t+1} \quad (4.1)$$

- Đáy (trough) tại thời điểm t được xác định khi:

$$y_{t-1} > y_t \text{ và } y_t < y_{t+1} \quad (4.2)$$

Tuy nhiên, không phải mọi điểm cực trị đều được chấp nhận là điểm chuyển pha.

Để đảm bảo rằng các đỉnh và đáy phản ánh các chuyển động chu kỳ thực sự, thuật toán Bry–Boschan bổ sung một số điều kiện ràng buộc:

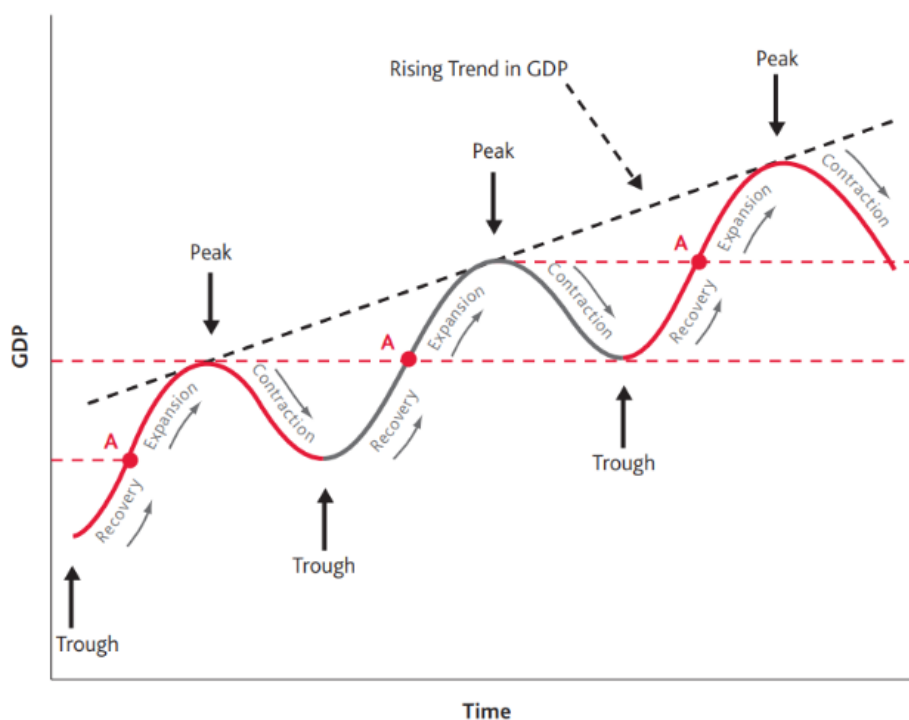
- Độ dài tối thiểu của pha: yêu cầu khoảng thời gian giữa một đỉnh và đáy (hoặc ngược lại) phải lớn hơn một số kỳ nhất định (thường là 2–3 kỳ đối với chuỗi quý, và 1–2 kỳ đối với chuỗi năm).

- Chu kỳ tối thiểu: độ dài từ một đỉnh đến đỉnh kế tiếp, hoặc từ đáy đến đáy kế tiếp, phải lớn hơn hoặc bằng một ngưỡng định trước (thường là 5–6 năm).

Các điều kiện trên nhằm loại bỏ các dao động ngắn hạn và tránh phân pha sai do nhiễu số liệu hoặc biến động ngẫu nhiên.

Sau khi xác định được dãy các đỉnh và đáy hợp lệ, thuật toán tiến hành đánh nhãn các giai đoạn giữa các điểm cực trị:

- Giai đoạn suy thoái: GDP từ đỉnh đến đáy;
- Giai đoạn phục hồi: GDP từ đáy về lại đỉnh trước đó;
- Giai đoạn tăng trưởng: từ khôi phục 100% đến khôi phục (thường là) 150%;
- Giai đoạn tăng trưởng mạnh: từ tăng trưởng đến đỉnh kế tiếp.



Hình 4. 2 Phân pha chu kỳ kinh tế¹

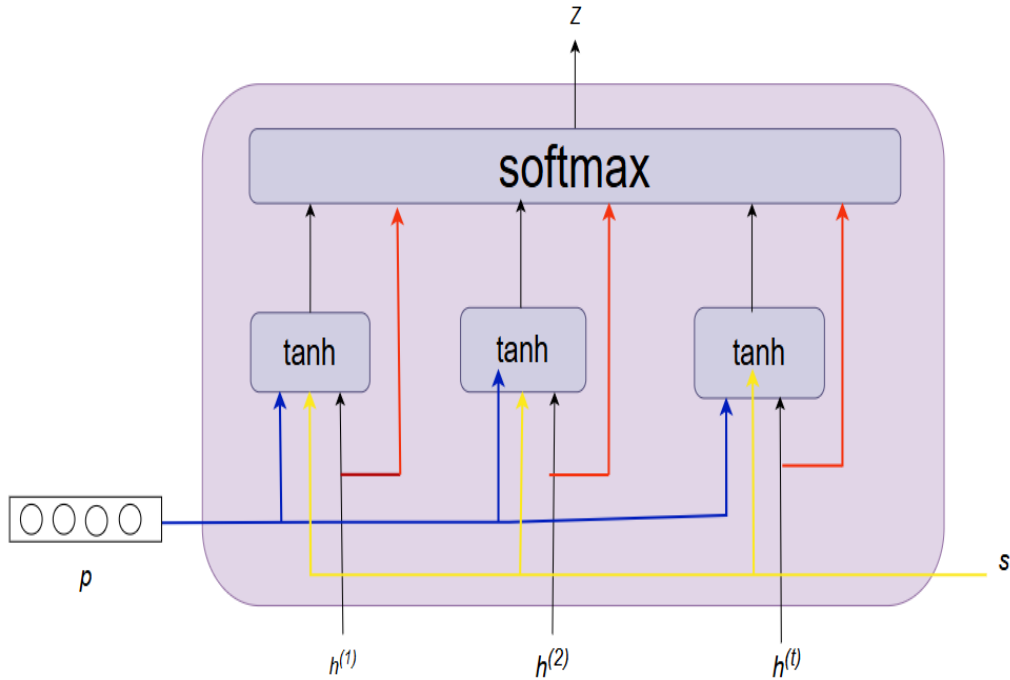
Những nhãn này sau đó được mã hóa lại thành vector one-hot/embedding để sử dụng làm đặc trưng đầu vào trong attention của PAA-LSTM.

¹ <https://www.linkedin.com/pulse/5-phases-business-cycle-nivethan-raj-/>

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

Trong các mô hình LSTM thông thường, cơ chế attention được sử dụng nhằm tập trung vào các thông tin quan trọng nhất trong chuỗi dữ liệu đầu vào khi dự báo tại thời điểm t . Tuy nhiên, cơ chế này thường cố định cho toàn bộ chuỗi, không phản ánh được sự thay đổi cấu trúc kinh tế theo thời gian.

Trong mô hình PAA-LSTM, cơ chế attention được thiết kế để thích ứng linh hoạt với từng pha của chu kỳ kinh tế nhằm cho phép mô hình học được sự thay đổi của trọng số chú ý giữa các biến kinh tế vĩ mô trong từng pha của chu kỳ kinh tế. Khác với cơ chế attention truyền thống chỉ dựa trên trạng thái ẩn của encoder và decoder, PAA-LSTM đưa thêm vector pha chu kỳ (p) vào công thức tính điểm attention nhằm điều chỉnh sự tập trung của mô hình theo đặc điểm của từng pha như hình 4.3.



Hình 4. 3 Kiến trúc Attention thích ứng theo pha (p)

Cụ thể, tại mỗi thời điểm t , điểm attention $e_t^{(p)}$ ứng với pha p được xác định như sau:

$$e_t^{(p)} = v_p^T \cdot \tanh(W_h h_t + W_s s_t + W_p p_t + b) \quad (4.3)$$

Trong đó:

- $h_t \in \mathbb{R}^d$: trạng thái ẩn của encoder LSTM tại thời điểm t ;
- $s_t \in \mathbb{R}^d$: trạng thái hiện tại của decoder (decoder state), phản ánh ngữ cảnh dự báo;
- $p \in \mathbb{R}^k$: vector biểu diễn pha chu kỳ;
- W_h, W_s, W_p : các ma trận trọng số chiều không gian của h_t, s_t , và p tương ứng về cùng một không gian attention;
- $v_p \in \mathbb{R}^{d_a}$: vector trọng số đặc thù của pha p , được huấn luyện riêng cho từng pha;
- $b \in \mathbb{R}^{d_a}$ da: vector bias
- $\tanh(\cdot)$: hàm kích hoạt phi tuyến giúp mô hình học được quan hệ phức tạp giữa các thành phần.

Sau khi tính được điểm attention $e_t^{(p)}$ tại từng thời điểm, mô hình sử dụng hàm softmax để chuẩn hóa và thu được trọng số attention:

$$\alpha_t^p = \text{softmax}(e_t^{(p)}) \quad (4.4)$$

Trong đó:

- α_t^p : là trọng số attention tại thời điểm t trong pha p ;
- $e_t^{(p)}$: điểm của attention ứng với pha p tại thời điểm t ;

Sau khi tính được các trọng số attention, chúng được sử dụng để tính trung bình có trọng số (weighted average) trên toàn bộ chuỗi các trạng thái ẩn của LSTM theo công thức:

$$c_t = \sum_t \alpha_t^p \cdot h_t \quad (4.5)$$

Vector ngữ cảnh c_t thu được sẽ được kết hợp với đầu ra hiện tại của LSTM, sau đó chuyển qua một lớp mạng dày (dense layer) để tạo ra dự báo tại bước thời gian tiếp theo.

Như vậy, cụ thể mỗi pha chu kỳ được xác định bằng thuật toán Bry–Boschan, từ đó gán nhãn pha $p_t \in \{1,2,3,4\}$ cho từng quan sát. Trong quá trình huấn luyện, cơ chế

attention tạo ra ma trận trọng số, thể hiện mức độ đóng góp của biến i tại thời điểm t trong pha p .

Giá trị trung bình pha được tính như:

$$w_i^{(p)} = \frac{1}{T_i} \sum \alpha_{t,i}^p \quad (4.6)$$

với T là số bước thời gian thuộc pha p .

Việc phân rã ma trận attention $\alpha_{t,i}^p$ thành các trọng số trung bình w_i^p cho phép trực quan hoá mức độ đóng góp của từng biến trong từng pha giúp giải thích vì sao mô hình đạt được hiệu suất dự báo cao hơn. Các trọng số attention theo pha cung cấp thông tin định lượng về biến số nào quan trọng nhất trong mỗi trạng thái kinh tế, hỗ trợ các nhà hoạch định chính sách xác định trọng tâm điều hành. Ví dụ: tập trung vào việc làm trong pha suy thoái, hay thúc đẩy đầu tư và năng suất trong pha phục hồi và tăng trưởng.

Cấu trúc tổng quát của cơ chế attention theo pha

Bước 1 – Phân pha chu kỳ: xác định pha p cho mỗi thời điểm bằng thuật toán Bry–Boschan.

Bước 2 – LSTM Encoder: trích xuất đặc trưng động học từ chuỗi đầu vào, thu được vector trạng thái h_t .

Bước 3 – Attention theo pha: lựa chọn bộ attention ứng với pha hiện tại để tính trọng số chú ý cho các biến.

Bước 4 – Tổng hợp pha: ghép các attention output của bốn pha thành ma trận $C^* = [C^{(1)}, C^{(2)}, C^{(3)}, C^{(4)}]$

Bước 5 – Dự báo GDP: đưa đầu ra qua lớp Fully Connected để ước lượng giá trị GDP hoặc tốc độ tăng trưởng dự báo.

Các bước chi tiết tính toán các trọng số attention đặc thù theo pha được trình bày trong hình 4.4.

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

Beggin

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

End

Hình 4. 4 Giả mã của thuật toán attention trong PAA-LSTM

Để hiểu rõ hơn về cách thức hoạt động của thuật toán PAA-LSTM đã trình bày trong Hình 4.4, phần tiếp theo sẽ minh họa một ví dụ cụ thể.

