

**BỘ GIÁO DỤC
VÀ ĐÀO TẠO**

**VIỆN HÀN LÂM KHOA HỌC
VÀ CÔNG NGHỆ VIỆT NAM**

HỌC VIỆN KHOA HỌC VÀ CÔNG NGHỆ



Phạm Thành Đạt

**NGHIÊN CỨU MỘT SỐ PHƯƠNG PHÁP DỰ
BÁO NGẮN HẠN VÀ HỖ TRỢ RA QUYẾT ĐỊNH
TRONG GIÁM SÁT, ĐIỀU TRỊ HIV DỰA TRÊN
HỌC SÂU**

LUẬN ÁN TIẾN SĨ MÁY TÍNH

Hà Nội – 2026

BỘ GIÁO DỤC
VÀ ĐÀO TẠO

VIỆN HÀN LÂM KHOA HỌC
VÀ CÔNG NGHỆ VIỆT NAM

HỌC VIỆN KHOA HỌC VÀ CÔNG NGHỆ

Phạm Thành Đạt

NGHIÊN CỨU MỘT SỐ PHƯƠNG PHÁP DỰ
BÁO NGẮN HẠN VÀ HỖ TRỢ RA QUYẾT ĐỊNH
TRONG GIÁM SÁT, ĐIỀU TRỊ HIV
DỰA TRÊN HỌC SÂU

LUẬN ÁN TIẾN SĨ MÁY TÍNH

Ngành: Hệ thống thông tin

Mã số: 9 48 01 04

Xác nhận của Học viện Người hướng dẫn 1 Người hướng dẫn 2
Khoa học và Công nghệ (Ký, ghi rõ họ tên) (Ký, ghi rõ họ tên)

Hà Nội – 2026

LỜI CAM ĐOAN

Tôi xin cam đoan luận án: "Nghiên cứu một số phương pháp dự báo ngắn hạn và hỗ trợ ra quyết định trong giám sát, điều trị HIV dựa trên học sâu" là công trình nghiên cứu của chính mình dưới sự hướng dẫn khoa học của tập thể hướng dẫn. Luận án sử dụng thông tin trích dẫn từ nhiều nguồn tham khảo khác nhau và các thông tin trích dẫn được ghi rõ nguồn gốc. Các kết quả nghiên cứu của tôi được công bố chung với các tác giả khác đã được sự nhất trí của đồng tác giả khi đưa vào luận án. Các số liệu, kết quả được trình bày trong luận án là hoàn toàn trung thực và chưa từng được công bố trong bất kỳ một công trình nào khác ngoài các công trình công bố của tác giả. Luận án được hoàn thành trong thời gian tôi làm nghiên cứu sinh tại Học viện Khoa học và Công nghệ, Viện Hàn lâm Khoa học và Công nghệ Việt Nam.

Hà Nội, ngày tháng năm 2026

Tác giả luận án
(Ký và ghi rõ họ tên)

Phạm Thành Đạt

LỜI CẢM ƠN

Trong suốt quá trình học tập và thực hiện luận án, tôi đã nhận được sự quan tâm, hỗ trợ và động viên quý báu từ nhiều cá nhân và đơn vị. Đây là nguồn động lực quan trọng giúp tôi vượt qua những khó khăn và hoàn thành công trình nghiên cứu này.

Trước hết, tôi xin bày tỏ lòng biết ơn sâu sắc tới PGS.TS Nguyễn Trường Thắng và PGS.TS Nguyễn Việt Anh, những người đã tận tình định hướng nghiên cứu, góp ý chuyên môn và hỗ trợ tôi trong suốt quá trình triển khai đề tài. Sự nghiêm túc trong học thuật, tinh thần trách nhiệm và kinh nghiệm nghiên cứu của các thầy là nền tảng quan trọng giúp tôi hoàn thiện luận án.

Tôi xin trân trọng cảm ơn Học viện Khoa học và Công nghệ, Viện Hàn lâm Khoa học và Công nghệ Việt Nam, cùng các thầy cô, cán bộ tại Viện Công nghệ Thông tin đã tạo điều kiện thuận lợi về môi trường học thuật và nghiên cứu.

Tôi cũng xin gửi lời cảm ơn tới các cơ quan, đơn vị và cộng sự đã hỗ trợ cung cấp dữ liệu, trao đổi chuyên môn và tạo điều kiện để tôi thực hiện các nội dung nghiên cứu liên quan đến dự báo dịch bệnh HIV, phân tích sống sót và suy luận nhân quả.

Cuối cùng, tôi xin bày tỏ lòng biết ơn sâu sắc tới gia đình, bố mẹ, anh chị em và những người thân yêu, những người luôn là chỗ dựa tinh thần vững chắc, động viên và ủng hộ tôi trong suốt chặng đường học tập và nghiên cứu.

Xin trân trọng cảm ơn.