Giả sử $T = 3$ và đầu ra LSTM là các giá trị vô hướng như sau:

$t = [1, 2, 3]; h = [0.5, 1.0, -0.2]$

Pha suy thoái:

Let $w^{\text{recession}} = 0.8$, and $b^{\text{recession}} = 0.1$

Tính toán e_t :

$$e_1 = \tanh(0.8 \times 0.5 + 0.1) = \tanh(0.5) = 0.462$$

$$e_2 = \tanh(0.8 \times 1.0 + 0.1) = \tanh(0.9) = 0.716$$

$$e_3 = \tanh(0.8 \times (-0.2) + 0.1) = \tanh(-0.06) = -0.06$$

Trọng số attention (softmax):

$$\alpha_1 = \frac{e^{0.462}}{e^{0.462} + e^{0.716} + e^{-0.06}} \approx 0.30$$

$$\alpha_2 \approx 0.47$$

$$\alpha_3 \approx 0.23$$

Vector ngữ cảnh:

$$C^{\text{recession}} \approx 0.30 \times 0.5 + 0.47 \times 1.0 + 0.23 \times (-0.2) \approx 0.41$$

Pha tăng trưởng:

Let $w^{\text{expansion}} = 0.2$, and $b^{\text{expansion}} = 0$

Tính toán e_t :

$$e_1 = \tanh(0.2 \times 0.5) = \tanh(0.1) = 0.1$$

$$e_2 = \tanh(0.2 \times 1.0) = \tanh(0.2) = 0.197$$

$$e_3 = \tanh(0.2 \times (-0.2)) = \tanh(-0.04) = -0.04$$

Trọng số attention (softmax):

$$\alpha \approx [0.31, 0.34, 0.35]$$

Vector ngữ cảnh:

$$C^{\text{expansion}} \approx 0.31 \times 0.5 + 0.34 \times 1.0 + 0.35 \times (-0.2) \approx 0.42$$

Trong pha suy thoái, vector trọng số $w^{\text{recession}}$ có giá trị dương lớn, khiến mô hình phân bổ attention cao hơn cho các bước thời gian có giá trị trạng thái ẩn lớn hơn (ví dụ: tại $t=2$). Trong pha tăng trưởng, các điểm số attention được phân bổ đồng đều hơn giữa các bước thời gian. Điều này phản ánh sự ổn định trong cơ chế attention của mô hình trong các giai đoạn tăng trưởng, khi mà động thái GDP thường diễn ra một cách từ tốn và ít biến động đột ngột.

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

Phân tích độ phức tạp tính toán là cơ sở quan trọng nhằm đánh giá mức độ tài nguyên cần thiết cho quá trình huấn luyện và suy luận, đồng thời chứng minh khả năng mở rộng khi mô hình được áp dụng trên các bộ dữ liệu có quy mô lớn hơn. Việc xác định độ phức tạp cũng giúp so sánh hiệu quả giữa PAA-LSTM và các mô hình dự báo khác.

Mô hình PAA-LSTM bao gồm hai thành phần chính: khối LSTM chuẩn và cơ chế Phase-Adaptive Attention (PAA). Với số đặc trưng đầu vào d , số đơn vị ẩn h , độ dài chuỗi T , số lớp L và kích thước lô B , chi phí tính toán của một lớp LSTM biểu diễn như sau:

$$O(BTL(dh + h^2)) \quad (4.7)$$

trong khi số tham số cho một lớp LSTM xấp xỉ:

$$\# \theta_{\text{LSTM}} \approx 4(dh + h^2 + h) \quad (4.8)$$

Đối với cơ chế Phase-Adaptive Attention, thay vì phải tính toán attention toàn cục với độ phức tạp $O(T^2)$, mô hình chỉ tính attention trong từng pha của chu kỳ kinh tế. Nhờ đó, chi phí của PAA tăng tuyến tính theo độ dài chuỗi: $O(BTLTh)$ giúp mô hình duy trì hiệu năng tốt ngay cả khi dữ liệu có tần suất cao (quý, tháng) hoặc chiều dài chuỗi lớn.

Tổng hợp hai thành phần trên, độ phức tạp tính toán của PAA-LSTM biểu diễn như sau:

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

Độ phức tạp này tăng tuyến tính theo độ dài chuỗi và số mẫu, cho phép mô hình dễ dàng mở rộng quy mô dữ liệu. Khi mở rộng sang nhiều quốc gia hoặc sử dụng dữ liệu có tần suất cao hơn, PAA-LSTM vẫn duy trì chi phí huấn luyện và suy luận ở mức hợp lý.

Tóm lại, PAA-LSTM có độ phức tạp tính toán hợp lý, vừa đảm bảo hiệu quả mô hình hóa động lực chu kỳ kinh tế, vừa duy trì chi phí khả thi trong huấn luyện và suy luận. Đây là cơ sở quan trọng để khẳng định tính khả thi của việc ứng dụng mô

hình trong các kịch bản dự báo ở quy mô đơn vị và quốc gia với nguồn dữ liệu phong phú hơn.

4.5. Thực nghiệm

4.5.1. Dữ liệu và Phạm vi

Thực nghiệm sử dụng bộ dữ liệu kinh tế vĩ mô được xây dựng ở chương 2.

Phạm vi: Sáu quốc gia, đại diện cho ba nhóm nền kinh tế khác nhau:

- Các nền kinh tế mới nổi: Trung Quốc, Nga;
- Các nền kinh tế đang phát triển: Việt Nam, Ấn Độ;
- Các nền kinh tế phát triển: Hoa Kỳ, Canada.

Thời gian: Từ năm 1980 đến 2019; riêng đối với Nga, dữ liệu chỉ có từ năm 1991.

Chiến lược chia tập huấn luyện/kiểm tra:

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).

Phương pháp kiểm định chéo theo chuỗi thời gian:

Bên cạnh chia cố định train/test, nghiên cứu còn áp dụng kiểm định chéo với cửa sổ mở rộng để đánh giá mô hình trong điều kiện chuỗi thời gian. Cấu trúc các folds như sau:

- **Đối với các quốc gia có dữ liệu từ 1980–2019:**
 - Fold 1: Huấn luyện = 1980–1999, Kiểm tra = 2000–2004;
 - Fold 2: Huấn luyện = 1980–2004, Kiểm tra = 2005–2009;
 - Fold 3: Huấn luyện = 1980–2009, Kiểm tra = 2010–2014;
 - Fold 4: Huấn luyện = 1980–2014, Kiểm tra = 2015–2019.
- **Đối với Nga (dữ liệu từ 1991–2019):**

- Fold 1: Huấn luyện = 1991–2010, Kiểm tra = 2011-2014;
- Fold 2: Huấn luyện = 1991–2014, Kiểm tra = 2015-2019;

Phương pháp kiểm định chéo có ý thức thời gian này đảm bảo rằng mô hình được đánh giá trong các điều kiện kinh tế khác nhau đồng thời duy trì quan hệ nhân quả.

4.5.2. Cài đặt môi trường thực nghiệm

Toàn bộ thực nghiệm được thực hiện trên một máy trạm chuyên dụng với cấu hình phần cứng: 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) .

Quy trình thực hiện

Trình tự các bước được thực hiện trong phương pháp này như sau:

- (1) Trích xuất dữ liệu từ cơ sở dữ liệu kinh tế vĩ mô xây dựng ở chương 2.
- (2) Phân đoạn chu kỳ kinh tế thành bốn pha: suy thoái, phục hồi, tăng trưởng và tăng trưởng mạnh, dựa trên GDP.
- (3) Huấn luyện mô hình PAA-LSTM bằng chiến lược tinh chỉnh theo pha, nhằm tối ưu hóa khả năng thích ứng của mô hình với đặc điểm riêng của từng pha kinh tế.
- (4) Dự báo và so sánh hiệu suất dự báo của mô hình với các mô hình đối chứng: ARIMA, XGBoost, Transformer, LSTM, Bi-LSTM.

Cài đặt các tham số

Đối với cấu hình tham số của mô hình, 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 trong Bảng 4.1 sau:

Bảng 4. 1 Kết quả của điều chỉnh tham số trong mô hình theo tập dữ liệu

Tham số	Giá trị
Số đơn vị ẩn (Số nơ-ron)	10
Hàm kích hoạt (activation)	Relu
Số vòng lặp (epochs)	500
Số lô (batch_size)	64
Bộ tối ưu (optimizer)	Adam
Hàm mất mát (loss)	mse

Cài đặt các mô hình so sánh

ARIMA: Đối với mỗi quốc gia, thuật toán auto-ARIMA được sử dụng để tự động xác định cấu hình mô hình tối ưu dựa trên tiêu chí thông tin Akaike (AIC). Phương pháp này đảm bảo tính đơn giản của mô hình đồng thời vẫn nắm bắt được các phụ thuộc theo thời gian trong chuỗi dữ liệu.

XGBoost: Mô hình tăng cường độ dốc được tinh chỉnh bằng cách tìm kiếm lưới (grid search) trên các siêu tham số quan trọng như max_depth, learning_rate, và n_estimators, kết hợp với kỹ thuật kiểm định chéo chuỗi thời gian (time-series cross-validation) để tránh quá khớp và đảm bảo khả năng tổng quát hóa.

Transformer: Kiến trúc Transformer được triển khai với hai lớp encoder, mỗi lớp có bốn đầu attention. Một lớp fully-connected được thêm vào cùng với kỹ thuật dropout để giảm thiểu overfitting. Các siêu tham số được hiệu chỉnh nhằm cân bằng giữa độ phức tạp của mô hình và hiệu suất dự báo, đặc biệt trong bối cảnh dữ liệu kinh tế vĩ mô có kích thước tương đối nhỏ.

LSTM, Bi-LSTM: Các mô hình này được triển khai với cùng độ dài chuỗi đầu vào, kích thước lớp ẩn, và quy trình tối ưu hóa thống nhất.

4.5.3. 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, luận án sử dụng chỉ số đánh giá chính là RMSE theo công thức (1.9).

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. Hơn nữa, RMSE giữ nguyên đơn vị của biến gốc, giúp dễ dàng diễn giải và ứng dụng vào thực tiễn kinh tế.

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).

4.5.4. Kết quả thực nghiệm và so sánh

Kết quả thực nghiệm được trình bày thông qua bảng và biểu đồ đánh giá hiệu năng của các mô hình tại sáu quốc gia. Các mô hình được so sánh bao gồm: ARIMA, XGBoost, LSTM, Bi-LSTM và Transformer.

Bảng 4.2 cùng với Hình 4.5 và Hình 4.6 trình bày chi tiết kết quả dự báo GDP của Việt Nam theo từng mô hình. Trong số các mô hình được đánh giá, PAA-LSTM đạt kết quả tốt nhất với $RMSE = 0.65$, $MAE = 0.48$ và $R^2 = 0.94$ cho thấy năng lực mô hình hóa vượt trội của mô hình đề xuất trong bối cảnh dữ liệu kinh tế đầy biến động.

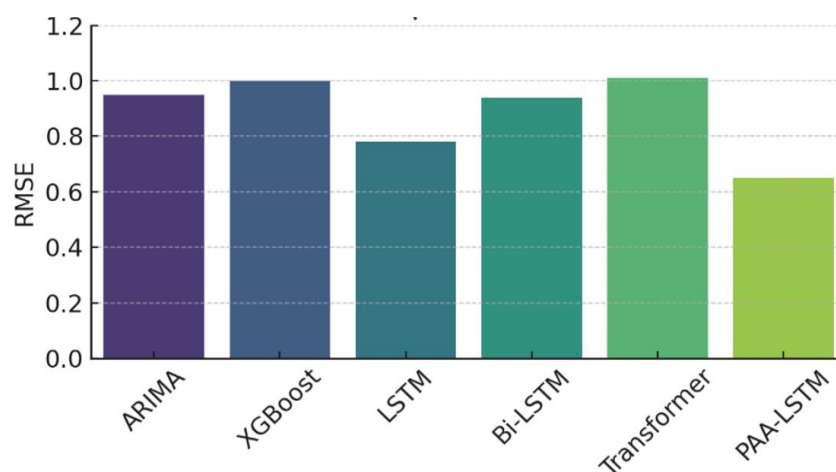
So với mô hình LSTM truyền thống ($RMSE = 0.78$), mô hình PAA-LSTM cải thiện rõ rệt độ chính xác nhờ vào cơ chế attention thích ứng theo pha, cho phép mô hình tập trung hiệu quả hơn vào các thời điểm then chốt. Bên cạnh đó, PAA-LSTM cũng vượt trội hơn so với các mô hình học máy truyền thống như ARIMA ($RMSE = 0.95$) và XGBoost ($RMSE = 1.00$) – vốn không có khả năng học các phụ thuộc phi tuyến và dài hạn.

Dù Transformer là kiến trúc tiên tiến nhưng vẫn có $RMSE = 1.01$ cao hơn đáng kể so với PAA-LSTM. Điều này cho thấy việc tích hợp thông tin pha kinh tế vào cơ chế attention giúp các mô hình học sâu nắm bắt tín hiệu vĩ mô hiệu quả hơn trong các giai đoạn kinh tế khác nhau. Kết quả tại Việt Nam thể hiện rõ lợi thế của cơ chế

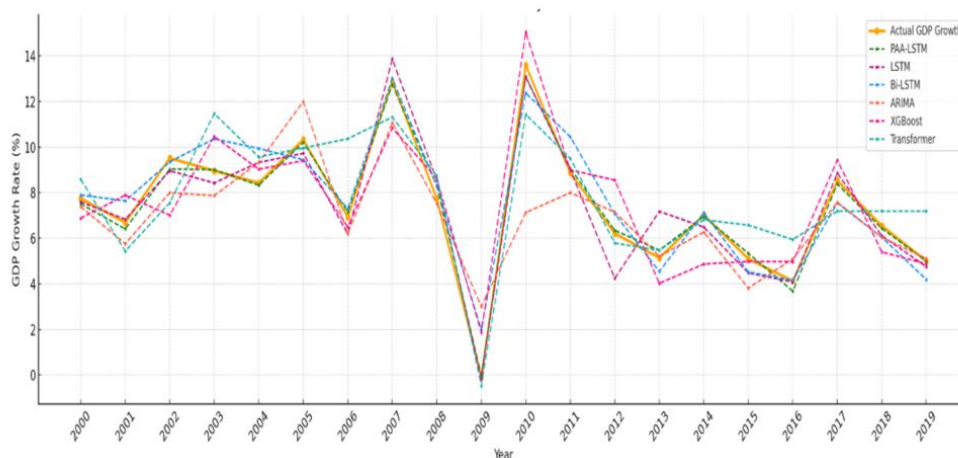
attention thích ứng theo pha trong bối cảnh các nền kinh tế đang phát triển, nơi dữ liệu thường thiếu ổn định và chịu ảnh hưởng mạnh bởi các yếu tố chu kỳ.

Bảng 4. 2 Kết quả dự báo GDP của Việt Nam

Model	RMSE	MAE	R ²
ARIMA	0.95	0.73	0.76
XGBoost	1.0	0.97	0.84
LSTM	0.78	0.61	0.75
Bi-LSTM	0.94	1.00	0.79
Transformer	1.01	0.83	0.82
PAA-LSTM	0.65	0.48	0.94



Hình 4. 5 So sánh RMSE của các mô hình khi dự báo GDP của Việt Nam



Hình 4. 6 Kết quả dự báo GDP của Việt Nam

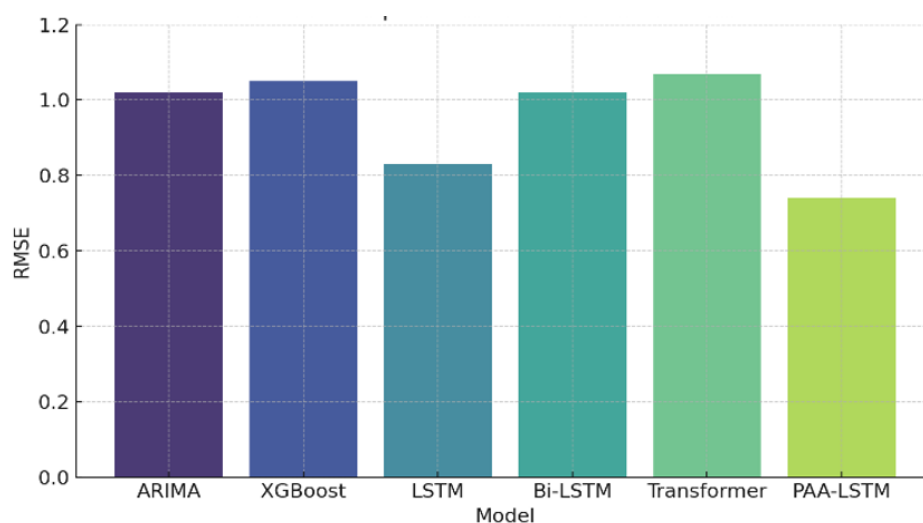
Sau khi đánh giá hiệu suất của mô hình tại Việt Nam, chúng tôi tiếp tục tiến hành kiểm định khả năng dự báo của mô hình tại Ấn Độ - một quốc gia đang phát triển và cũng là đại diện tiêu biểu cho nhóm nền kinh tế mới nổi. Bảng 4.3, cùng với Hình 4.7 và Hình 4.8, trình bày chi tiết kết quả dự báo trong bối cảnh này. PAA-LSTM tiếp tục thể hiện hiệu năng vượt trội, đạt $RMSE = 0.74$, $MAE = 0.48$ và $R^2 = 0.95$ vượt trội hơn tất cả các mô hình đối chứng.

Đáng chú ý, so với Bi-LSTM ($RMSE = 1.02$) và Transformer ($RMSE = 1.07$), mô hình đề xuất đã giảm đáng kể sai số dự báo. Điều này nhấn mạnh vai trò quan trọng của cơ chế attention thích ứng theo pha, giúp mô hình học được các hành vi kinh tế có tính chu kỳ và biến động theo thời gian.

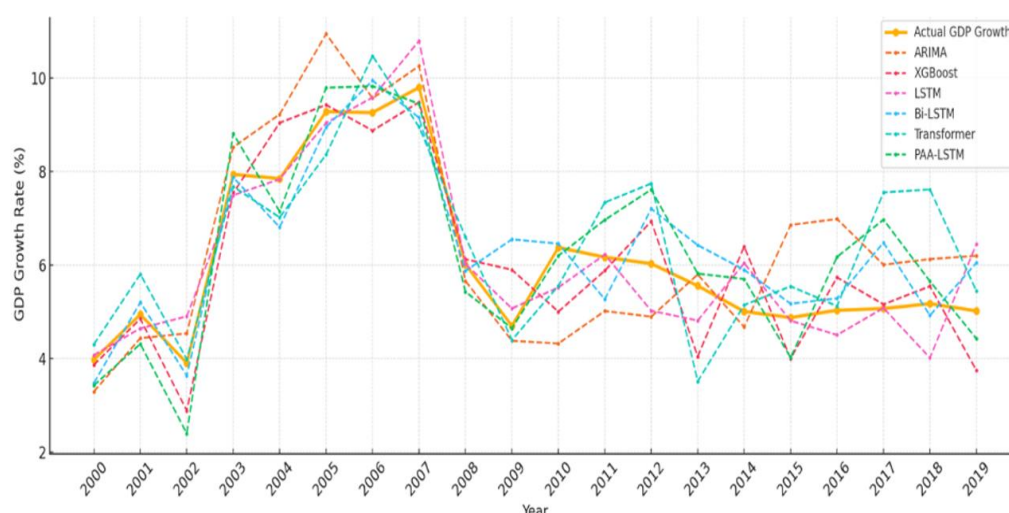
Với khả năng linh hoạt trong thích nghi với các thay đổi cấu trúc, PAA-LSTM đặc biệt phù hợp với các môi trường kinh tế năng động như Ấn Độ, nơi các chỉ số như đầu tư, sản xuất và tiêu dùng thường xuyên biến động dưới tác động của cả chính sách và lực lượng thị trường.

Bảng 4. 3 Kết quả dự báo GDP của Ấn Độ

Model	RMSE	MAE	R²
ARIMA	1.02	0.78	0.83
XGBoost	1.05	0.95	0.87
LSTM	0.83	0.66	0.77
Bi-LSTM	1.02	0.86	0.84
Transformer	1.07	0.91	0.84
PAA-LSTM	0.74	0.48	0.95



Hình 4. 7 So sánh RMSE của các mô hình khi dự báo GDP của Ấn Độ



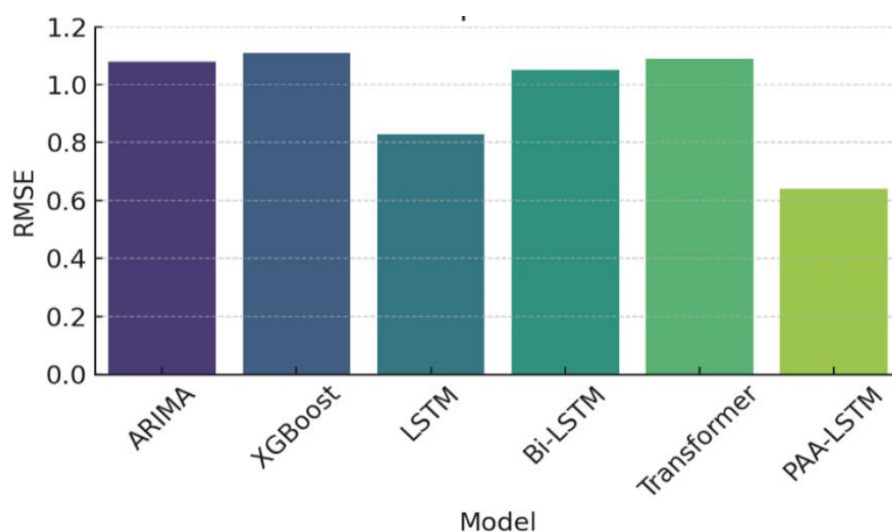
Hình 4. 8 Kết quả dự báo GDP của Ấn Độ

Tiếp theo, kết quả dự báo GDP cho Trung Quốc và Nga tiếp tục củng cố hiệu quả của mô hình PAA-LSTM được đề xuất. Tại Trung Quốc, như được trình bày trong Bảng 4.4 và Hình 4.9, Hình 4.10, mô hình PAA-LSTM đạt $RMSE = 0.64$, thấp hơn rõ rệt so với các mô hình khác như Transformer ($RMSE = 1.09$), Bi-LSTM (1.05) và XGBoost (1.11).

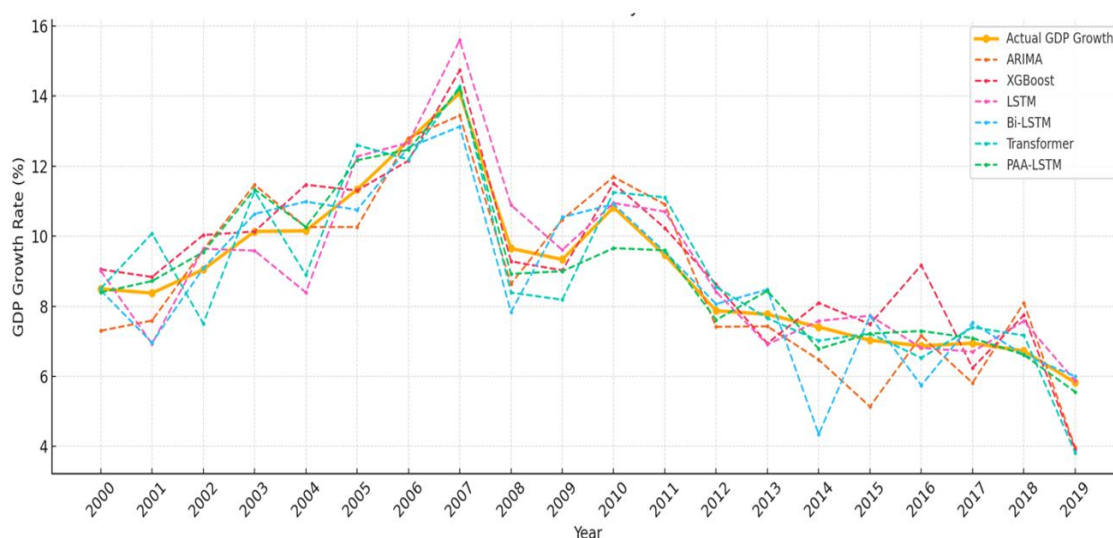
Những kết quả này cho thấy khả năng mạnh mẽ của mô hình trong việc nắm bắt các mẫu chu kỳ đặc trưng của một nền kinh tế đang tăng trưởng nhanh và có cấu trúc phức tạp như Trung Quốc.

Bảng 4. 4 Kết quả dự báo GDP của Trung Quốc

Model	RMSE	MAE	R ²
ARIMA	1.08	0.9	0.78
XGBoost	1.11	0.71	0.85
LSTM	0.83	0.55	0.85
Bi-LSTM	1.05	0.78	0.8
Transformer	1.09	0.86	0.82
PAA-LSTM	0.64	0.5	0.8



Hình 4. 9 So sánh RMSE của các mô hình khi dự báo GDP của Trung Quốc



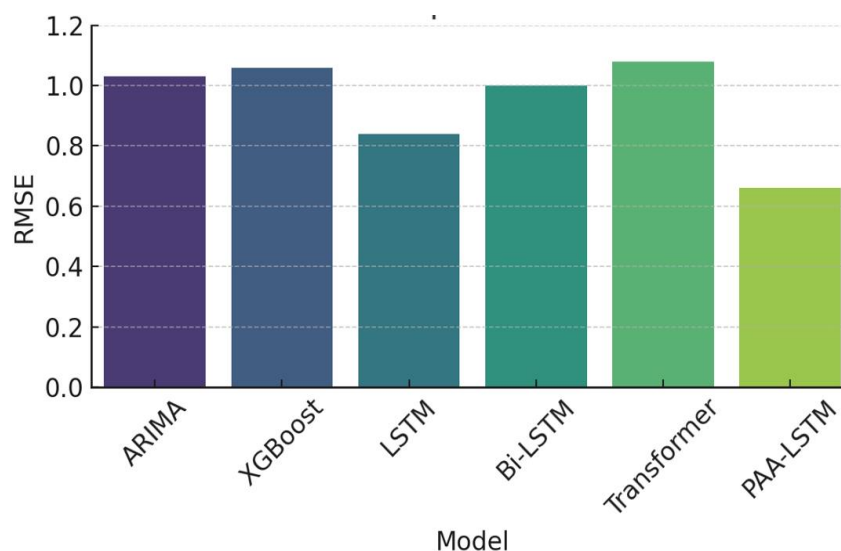
Hình 4. 10 Kết quả dự báo GDP của Trung Quốc

Đối với Nga, Bảng 4.5, Hình 4.11 và Hình 4.12 trình bày kết quả dự báo GDP cho một nền kinh tế mới nổi khác. Mô hình PAA-LSTM đạt $RMSE = 0.81$, thấp hơn so với các mô hình Transformer (1.08), Bi-LSTM (1.00) và ARIMA (1.03).

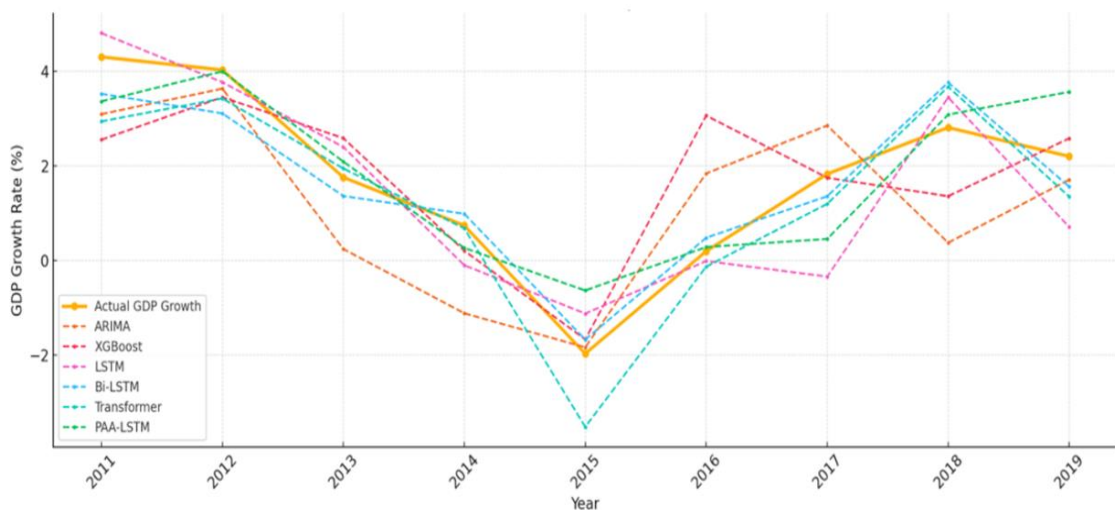
Những phát hiện này phù hợp với kỳ vọng rằng các nền kinh tế dễ bị tác động bởi các cú sốc bên ngoài, chẳng hạn như biến động giá năng lượng và căng thẳng địa chính trị, sẽ hưởng lợi đáng kể từ các mô hình có khả năng thích ứng linh hoạt với các pha khác nhau của chu kỳ kinh tế.

Bảng 4. 5 Kết quả dự báo GDP của Nga

Model	RMSE	MAE	R²
ARIMA	1.03	0.9	0.77
XGBoost	1.06	0.83	0.78
LSTM	0.84	0.61	0.78
Bi-LSTM	1.0	0.9	0.91
Transformer	1.08	0.96	0.83
PAA-LSTM	0.69	0.49	0.89



Hình 4. 11 So sánh RMSE của các mô hình khi dự báo GDP của Nga



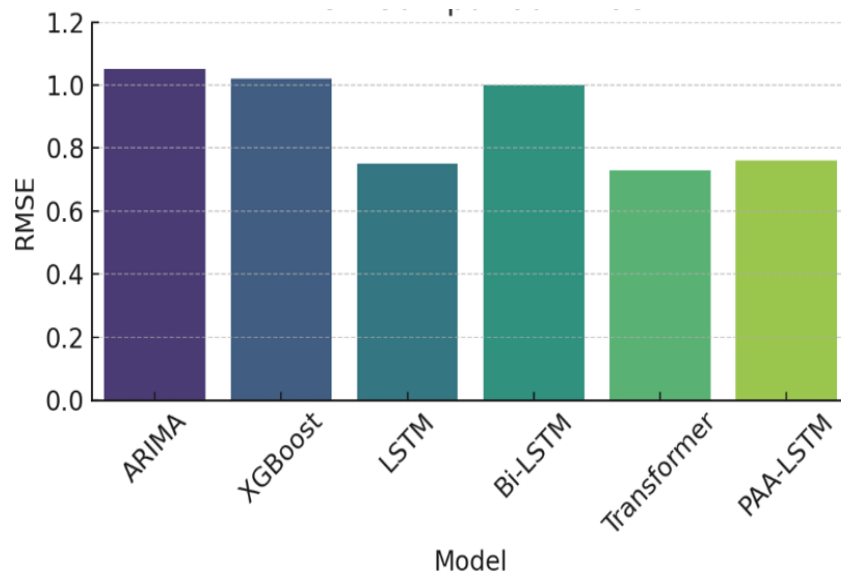
Hình 4. 12 Kết quả dự báo GDP của Nga

Cuối cùng, kết quả dự báo GDP cho Hoa Kỳ và Canada - hai nền kinh tế phát triển - cho thấy một xu hướng khác biệt. Đối với Hoa Kỳ, Bảng 4.6, hình 4.13 và hình 4.14 trình bày kết quả của mô hình PAA-LSTM, với $RMSE = 0.76$, thấp hơn một chút so với Transformer ($RMSE = 0.73$) và LSTM ($RMSE=0.75$) về độ lỗi.

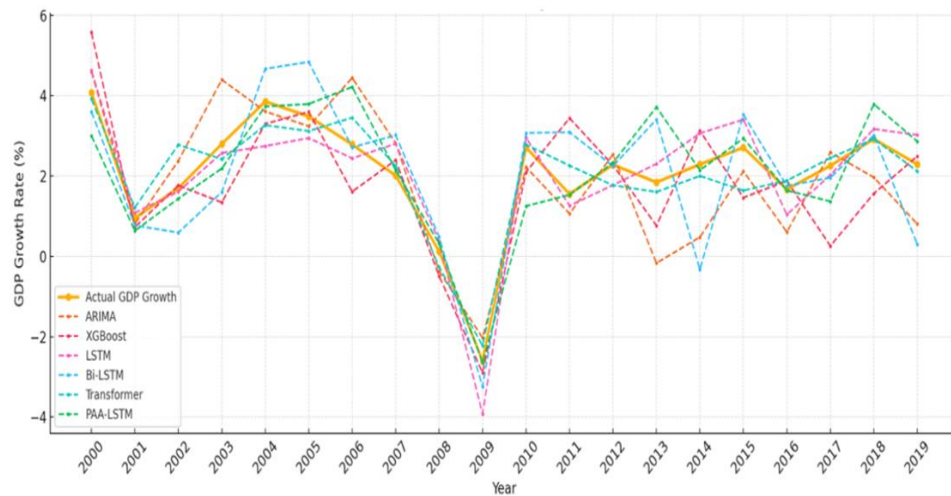
Trong khi đó, các mô hình truyền thống như ARIMA ($RMSE = 1.05$) và XGBoost ($RMSE = 1.02$) cho thấy mức sai số cao hơn đáng kể. Điều này gợi ý rằng trong môi trường dữ liệu ổn định, các kiến trúc cơ bản như LSTM hoặc Transformer vẫn có thể hoạt động tốt mà không nhất thiết cần đến attention thích ứng theo pha.

Bảng 4. 6 Kết quả dự báo GDP của Mỹ

Model	RMSE	MAE	R^2
ARIMA	1.05	0.8	0.79
XGBoost	1.02	0.79	0.87
LSTM	0.75	0.57	0.81
Bi-LSTM	1	0.9	0.86
Transformer	0.73	0.98	0.89
PAA-LSTM	0.76	0.47	0.76



Hình 4. 13 So sánh RMSE của các mô hình khi dự báo GDP của Mỹ

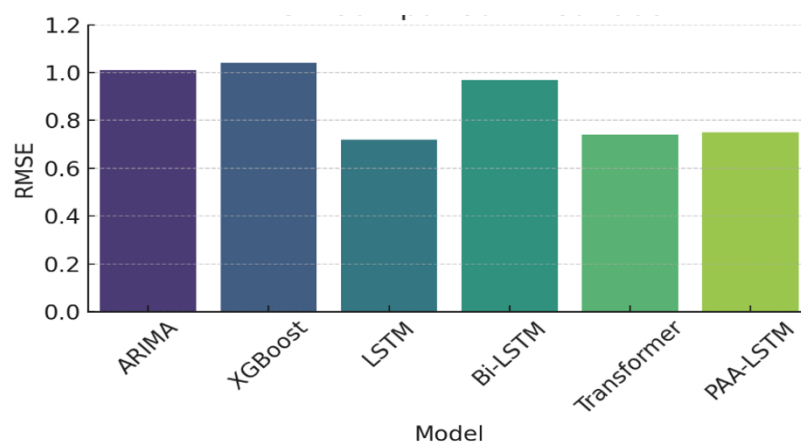


Hình 4. 14 Kết quả dự báo GDP của Mỹ

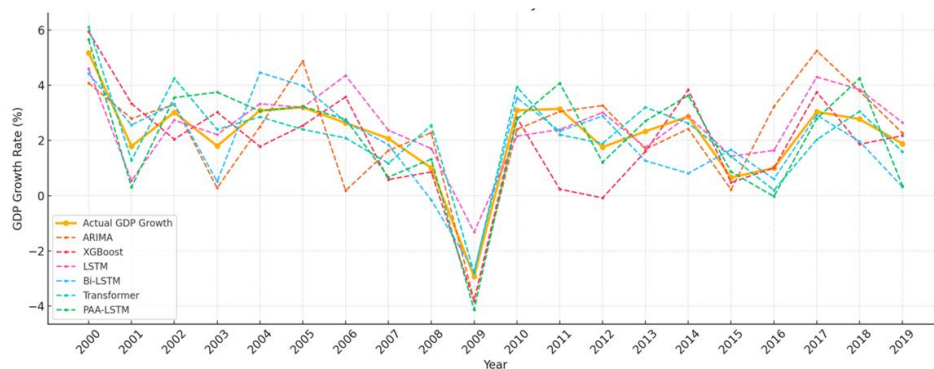
Trong trường hợp của Canada, bảng 4.7, hình 4.15 và hình 4.16 trình bày kết quả dự báo cho thấy một bức tranh cạnh tranh giữa các mô hình được đánh giá. Mô hình LSTM đạt giá trị RMSE thấp nhất là 0.72, theo sát là Transformer (0.74) và PAA-LSTM (0.75). Kết quả thể hiện khoảng cách giữa các mô hình tại Canada không lớn như ở nhóm nước mới nổi.

Bảng 4. 7 Kết quả dự báo GDP của Canada

Model	RMSE	MAE	R ²
ARIMA	1.01	0.73	0.76
XGBoost	1.04	0.79	0.82
LSTM	0.72	0.59	0.82
Bi-LSTM	0.97	0.83	0.77
Transformer	0.74	0.98	0.82
PAA-LSTM	0.75	0.48	0.9



Hình 4. 15 So sánh RMSE của các mô hình khi dự báo GDP của Canada



Hình 4. 16 Kết quả dự báo GDP của Canada.

Những kết quả này cho thấy rằng mặc dù PAA-LSTM không phải lúc nào cũng vượt trội về mặt sai số tuyệt đối so với các mô hình khác, nhưng mô hình này liên tục thể hiện sự cạnh tranh trong dự báo kinh tế vĩ mô. Đối với các quốc gia có

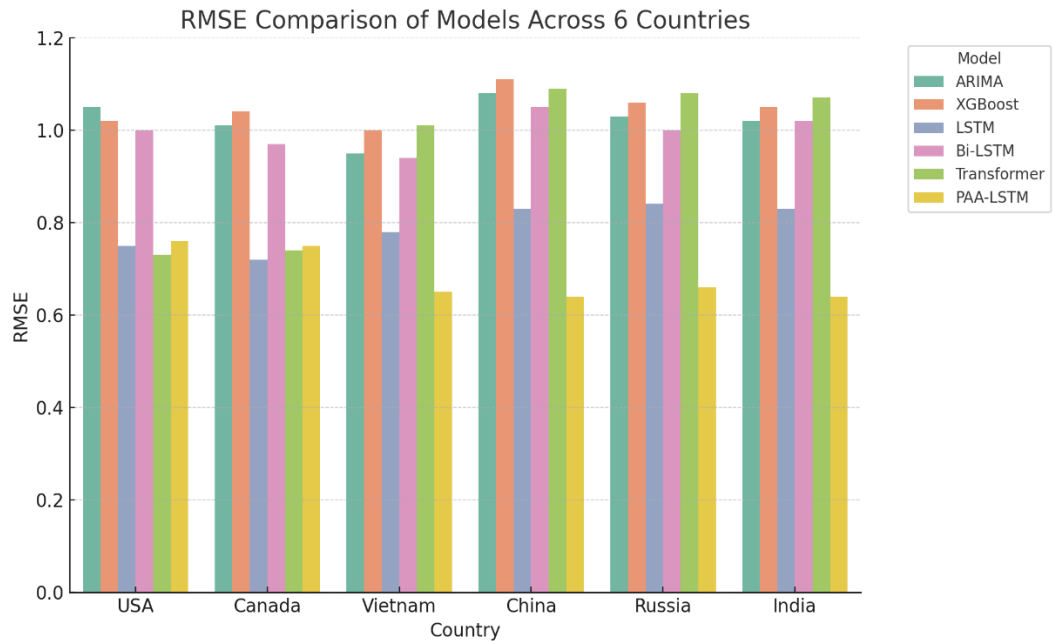
nền kinh tế không ổn định, PAA-LSTM cho kết quả dự báo tốt hơn rõ rệt. Đối với các quốc gia có nền kinh tế ổn định, mô hình PAA-LSTM vẫn đạt được kết quả rất khả quan, đặc biệt là về khả năng khái quát hóa và tính ổn định. Nhìn chung, PAA-LSTM thể hiện năng lực dự báo GDP mạnh mẽ và linh hoạt trong nhiều điều kiện kinh tế khác nhau, qua đó củng cố giá trị của việc tích hợp động lực chu kỳ kinh tế vào cơ chế attention trong các mô hình dự báo kinh tế vĩ mô.

4.6. Phân tích hiệu năng và ý nghĩa thực tiễn

Kết quả thực nghiệm trên sáu quốc gia - đại diện cho các nền kinh tế phát triển, mới nổi và đang phát triển - như thể hiện trong bảng 4.8 và hình 4.17, cho thấy hiệu suất mạnh mẽ và khả năng thích ứng cao của mô hình PAA-LSTM. Phần này sẽ mở rộng phân tích các kết quả đó bằng cách so sánh với các nghiên cứu trước đây, đồng thời đánh giá mức độ mà cách tiếp cận hiện tại phù hợp hoặc khác biệt so với các công trình nghiên cứu đã có.

Bảng 4. 8 Bảng RMSE của các mô hình với 6 quốc gia

Model	Việt Nam	Ấn Độ	Trung Quốc	Nga	Mỹ	Canada
ARIMA	0.95	1.02	1.08	1.03	1.05	1.01
XGBoost	1.0	1.05	1.11	1.06	1.02	1.04
LSTM	0.78	0.83	0.83	0.84	0.75	0.72
Bi-LSTM	0.94	1.02	1.05	1.0	1	0.97
Transformer	1.01	1.07	1.09	1.08	0.73	0.74
PAA-LSTM	0.65	0.74	0.64	0.69	0.76	0.75



Hình 4. 17 So sánh RMSE của các mô hình với 6 quốc gia

Mô hình PAA-LSTM liên tục đạt được giá trị RMSE thấp nhất hoặc gần thấp nhất ở các nền kinh tế biến động mạnh như Việt Nam, Trung Quốc và Ấn Độ. Điều này xác nhận các quan sát trước đây trong nghiên cứu của Qiu và cộng sự năm 2020 [112_1] và nghiên cứu của Qin và cộng sự năm 2017 [112], những người đã nhấn mạnh giá trị của cơ chế attention trong việc nắm bắt các động học thời gian phức tạp. Tuy nhiên, trong khi các mô hình trước đó sử dụng attention tĩnh, kết quả của chúng tôi chứng minh rằng việc tích hợp thông tin về pha kinh tế giúp cải thiện độ chính xác dự báo, đặc biệt trong các giai đoạn chuyển tiếp như suy thoái hoặc phục hồi.

Khác với các mô hình truyền thống như ARIMA và XGBoost, vốn giả định tuyến tính và tính bất biến theo thời gian, mô hình của chúng tôi thể hiện hiệu suất vượt trội trong hầu hết các trường hợp. Điều này phù hợp với các phát hiện trong nghiên cứu của Oancea công bố năm 2025 [104], nghiên cứu của Chen và cộng sự năm công bố 2025 [38] và nghiên cứu của Jallow và cộng sự công bố năm 2025 [80], những người cho rằng các mô hình học sâu thường vượt trội hơn so với kỹ thuật kinh tế lượng tuyến tính. Đáng chú ý, trong thí nghiệm của chúng tôi, ARIMA và XGBoost

liên tục nằm trong nhóm mô hình có hiệu suất kém hơn, với RMSE thường vượt quá 1.0.

Khi so sánh với các nghiên cứu gần đây ứng dụng mô hình dựa trên LSTM, kết quả của chúng tôi củng cố lợi thế của attention thích ứng theo pha. Ví dụ, nghiên cứu của Xie và cộng sự công bố năm 2024 [136] chỉ ra rằng hiệu suất của mô hình học sâu thay đổi theo từng quốc gia tùy thuộc vào độ phức tạp kinh tế và sự ổn định của dữ liệu. Mô hình của chúng tôi mở rộng điều này bằng cách điều chỉnh động theo các mẫu chu kỳ đặc trưng theo từng quốc gia, mang lại chỉ số RMSE nhỏ hơn trong nhiều bối cảnh khác nhau.

Mô hình Transformer, dù rất thành công trong xử lý ngôn ngữ tự nhiên và các lĩnh vực khác, lại cho kết quả không nhất quán trong dự báo kinh tế vĩ mô theo nghiên cứu của chúng tôi. Điều này phù hợp với nhận định của Li cộng sự công bố năm 2025 [91], rằng Transformer đòi hỏi tập dữ liệu lớn và có thể gặp khó khăn khi dữ liệu kinh tế thưa hoặc gián đoạn.

Cần nhấn mạnh rằng các pha kinh tế vốn là biến tiềm ẩn không thể quan sát trực tiếp theo thời gian thực mà chỉ có thể xấp xỉ dựa trên các chỉ báo có sẵn. Hạn chế này đã được thừa nhận rộng rãi trong các mô hình chuyển chế độ cổ điển, điển hình là mô hình của Hamilton (1989) [66], vốn coi các chế độ kinh tế là không quan sát được và được suy diễn thông qua các cấu trúc xác suất. Dựa trên cơ sở lý thuyết này, nghiên cứu của chúng tôi áp dụng cách tiếp cận thực tiễn bằng cách xác định các pha kinh tế thông qua ngưỡng quy tắc đơn giản áp dụng cho tốc độ tăng trưởng GDP thực. Những tín hiệu thực nghiệm này đóng vai trò là đại diện hiệu quả cho các trạng thái chu kỳ tiềm ẩn, đồng thời dễ dàng tích hợp vào quy trình học sâu. Khác với các mô hình chuyển chế độ truyền thống vốn đòi hỏi các thủ tục suy luận phức tạp, cách tiếp cận của chúng tôi cho phép mô hình hóa động lực chu kỳ một cách hiệu quả và minh bạch, từ đó tăng cường khả năng diễn giải và tính ứng dụng trong các bài toán dự báo thực tế.

Một điểm đáng lưu ý khác là kết quả cho thấy Bi-LSTM không mang lại lợi thế rõ rệt so với LSTM tiêu chuẩn trong bối cảnh kinh tế vĩ mô. Điều này tương đồng

với kết luận trong nghiên cứu của Fan¹ và cộng sự năm 2024, rằng quan hệ nhân quả có hướng trong các chỉ số kinh tế vĩ mô có thể hạn chế lợi ích của kiến trúc hai chiều.

Đối với các nền kinh tế phát triển có độ ổn định cao như Hoa Kỳ và Canada, mô hình PAA-LSTM vẫn đạt được mức độ chính xác dự báo tương đối tốt; tuy nhiên, sự chênh lệch về hiệu suất so với các mô hình đơn giản hơn như LSTM cơ bản hoặc Transformer có xu hướng thu hẹp đáng kể. Kết quả này phản ánh đặc điểm của môi trường dữ liệu có mức độ biến động thấp, nơi các chỉ tiêu kinh tế thay đổi tương đối mượt mà theo thời gian và hiếm khi xuất hiện các cú sốc cấu trúc rõ rệt. Trong những bối cảnh như vậy, cơ chế attention thích ứng theo pha vốn được thiết kế để tận dụng các đặc điểm chuyển tiếp giữa các giai đoạn của chu kỳ kinh tế không còn phát huy được lợi thế rõ ràng. Việc bổ sung thêm độ phức tạp kiến trúc trong mô hình PAA-LSTM, chẳng hạn như phân pha chu kỳ và tái phân bổ trọng số attention theo thời gian, có thể không mang lại sự cải thiện hiệu suất tương xứng, thậm chí làm gia tăng nguy cơ quá khớp. Do đó, trong các môi trường ổn định này, các mô hình có cấu trúc đơn giản hơn nhưng có năng lực học ổn định chuỗi thời gian, như LSTM hoặc Transformer thông thường, có thể đạt được kết quả dự báo tương đương hoặc tốt hơn về mặt RMSE.

Tổng thể, chương này mở rộng lĩnh vực học sâu trong dự báo KTVM bằng cách chứng minh rằng attention thích ứng theo pha đặc biệt có giá trị trong các môi trường kinh tế vĩ mô biến động mạnh. Đồng thời cũng gợi ý rằng các nghiên cứu trong tương lai nên xem xét chiến lược chọn mô hình thích ứng, tức là lựa chọn mô hình dựa trên đặc điểm cấu trúc của nền kinh tế được dự báo.

4.7. Kết luận chương

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

¹ <https://pubmed.ncbi.nlm.nih.gov/38931746/>

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

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 đánh giá các phương pháp và đề xuất mô hình 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:

1. Mở rộng mô hình sang các chỉ tiêu kinh tế khác như lạm phát, cán cân thương mại, tỷ lệ thất nghiệp, v.v. nhằm kiểm chứng tính linh hoạt và độ khái quát hóa của mô hình PAA-LSTM.
2. 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.
3. 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 (ICACT)*, page 22-233, IEEE, 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, "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)

TÀI LIỆU THAM KHẢO

Tiếng Việt

1. Phạm Thành Công, Lý Quang Hùng (2025), “Dự báo tăng trưởng kinh tế Việt Nam năm 2025: So sánh mô hình GARCH và MIDAS,” *Tạp chí Kinh tế và Dự báo*. Truy cập từ: <https://kinhtevadubao.vn/du-bao-tang-truong-kinh-te-viet-nam-nam-2025-so-sanh-mo-hinh-garch-va-midas1-30852.html>.
2. Nguyễn Thị Hiên, Lê Mai Trang, Phạm Long, Vũ Phạm Văn Duy, Hoàng Lê Quang Huy (2023), “Ứng dụng mô hình ARIMA để dự báo lạm phát của Việt Nam và một số khuyến nghị,” *Tạp chí Ngân hàng*. Truy cập từ: <https://tapchinganhang.gov.vn/ung-dung-mo-hinh-arima-de-du-bao-lam-phat-cua-viet-nam-va-mot-so-khuyen-nghi-12311.html>.
3. Nguyễn Thị Hiên, Đặng Thị Minh Nguyệt (2022), “Ứng dụng mô hình ARCH, GARCH phân tích độ biến động của hợp đồng tương lai VN30F1M trên thị trường chứng khoán phái sinh Việt Nam,” *Tạp chí Ngân hàng*. Truy cập từ: <https://tapchinganhang.gov.vn/ung-dung-mo-hinh-arch-garch-phan-tich-do-bien-dong-cua-hop-dong-tuong-lai-vn30f1m-tren-thi-truong-chung-khoan-phai-sinh-viet-nam-10638.html>.
4. Nguyễn Thanh Hương, Bùi Quang Trung (2015), “Ứng dụng mô hình kết hợp arima-garch để dự báo chỉ số vn-index,” *Tạp chí Khoa học và Công nghệ Đại học Đà Nẵng*, 4(89), 118–122.
5. Trần Minh Quân (2021), “Ứng dụng mô hình MIDAS để dự báo lạm phát ở Việt Nam,” *Tạp chí Tài chính*, 29(2), 112–124.
6. Phùng Duy Quang (2024), “Ứng dụng mô hình arch và garch dự báo giá gạo tại Thái Lan,” *FTU Working Paper Series*. Truy cập từ: <https://fwps.ftu.edu.vn/2024/07/22/ung-dung-mo-hinh-arch-va-garch-du-bao-gia-gao-tai-thai-lan/>.
7. Đỗ Văn Thành (2017), “Ứng dụng mô hình GARCH trong kiểm định hiệu quả thị trường và dự báo giá cổ phiếu,” *Tạp chí Phát triển Kinh tế*, số 9, 2017.
8. Đỗ Văn Thành, Nguyễn Minh Hải (2016A), “Ứng dụng mô hình ARDL trong dự báo VN-Index theo ngày bằng phát hiện nhân quả Granger,” *Tạp chí Khoa học và Đào tạo Ngân hàng*, 2016A.
9. Đỗ Văn Thành, Nguyễn Minh Hải (2016B), “Giảm chiều dữ liệu dựa trên phân tích thành phần chính trong dự báo VN-Index,” *Tạp chí Kinh tế và Dự báo*, 2016B.
10. Cù Thu Thủy, Trần Quang Huy, Nguyễn Quốc Hưng (2017), “Ứng dụng PCA và hồi quy tuyến tính trong dự báo chỉ số giá tiêu dùng,” *Tạp chí Tài chính*, số 10, 2017.

11. Lê Mai Trang, Hoàng Anh Tuấn, Nguyễn Thị Hiền Đinh Thị Hà, Trần Kim Anh (2022), “Ứng dụng mô hình MIDAS để dự báo tăng trưởng xuất khẩu của Việt Nam,” *Tạp chí Ngân hàng*. Truy cập từ: <https://tapchinganhang.gov.vn/ung-dung-mo-hinh-midas-de-du-bao-tang-truong-xuat-khau-cua-viet-nam-13382.html>.
12. Lê Văn Tuấn, Phùng Duy Quang (2020), “Áp dụng mô hình GARCH dự báo ảnh hưởng của đại dịch Covid-19 đến thị trường chứng khoán Việt Nam,” *Tạp chí Công thương*. Truy cập từ: <https://tapchicongthuong.vn/ap-dung-mo-hinh-garch-du-bao-anh-huong-cua-dai-dich-covid-19-den-thi-truong-chung-koan-viet-nam-75518.htm>.
13. Nguyễn Thanh Tùng, Phạm Thị Lan Anh (2022), “Dự báo tăng trưởng GDP Việt Nam bằng mô hình MIDAS kết hợp dữ liệu tần suất cao,” *Tạp chí Nghiên cứu Kinh tế*, 39(5), 67–79.

Tiếng Anh

14. Adolfson M., Laséen S., Lindé J., Villani M. (2007), “Bayesian estimation of an open economy DSGE model with incomplete pass-through”, *Journal of International Economics*, 72(2), pp. 481–511. [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S0022199607000219>
15. Aisy R. R., Zulfa L., Rahim Y., Ahsan M. (2025), “Residual XGBoost regression - Based individual moving range control chart for Gross Domestic Product growth monitoring,” *PLOS ONE*. [Online]. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC12063872/>
16. Akkemik K.A. (2007), “The response of employment to GDP growth in Turkey: an econometric estimation”, *Applied Econometrics and International Development* 7(1), pp. 1–7.
17. Atashbar T., Shi R. A. (2022), *Deep Reinforcement Learning: Emerging Trends in Macroeconomics and Future Prospects*, IMF Working Paper, WP/22/259, International Monetary Fund, Washington, DC, USA, December.
18. Athey S. (2018), “The impact of machine learning on economics”, in *The economics of artificial intelligence: An agenda*, University of Chicago Press, pp. 507–547.
19. Azman S., Pathmanathan D., Balakrishnan V. (2025), “Leveraging Machine Learning for Enhanced Predictive Accuracy in Time Series Forecasting: A Comparative Analysis of LSTM and GRU Models”, *PLOS ONE*, 20(5), e0323015. [Online]. Available: <https://www.researchgate.net/publication/386461159>
20. Bahdanau D., Cho K., Bengio Y. (2014), “Neural machine translation by jointly learning to align and translate”, *arXiv preprint*, arXiv:1409.0473.

21. Baldo F. (2024), *Informed Machine Learning for Epidemics from Data Analysis to Time-Series Forecasting*, University of Bologna, Italy.
22. Banerjee M., Marcellino M., Masten I. (2005), “Forecasting macroeconomic variables using diffusion indexes in short samples with structural change”, *CEPR Discussion Paper*, no. 5315.
23. Bantis E. (2024), *Essays on Economic Forecasting with Alternative Datasets*, Ph.D. thesis, University of Reading, January.
24. Batsuuri T., He S., Hu R., Leslie J., Lutz F. (2024), “Predicting IMF Supported Programs: A Machine Learning Approach,” *IMF Working Paper*, no. 2024/054, March. [Online]. Available: <https://www.elibrary.imf.org/downloadpdf/view/journals/001/2024/054/001.2024.issue-054-en.pdf>
25. Beg K., Raza M. T., Padmapriya B., Shajar S. N. (2025), “Comparative analysis of machine learning models for sectoral volatility prediction in financial markets,” *Journal of Information Systems Engineering and Management*, vol. 10, no. 31s. [Online]. Available: <https://www.jisem-journal.com/index.php/journal/article/download/5137/2418/8560>
26. Bekaert G., Harvey C.R., Lundblad C., Siegel S. (2007), “Global growth opportunities and market integration”, *The Journal of Finance* 62(3), pp. 1081–1137.
27. Bohlke-Schneider, M., Kapoor, S., & Januschowski, T. (2022). *Resilient neural forecasting systems*. arXiv. <https://arxiv.org/abs/2203.08492>
28. Bollerslev T. (1986), *Generalized Autoregressive Conditional Heteroskedasticity*, *Journal of Econometrics*, vol. 31, no. 3, pp. 307–327.
29. Box G. E. P., Jenkins G. M. (1970), *Time Series Analysis: Forecasting and Control*, Holden-Day, San Francisco, CA, USA.
30. Breiman L. (2001), “Random forests,” *Machine Learning*, vol. 45, no. 1, pp. 5–32.
31. Breiman L., Friedman J. H., Olshen R. A., Stone C. J. (2017), *Classification and Regression Trees*, Routledge, Boca Raton, FL, USA.
32. Breunig M. M., Kriegel H.-P., Ng R. T., Sander J. (2000), “LOF: Identifying density-based local outliers”, in *Proc. ACM SIGMOD Int. Conf. Management of Data*, Dallas, TX, USA, pp. 93–104, doi: 10.1145/342009.335388.
33. Bry, G., & Boschan, C. (1971). *Cyclical Analysis of Time Series: Selected Procedures and Computer Programs*. New York: National Bureau of Economic Research.
34. Burns A., Mitchell W. (1946), *Measuring Business Cycles*, National Bureau of Economic Research, Cambridge, MA, USA.

35. Casey E. (2020), “Do macroeconomic forecasters use macroeconomics to forecast?”, *International Journal of Forecasting* 36(4), pp. 1439–1453.
36. Canova F., Ciccarelli M. (2004), “Forecasting and turning point predictions in a Bayesian panel VAR model”, *Journal of Econometrics*, 120(2), pp. 327–359. [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S0304407603002161>
37. Chawla N. V., Bowyer K. W., Hall L. O., Kegelmeyer W. P. (2002), “SMOTE: Synthetic minority over-sampling technique”, *Journal of Artificial Intelligence Research*, 16, pp. 321–357, doi: 10.1613/jair.953.
38. Chen X. S., Kim M. G., Lin C., Na H. J. (2025), “Development of Per Capita GDP Forecasting Model Using Deep Learning: Incorporating CPI and Unemployment,” *Sustainability*, vol. 17, no. 3, p. 843.
39. Chen C. H., Cheng W. T. (2020), “Forecasting GDP using LSTM with Markov switching”, *Economic Modelling*, 90, pp. 88–100.
40. Chen T., Guestrin C. (2016), “XGBoost: A scalable tree boosting system,” in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Francisco, CA, USA, pp. 785–794.
41. Chung S. (2024), *Inside the black box: Neural network based real time prediction of US recessions*, arXiv:2310.17571, revised May 2024. [Online]. Available: <https://arxiv.org/abs/2310.17571> (Accessed: Jun. 10, 2025)
42. Cook T. R., Smalter Hall A. (2017), *Macroeconomic Indicator Forecasting with Deep Neural Networks*, Research Working Paper RWP 17-11, Federal Reserve Bank of Kansas City, September 29. [Online]. Available: <https://www.kansascityfed.org/documents/4065/rwp17-11cooksmalterhall.pdf>
43. Cortes C., Vapnik V., (1995), “Support-vector networks,” *Machine Learning*, vol. 20, no. 3, pp. 273–297. DOI: 10.1007/BF00994018.
44. Engle R. F. (1982), “Autoregressive Conditional Heteroscedasticity with Estimates of the Variance of United Kingdom Inflation”, *Econometrica*, 50(4), pp. 987–1007.
45. Evans G.W., Honkapohja S. (2012), “Learning and expectations in macroeconomics”, in *Learning and Expectations in Macroeconomics*, Princeton University Press, Princeton.
46. European Central Bank (2023), *ECB Economic Bulletin*, Issue 2/2023, European Central Bank, Frankfurt am Main.
47. European Commission (2023), *European Economic Forecast – Autumn 2023*, Institutional Paper 220, Publications Office of the European Union, Luxembourg. [Online]. Available: https://economy-finance.ec.europa.eu/publications/european-economic-forecast-autumn-2023_en.

48. Ewing R., Gruen D., Hawkins J. (2005), "Forecasting the macro economy", *Economic Roundup*, (2), pp. 11–25.
49. Federal Reserve Board (2015), *Measurement error in macroeconomic data and economics research* (FEDS Working Paper No. 2015-102). Board of Governors of the Federal Reserve System
50. Feldkircher M., Huber F., Schreiner J., Tirpák M., Tóth P., Wörz J. (2015), "Bridging the information gap: Small-scale nowcasting models of GDP growth for selected CEE countries", *Focus on European Economic Integration*, pp. 56–75.
51. Fischer M., Krauss C. (2018), "Deep learning with long short-term memory networks for financial market predictions", *European Journal of Operational Research*, 270(2), pp. 654–669.
52. Garboden P.M. (2020), "Sources and types of big data for macroeconomic forecasting", in *Macroeconomic Forecasting in the Era of Big Data*, Springer, pp. 3–23
53. Gadde N., Mohapatra A., Kusunam I., Uppaladinni S. (2024), "How Machine Learning Models Impacting Economic Predictions: A Review," *Research Advances in Engineering Research & Technology*, vol. 1. [Online]. Available: https://www.researchgate.net/profile/Avaneesh-Mohapatra-2/publication/384249542_How_Machine_learning_Models_Impacting_Economic_Predictions_A_Review/links/66f107df19c9496b1fb7ded7/How-Machine-learning-Models-Impacting-Economic-Predictions-A-Review.pdf
54. Garnitz J., Lehmann R., Wohlrabe K. (2019), "Forecasting GDP all over the world using leading indicators based on comprehensive survey data", *Applied Economics* 51(54), pp. 5802–5816.
55. Ghazo A. (2021), "Applying the ARIMA Model to the Process of Forecasting GDP and CPI in the Jordanian Economy", *International Journal of Financial Research*, 12(3), pp. 70–77.
56. Giannone D., Reichlin L., Small D. (2008), "Nowcasting: The real-time informational content of macroeconomic data", *Journal of Monetary Economics* 55(4), pp. 665–676.
57. Gogas P., Papadimitriou T., Matthaiou M., Chrysanthidou E. (2014), "Yield Curve and Recession Forecasting in a Machine Learning Framework", *Computational Economics*, vol. 45, pp. 635–645.
58. Goulet Coulombe P., Leroux M., Stevanovic D., Surprenant S. (2020), "How is Machine Learning Useful for Macroeconomic Forecasting?," *GCLSS Workshop*, University of Pennsylvania, first version October 2019; current version August 24, 2020. [Online]. Available: https://www.stevanovic.uqam.ca/GCLSS_MC_MacroFcst.pdf

59. Greene W. H. (2018), *Econometric Analysis*, 8th ed., Pearson, New York, NY, USA.
60. Gujarati D. N. (2004), *Basic Econometrics*, 4th ed., McGraw Hill/Irwin, New York, NY, USA.
61. Guo Y. (2023), *Three Essays in Macroeconomic Forecasting Using Dimensionality Reduction Methods*, Ph.D. thesis, University of Glasgow.
62. Gupta A. (2025), "Neural Networks for Financial Time Series Forecasting ", *GeeksforGeeks*. [Online]. Available: <https://www.geeksforgeeks.org/time-series-forecasting-with-support-vector-regression/>
63. Gürses-Tran G. (2024), *Machine Learning Techniques for Time Series Forecasting in Power Systems Operation*, RWTH Aachen University, Germany.
64. Guyon I., Elisseeff A. (2003), "An introduction to variable and feature selection", *Journal of Machine Learning Research*, 3, pp. 1157–1182. [Online]. Available: <https://www.jmlr.org/papers/volume3/guyon03a/guyon03a.pdf>
65. Hamilton J. D. (1994), *Time Series Analysis*, Princeton University Press, Princeton, NJ, USA.
66. Hamilton J. D. (1989), *A New Approach to the Economic Analysis of Nonstationary Time Series and the Business Cycle*, *Econometrica*, vol. 57, no. 2, pp. 357–384.
67. Hassan M.M. (2022), "Neural Networks for Financial Time Series Forecasting", *PubMed Central*. [Online]. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC9141105/>
68. He H., Garcia E. A. (2009), "Learning from imbalanced data", *IEEE Transactions on Knowledge and Data Engineering*, 21(9), pp. 1263–1284, doi: 10.1109/TKDE.2008.239.
69. Hochreiter S., Schmidhuber J. (1997), "Long short-term memory", *Neural Comput.*, 9, pp. 1735–1780.
70. Hodrick R. J., Prescott E. C. (1997), *Postwar U.S. Business Cycles: An Empirical Investigation*, *Journal of Money, Credit and Banking*, vol. 29, no. 1, pp. 1–16.
71. Holmes G., Pfahringer B., Kirkby R., Frank E., Hall M. (2002), "Multiclass alternating decision trees," in *Proceedings of the European Conference on Machine Learning*, pp. 161–172.
72. Hull I. (2021), *Machine Learning for Economics and Finance in TensorFlow 2: Deep Learning Models for Research and Industry*, 1st ed., Apress, Berkeley, CA. doi: 10.1007/978-1-4842-6373-0.
73. Hyndman R. J., Athanasopoulos G. (2018), *Forecasting: Principles and Practice*, 2nd ed., OTexts, Melbourne, Australia. [Online]. Available: <https://otexts.com/fpp2/>

74. Ibrahim Y. A., Nweze N. O., Adehi M. U., Chaku S. E. (2024), “Development of a hybrid macroeconomic model for forecast of economic indicators”, *Scientific World Journal*, 20(1), pp. 44–56.
75. International Monetary Fund (2012), *Gross domestic product: An economy's all*, IMF Publications, Washington, DC.
76. International Monetary Fund (2020), *World Economic Outlook Update: January 2020*, IMF, Washington, DC, USA:. [Online]. Available: <https://www.imf.org/en/Publications/WEO/Issues/2020/01/20/weo-update-january2020>
77. International Monetary Fund (2020), *World Economic Outlook Update: June 2020*, IMF, Washington, DC, USA. [Online]. Available: <https://www.imf.org/en/Publications/WEO/Issues/2020/06/24/WEOUpdateJune2020>.
78. International Monetary Fund (2023), *World Economic Outlook: Navigating Global Divergences*, IMF, Washington, DC, USA. [Online]. Available: <https://www.imf.org/en/Publications/WEO/Issues/2023/10/10/world-economic-outlook-october-2023>.
79. Izzo C. (2022), *On Explainable Deep Learning for Macroeconomic Forecasting and Finance*, Ph.D. dissertation, Department of Computer Science, University College London, London, U.K.
80. Jallow H., Mwangi R. W., Gibba A., Imboga H. (2025), “Transfer learning for predicting of gross domestic product growth based on remittance inflows using RNN-LSTM hybrid model: a case study of The Gambia”, *Frontiers in Artificial Intelligence*, 8, Art. no. 1510341, Feb.
81. Jansen W.J., de Winter J.M. (2018), “Combining model-based near-term GDP forecasts and judgmental forecasts: A real-time exercise for the G7 countries”, *Oxford Bulletin of Economics and Statistics* 80(6), pp. 1213–1242.
82. Johansen S. (1991), “Estimation and hypothesis testing of cointegration vectors in Gaussian vector autoregressive models”, *Econometrica*, 59(6), pp. 1551–1580.
83. Kim J., Kim H., Kim H. G., Lee D., Yoon S. (2025), *A Comprehensive Survey of Deep Learning for Time Series Forecasting: Architectural Diversity and Open Challenges*, *Artificial Intelligence Review*, vol. 58, no. 216, April. doi: 10.1007/s10462-025-11223-9.
84. Kim Y., Kim H. (2020), “Forecasting stock prices with a feature fusion CNN-LSTM model”, *Expert Systems with Applications*, 158, p. 113510.
85. Kidane W., Zhu W. (2013), “China-Africa investment treaties: Old rules, new challenges”, *Fordham International Law Journal* 37, pp. 1035.

86. Kontogeorgos G. (2023), *Essays in Macroeconomic Forecasting with Linear and Non-linear Time Series Models Using Bayesian Techniques*, Ph.D. thesis, University of Strathclyde.
87. Kwiatkowski D., Phillips P. C. B., Schmidt P., Shin Y. (1992), *Testing the Null Hypothesis of Stationarity Against the Alternative of a Unit Root: How Sure Are We That Economic Time Series Have a Unit Root?*, *Journal of Econometrics*, vol. 54, no. 1–3, pp. 159–178. doi: 10.1016/0304-4076(92)90104-Y.
88. Le T., Tran S., Ngo T., Bui H. D. (2024), “Predicting Vietnam’s economic growth using machine learning approaches”, *The Empirical Economics Letters*, 23(SI 2), Feb. doi: 10.5281/zenodo.10893067
89. Lepenies P. (2016), *The power of a single number: a political history of GDP*, Columbia University Press, New York.
90. Li S., Schulwolf Z. B., Miikkulainen R. (2025), “Transformer based time-series forecasting for stock”, *arXiv preprint*, arXiv:2502.09625. doi: 10.48550/arXiv.2502.09625
91. Li X., Xue L., Liang J. (2025), "Macro Financial Condition Index Construction and Forecasting Based on Machine Learning Techniques: Empirical Evidence from China", *Symmetry*, vol. 17, no. 6, art. 904. [Online]. Available: <https://www.mdpi.com/2073-8994/17/6/904>
92. Lim Z., Zohren S. (2021), *Time-Series Forecasting with Deep Learning: A Survey*, *Philosophical Transactions of the Royal Society A*, vol. 379. doi: 10.1098/rsta.2020.0209.
93. Litterman R. B. (1986), “Forecasting with Bayesian Vector Autoregressions — Five Years of Experience”, *Journal of Business & Economic Statistics*, 4(1), pp. 25–38.
94. Liu F. T., Ting K. M., Zhou Z.-H. (2008), “Isolation forest”, in *Proc. 8th IEEE Int. Conf. Data Mining (ICDM)*, Pisa, Italy, pp. 413–422, doi: 10.1109/ICDM.2008.17.
95. Lütkepohl H. (2005), *New Introduction to Multiple Time Series Analysis*, Springer, Berlin, Germany.
96. Lyhagen J., Ekberg S., Eidestet R. (2010), “Beating the war: Improving Swedish GDP forecasts using art and concert-effect corrections”, *Journal of Forecasting* 31(5), pp. 354–361.
97. Mankiw N.G. (2020), *Principles of Macroeconomics*, 8th ed., Cengage Learning, Boston, MA, USA.
98. Marcellino M., Schumacher C. (2010), “Factor midas for nowcasting and forecasting with ragged-edge data: A model comparison for German GDP”, *Oxford Bulletin of Economics and Statistics* 72(4), pp. 518–550

99. Masini R. P., Medeiros M. C., Mendes E. F. (2023), "Machine learning advances for time series forecasting", *Journal of Economic Surveys*, 37(1), pp. 76–111.
100. Mills, T. C. (2013). *A very British affair: Six Britons and the development of time series analysis during the 20th century*. Palgrave Macmillan. <https://doi.org/10.1057/9781137291264>
101. Montaña Moreno J.J., Palmer Pol A., Muñoz Gracia P. (2011), "Artificial Neural Networks Applied to Forecasting Time Series", *Psicothema*, vol. 23, no. 2, pp. 322–329. [Online]. Available: <https://www.psicothema.com/pdf/3889.pdf>
102. Nicholson W.B., Matteson D.S., Bien J. (2017), "Varx-l: Structured regularization for large vector autoregressions with exogenous variables", *International Journal of Forecasting* 33(3), pp. 627–651
103. Montgomery D.C., Peck E.A., Vining G.G. (2012), *Introduction to linear regression analysis*, Wiley, Hoboken.
104. Oancea B. (2025), "Advancing GDP Forecasting: The Potential of Machine Learning Techniques in Economic Predictions", *CoRR*, abs/2502.19807. [Online]. Available: <https://arxiv.org/abs/2502.19807>
105. OECD (2021), *Tax Administration 2021: Comparative Information on OECD and Other Advanced and Emerging Economies*, OECD Publishing. [Online]. Available: https://www.oecd.org/content/dam/oecd/en/publications/reports/2021/09/tax-administration-2021_72b221d1/cef472b9-en.pdf.
106. Okasha M.K. (2014), "Using Support Vector Machines in Financial Time Series Forecasting", *Journal of Statistics*, vol. 2, no. 1, pp. 28–39. doi: 10.5923/j.statistics.20140401.03.
107. Okasha M. K., Yassen A. A. (2013), "The Application of Artificial Neural Networks in Forecasting Economic Time Series", *Journal of Statistics and Mathematics*, (6), [Online]. Available: https://www.researchgate.net/publication/249008657_The_Application_of_Artificial_Neural_Networks_In_Forecasting_Economic_Time_Series
108. Ollivaud P., Pionnier P.-A., Rusticelli E., Schwellnus C., Koh S.-H. (2016), "Forecasting GDP during and after the great recession", *OECD Economic Policy Paper*, no. 1313, OECD Publishing, Paris
109. Paloni A., Zanardi M. (2012), *The IMF, World Bank and policy reform*, Routledge, London.
110. Paul J. (2024), "Financial Time Series Analysis with Transformer Models", December. [Online]. Available: https://www.researchgate.net/publication/387524930_Financial_Time_Series_Analysis_with_Transformer_Models

111. Perron P. (1989), *The Great Crash, the Oil Price Shock, and the Unit Root Hypothesis*, *Econometrica*, vol. 57, no. 6, pp. 1361–1401.
112. Qin Y., Song D., Cheng H., Cheng W., Jiang G., Cottrell G. W. (2017), “A dual-stage attention-based recurrent neural network for time series prediction”, in *Proc. 26th Int. Joint Conf. Artif. Intell. (IJCAI)*, pp. 2627–2633. doi: 10.24963/IJCAI.2017/366
113. Qiu, J.; Wang, B.; Zhou, C. "Forecasting stock prices with long-short term memory neural network based on attention mechanism." *PLoS ONE* 2020, 15, e0227222. DOI: 10.1371/journal.pone.0227222
114. Quinlan J. R. (1986), *Induction of Decision Trees*, *Machine Learning*, vol. 1, no. 1, pp. 81–106. doi: 10.1023/A:1022643204877.
115. Samavati A., Kung J., Rock J., Azil M. (2024), "Leveraging Deep Learning for Economic Forecasting and Decision Making: Opportunities and Challenges", *EasyChair Preprint*, no. FzQm. [Online]. Available: <https://easychair.org/publications/preprint/FzQm/open>
116. Sharma A. (2018), “ML | Feature Scaling – Part 2”, *GeeksforGeeks*, Jul. 17. [Online]. Available: <https://www.geeksforgeeks.org/ml-feature-scaling-part-2/>
117. Siami-Namini S., Tavakoli N., Namin A. S. (2018), “A comparison of ARIMA and LSTM in forecasting time series”, in *Proc. IEEE Int. Conf. Big Data*, pp. 1394–1401.
118. Sims C. A. (1980), *Macroeconomics and Reality*, *Econometrica*, vol. 48, no. 1, pp. 1–48.
119. Smets F., Wouters R. (2003), “An Estimated Dynamic Stochastic General Equilibrium Model of the Euro Area”, *Journal of the European Economic Association*, 1(5), pp. 1123–1175.
120. Stock J. H., Watson M. W. (2001), *Vector Autoregressions*, *Journal of Economic Perspectives*, vol. 15, no. 4, pp. 101–115.
121. Stock J. H., Watson M. W. (2007), “Forecasting in Economics,” in Elliott G., Granger C. W. J., Timmermann A. (Eds.), *Handbook of Economic Forecasting*, vol. 1, Elsevier, Amsterdam, Netherlands, pp. 3–59.
122. Stock J. H., Watson M. W. (2002), “Forecasting using principal components from a large number of predictors”, *Journal of the American Statistical Association*, 97(460), pp. 1167–1179.
123. Stock J. H., Watson M. W. (2003), “Forecasting output and inflation: The role of asset prices”, *Journal of Economic Literature*, 41(3), pp. 788–829. doi: 10.1257/jel.41.3.788
124. Stock, J. H., Watson, M. W. (2002), "Macroeconomic forecasting using diffusion indexes," *Journal of Business & Economic Statistics*, 20(2), 147–162. <https://doi.org/10.1198/073500102317351921>

125. Suleiman H., Nguyen M.-T. T., Mendez C. (2025), "Predicting subnational GDP in Vietnam with remote sensing data: A machine learning approach", *Letters in Spatial and Resource Sciences*, 18, Art. 5, Mar. doi: 10.1007/s12076-025-00397-z
126. Tao Y., Yan J., Niu E., Zhai P., Zhang S. (2025), "An SVM Based Anomaly Detection Method for Power System Security Analysis Using Particle Swarm Optimization and t SNE for High Dimensional Data Classification", *Processes*, vol. 13, no. 2, art. 549.
127. Thu L. H., Leon-Gonzalez R. (2021), *Forecasting macroeconomic variables in emerging economies: An application to Vietnam*, GRIPS Discussion Paper 21-03, National Graduate Institute for Policy Studies, Tokyo, Japan, Jun. [Online]. Available: https://grips.repo.nii.ac.jp/?action=repository_uri&item_id=1824&file_id=20&file_no=1
128. Tran P., Le A., Ngo T. (2024), "Application of machine learning methods in forecasting economic growth and inflation of Vietnam", *ResearchGate*, Mar. [Online]. Available: <https://www.researchgate.net/publication/4735651>
129. UNCTAD (2023), *Trade and Development Report 2023: Growth, Debt, and Climate – Realigning the Global Financial Architecture*, United Nations, Geneva. [Online]. Available: https://unctad.org/system/files/official-document/tdr2023_en.pdf.
130. United Nations, European Commission, International Monetary Fund, OECD, W European Economic Forecast (2009), *System of National Accounts 2008*, UN Publications, New York, paras. 9.113–9.114.
131. United Nations Statistics Division. (2013). *Guidelines on integrated economic statistics*. United Nations.
132. Vladošić L., Živković A., Pantić N. (2020), "Macroeconomic analysis of GDP and employment in EU countries", *Ekonomika* 66(1), pp. 65–76.
133. Vaswani A., Shazeer N., Parmar N., Uszkoreit J., Jones L., Gomez A. N., Kaiser Ł., Polosukhin I. (2017), "Attention is all you need", in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 30, pp. 5998–6008. doi: 10.5555/3295222.3295349
134. Vreeland J.R. (2003), *The IMF and economic development*, Cambridge University Press, Cambridge.
135. World Bank (2024), *Global Economic Prospects*, World Bank. [Online]. Available: <https://thedocs.worldbank.org/en/doc/661f109500bf58fa36a4a46eeace6786-0050012024/original/GEP-Jan-2024.pdf>.

136. Xie H., Xu X., Yan F., Qian X., Yang Y. (2024), “Deep learning for multi-country GDP prediction: A study of model performance and data impact”, *CoRR*, abs/2409.02551, doi: 10.48550/arXiv.2409.02551
137. Zhang J., Zhang K. (2013), "A Short Term Forecasting Model with Inhibiting Normal Distribution Noise of Sale Series", *Journal of Applied Statistics*, vol. 39, no. 8, pp. 1800–1815. [Online]. Available: <https://www.tandfonline.com/doi/full/10.1080/08839514.2013.805599>.
138. Zhao T., Chen G., Suraphee S., Phoophiwfa T., Busababodhin P. (2025), "A Hybrid TCN-XGBoost Model for Agricultural Product Market Price Forecasting", *PLOS ONE*, vol. 20, no. 5, e0322496. [Online]. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC12047798/>

Phụ lục BỘ DỮ LIỆU KTVM

Bộ dữ liệu vĩ mô của Việt Nam và một số quốc gia đại diện cho các nền kinh tế khác nhau mà luận án xây dựng gồm các chỉ tiêu kinh tế chia thành các nhóm trong các bảng từ Bảng PL.1 đến Bảng PL.7 sau:

- (1) Các chỉ tiêu nhóm đo lường quy mô nền kinh tế, lực lượng lao động và năng suất theo thời gian.

Bảng PL. 1 Các chỉ tiêu đo lường quy mô nền kinh tế, lực lượng lao động và năng suất theo thời gian

STT	Mã chỉ tiêu	Tên chỉ tiêu	Đơn vị	Diễn giải/ Ý nghĩa
1	rgdpe	GDP thực tế theo phía chi tiêu	Triệu USD (2017)	Biến này phản ánh tổng chi tiêu cuối cùng trong nền kinh tế đã điều chỉnh theo sức mua tương đương. Được tính theo PPP chuỗi
2	rgdpo	GDP thực tế theo phía sản xuất	Triệu USD (2017)	Tổng giá trị sản lượng hàng hóa và dịch vụ. Cũng tính theo PPP chuỗi
3	pop	Dân số	Triệu người	Chỉ tiêu này có vai trò quan trọng trong phân tích kinh tế dài hạn và ước tính GDP bình quân đầu người.
4	Emp	Số người có việc làm	Triệu người	Đây là một chỉ số then chốt để đánh giá thị trường lao động và mức độ tận dụng nguồn nhân lực.
5	ahv	Số giờ làm việc trung bình mỗi năm của người lao động	Giờ/năm/người	Giúp phân tích năng suất lao động dựa trên thời gian lao động
6	hc	Chỉ số vốn con người		Được tính toán từ số năm đi học và hiệu suất sinh

				lợi từ giáo dục; phản ánh chất lượng lực lượng lao động
7	une_r	Tỷ lệ thất nghiệp	%	

- (2) Nhóm đo lường tổng sản lượng và tích lũy vốn theo mức giá hiện hành đã điều chỉnh theo PPP.

Bảng PL. 2 Các chỉ tiêu đo lường sản lượng và tích lũy vốn

STT	Mã chỉ tiêu	Tên chỉ tiêu	Đơn vị	Diễn giải/ Ý nghĩa
1	ccon	Tiêu dùng thực tế	Triệu USD (2017)	Tiêu dùng thực tế của hộ gia đình và chính phủ, thể hiện tổng cầu nội địa tiêu dùng
2	cda	Hấp thụ nội địa thực tế	Triệu USD (2017)	Bao gồm tiêu dùng và đầu tư; là một chỉ số tổng hợp của cầu trong nước
3	cgdpe	GDP thực tế theo phía chi tiêu tại mức PPP hiện hành	Triệu USD (2017)	Phản ánh quy mô nền kinh tế tại giá trị hiện tại của hàng hóa dịch vụ tiêu dùng cuối cùng
4	cgdpo	GDP thực tế phía sản xuất tại mức PPP hiện hành	Triệu USD (2017)	Tính theo tổng giá trị sản xuất
5	cn	Tổng vốn vật chất của nền kinh tế,	Triệu USD (2017)	Biến này đại diện cho tổng vốn vật chất tích lũy của nền kinh tế tại một thời điểm, sau khi đã trừ khấu hao. Bao gồm các tài sản cố định phục vụ sản xuất như: máy móc, nhà xưởng, thiết bị, cơ sở hạ tầng... Nó phản ánh quy mô tài sản cố định thực tế còn lại của nền kinh tế để tạo ra sản lượng trong tương lai.

- (3) Nhóm các biến dựa trên tài khoản quốc gia

Bảng PL. 3 Các chỉ tiêu dựa trên tài khoản quốc gia

<i>STT</i>	<i>Mã chỉ tiêu</i>	<i>Tên chỉ tiêu</i>	<i>Đơn vị</i>	<i>Diễn giải/ Ý nghĩa</i>
1	rgdpna	GDP thực tế theo giá quốc gia	Triệu USD (2017)	Phản ánh tổng sản lượng không điều chỉnh theo PPP mà theo giá nội địa – thích hợp cho phân tích nội bộ từng quốc gia. Cố định năm 2017.
2	rconna	Tiêu dùng thực tế theo giá quốc gia	Triệu USD (2017)	Cố định năm 2017 – bao gồm cả hộ gia đình và chính phủ.
3	rdana	Hấp thụ trong nước thực tế theo giá cố định	Triệu USD (2017)	Tổng hợp tiêu dùng và đầu tư trong nước theo giá cố định năm 2017.
4	rnna	Tổng vốn vật chất thực tế theo giá quốc gia	Triệu USD (2017)	Cố định năm 2017
5	rkna	Dịch vụ vốn theo giá quốc gia	Triệu USD (2017)	Phản ánh năng lực sản xuất từ vốn
6	rtpna	Năng suất nhân tố tổng hợp(TFP)		Phản ánh năng suất tổng thể của các yếu tố sản xuất (vốn và lao động) trong nền kinh tế. Theo giá quốc gia cố định, được chuẩn hóa tại năm 2017 = 1.
7	rwtfpna	TFP có ý nghĩa phúc lợi		Phản ánh hiệu quả phân bổ nguồn lực ảnh hưởng đến phúc lợi, cũng được tính tại giá quốc gia cố định (2017=1).
8	labsh	Tỷ trọng thu nhập lao động trong	%	Phân tích cơ cấu phân phối thu nhập.

		GDP (giá hiện hành)		
9	irr	Tỷ suất sinh lợi nội tại thực tế của vốn đầu tư	%	Thể hiện khả năng tạo giá trị từ vốn vật chất.
10	delta	Tỷ lệ khấu hao trung bình của vốn vật chất hàng năm.	%	Đây là phần giá trị của tài sản cố định bị hao mòn hoặc mất đi do sử dụng, lỗi thời hoặc suy giảm hiệu suất theo thời gian. Tỷ lệ này phản ánh mức độ cần thiết để bù đắp đầu tư nhằm duy trì năng lực sản xuất hiện tại

(4) Nhóm chỉ số mức giá/tỷ giá điều chỉnh sức mua theo PPP

Bảng PL. 4 Các chỉ tiêu điều chỉnh sức mua

STT	Mã chỉ tiêu	Tên chỉ tiêu	Đơn vị	Diễn giải/ Ý nghĩa
1	xr	Tỷ giá hối đoái	Nội tệ/USD	Tỷ giá hối đoái giữa nội tệ và USD, bao gồm dữ liệu thị trường và ước lượng.
2	pl_con	Mức giá tiêu dùng thực tế (CCON)		Là tỷ số giữa giá tiêu dùng nội địa và giá theo PPP tương ứng; chuẩn hóa theo mức giá của Mỹ trong năm 2017 = 1.
3	pl_da	Mức giá của hấp thụ nội địa thực tế (CDA)		tương tự pl_con, chuẩn hóa theo giá của Mỹ
4	pl_gdpo	Mức giá GDP theo phía sản xuất (CGDPo)		chuẩn hóa theo giá của Mỹ (2017 = 1), dùng để điều chỉnh sức mua và so sánh quốc tế

(5) Nhóm các biến mức giá theo PPP của các cấu phần chi tiêu và vốn

Bảng PL. 5 Nhóm các chỉ tiêu mức giá theo PPP của các cấu phần chi tiêu và vốn

<i>STT</i>	<i>Mã chỉ tiêu</i>	<i>Tên chỉ tiêu</i>	<i>Đơn vị</i>	<i>Diễn giải/ Ý nghĩa</i>
1	pl_c	Mức giá tiêu dùng	%	Mức giá tiêu dùng hộ gia đình so với Hoa Kỳ (năm 2017 = 1)
2	pl_i	Mức giá đầu tư	%	Nó cho biết giá hàng hóa đầu tư trong nước cao hay thấp hơn so với tiêu chuẩn quốc tế (thường là Mỹ = 1)
3	pl_g	Mức giá chi tiêu chính phủ	%	Mức giá chi tiêu công so với Hoa Kỳ (năm 2017 = 1)
4	pl_x	Mức giá xuất khẩu	%	Mức giá hàng hóa và dịch vụ xuất khẩu so với Hoa Kỳ (năm 2017 = 1)
5	pl_m	Mức giá nhập khẩu	%	Mức giá hàng hóa và dịch vụ nhập khẩu so với Hoa Kỳ (năm 2017 = 1)
6	pl_n	Mức giá vốn vật chất	%	Mức giá của tổng tài sản cố định (vốn vật chất) so với Hoa Kỳ
7	pl_k	Mức giá dịch vụ	%	Mức giá dịch vụ sử dụng vốn (chi phí vốn) so với Hoa Kỳ

(6) Nhóm các biến tỷ trọng trong GDP theo PPP hiện hành

Bảng PL. 6 Các biến tỷ trọng trong GDP theo PPP

<i>STT</i>	<i>Mã chỉ tiêu</i>	<i>Tên chỉ tiêu</i>	<i>Đơn vị</i>	<i>Diễn giải/ Ý nghĩa</i>
1	csn_c	Tỷ trọng tiêu dùng hộ gia đình	%	Tỷ trọng tiêu dùng của hộ gia đình trong GDP tại PPP hiện hành
2	csn_i	Tỷ trọng đầu tư	%	Tỷ trọng tổng đầu tư cố định và hàng tồn kho trong GDP tại PPP hiện hành
3	csn_g	Tỷ trọng chi tiêu chính phủ	%	Tỷ trọng chi tiêu của khu vực công trong GDP tại PPP hiện hành

4	csx_x	Tỷ trọng xuất khẩu hàng hóa	%	Tỷ trọng giá trị xuất khẩu hàng hóa trong GDP tại PPP hiện hành
5	csx_m	Tỷ trọng nhập khẩu hàng hóa	%	Tỷ trọng giá trị nhập khẩu hàng hóa trong GDP tại PPP hiện hành
6	csx_r	Tỷ trọng sai số thống kê và thương mại ròng	%	Phần còn lại giữa các cấu phần GDP, bao gồm sai số thống kê và thương mại ròng

(7) Nhóm chỉ tiêu giá tiêu dùng và lạm phát

Bảng PL. 7 Các chỉ tiêu giá tiêu dùng và lạm phát

STT	Mã chỉ tiêu	Tên chỉ tiêu	Đơn vị	Diễn giải/ Ý nghĩa
1	cpi	Chỉ số giá tiêu dùng	Chỉ số (năm gốc =100)	Phản ánh mức giá trung bình của giỏ hàng hóa và dịch vụ tiêu dùng.
2	inf_r	Tỷ lệ lạm phát	%/năm	Được tính từ tốc độ thay đổi của CPI giữa hai kỳ. Là biến số đầu ra được quan tâm trong các mô hình dự báo kinh tế vĩ mô liên quan đến ổn định giá cả.