Hà Nội, ngày tháng năm 2026

Tác giả luận án
(Ký và ghi rõ họ tên)

Phạm Thành Đạt

MỤC LỤC

DANH MỤC CÁC TỪ VIẾT TẮT	vii
MỞ ĐẦU	1
1 Bối cảnh nghiên cứu	1
2 Các nghiên cứu hiện có và khoảng trống nghiên cứu trong dự báo và hỗ trợ ra quyết định cho HIV	2
2.1 Các nghiên cứu hiện có và các hạn chế	2
2.2 Khoảng trống nghiên cứu	5
3 Khung nghiên cứu của luận án	5
4 Mục tiêu, đối tượng và phạm vi nghiên cứu của luận án	6
4.1 Mục tiêu nghiên cứu của luận án	6
4.2 Đối tượng nghiên cứu	6
4.3 Phạm vi nghiên cứu	7
5 Phương pháp nghiên cứu	7
6 Những đóng góp chính trong luận án	8
7 Cấu trúc luận án	10
1 CƠ SỞ LÝ THUYẾT	12
1.1 Cơ sở lý thuyết cho dự báo dịch bệnh ngắn hạn	12
1.1.1 Lý thuyết chung về dự báo dịch bệnh	12
1.1.2 Các mô hình học sâu cho dự báo chuỗi thời gian dịch bệnh . .	14
1.1.3 Biểu diễn đồ thị trong phân tích dịch bệnh	15
1.1.4 Các phương pháp gộp đồ thị (Graph Pooling)	16
1.1.5 Học sâu trên cấu trúc đồ thị cho dự báo dịch bệnh	17
1.2 Cơ sở lý thuyết về phân tích sống sót	19
1.2.1 lý thuyết chung về phân tích sống sót	19
1.2.2 Các phương pháp học sâu trong phân tích sống sót	21
1.2.3 Các mô hình không gian trạng thái	23
1.3 Bất định trong dự báo y tế	24
1.4 Cơ sở lý thuyết về ước lượng nhân quả	25
1.4.1 lý thuyết chung về ước lượng nhân quả	25
1.4.2 Cấu trúc đồ thị nhân quả và cơ chế nhận dạng tác động can thiệp	27
1.4.3 Học máy nhân quả trong dữ liệu quan sát y tế	28
1.5 Các chỉ số đánh giá	30
1.5.1 Các chỉ số đánh giá trong dự báo	30
1.5.2 Các chỉ số đánh giá mô hình sống sót	30
1.5.3 Các chỉ số đánh giá trong ước lượng nhân quả	31
2 DỰ BÁO SỐ CA DƯƠNG TÍNH HIV DỰA TRÊN ĐỒ THỊ CA NHIỄM	34
2.1 Mô tả bài toán	34

2.2	Phân tích bài toán và cơ sở thiết kế mô hình	35
2.3	Biểu diễn đồ thị ca dương tính theo không gian –thời gian cấp thành phố	36
2.3.1	Khung biểu diễn và động cơ mô hình hóa.	36
2.3.2	Xây dựng cạnh dựa trên các yếu tố không gian, thời gian và dịch tễ.	36
2.3.3	Khái niệm tiềm năng lan truyền không gian–thời gian.	37
2.3.4	Nhúng đặc trưng nút và cạnh về cùng không gian biểu diễn	38
2.4	Mô hình dự báo	39
2.4.1	Tổng quan mô hình đề xuất	39
2.4.2	Mạng nơ-ron đồ thị dựa trên chú ý cho đồ thị đồ thị	41
2.4.3	Trích xuất đồ thị con theo khu vực từ đồ thị ca dương tính cấp thành phố	43
2.4.4	Gộp đồ thị (Graph pooling) cho đồ thị con theo khu vực	43
2.4.5	Cập nhật biểu diễn và tổng hợp thông tin mức đồ thị con	44
2.4.6	Tầng dự báo và đầu ra của mô hình	45
2.4.7	Lựa chọn cấu hình mô hình tối ưu	45
2.4.8	Quy trình huấn luyện mô hình	45
2.5	Thực nghiệm và kết quả	46
2.5.1	Dữ liệu và tiền xử lý dữ liệu	46
2.5.2	Chiến lược chia dữ liệu	47
2.5.3	Chỉ số đánh giá	47
2.5.4	Các mô hình so sánh	47
2.5.5	Thiết lập thực nghiệm	47
2.5.6	Kết quả thực nghiệm	48
2.5.7	Phân tích hệ số xác định (R^2) trong quá trình kiểm định	51
2.6	Thảo luận về các lựa chọn tham số thiết kế, ưu điểm và hạn chế của mô hình	53
2.7	kết luận chương	54

3 DỰ BÁO NGUY CƠ BẰNG PHÂN TÍCH SỐNG SÓT TRONG ĐIỀU

TRỊ HIV	56
3.1 Mô tả bài toán	56
3.2 Phân tích bài toán và cơ sở thiết kế mô hình	57
3.3 Phương pháp đề xuất	58
3.3.1 Tổng quan mô hình đề xuất	58
3.3.2 Mã hóa đặc trưng lâm sàng	59
3.3.3 Thành phần học tri thức quần thể dựa trên bộ nhớ	60
3.3.4 Học trạng thái bệnh tiềm ẩn bằng Mô hình hóa không gian trạng thái có chọn lọc	62
3.3.5 Thành phần dự báo đa nhiệm hỗ trợ ra quyết định lâm sàng .	64
3.3.6 Giải thích kết quả dự báo bằng SHAP (Temporal Explainability Analysis)	67

3.3.7	Chiến lược huấn luyện mô hình	68
3.4	Thực nghiệm và kết quả	68
3.4.1	Dữ liệu nghiên cứu	69
3.4.2	Tiền xử lý dữ liệu và xây dựng đặc trưng	70
3.4.3	Chiến lược chia dữ liệu	71
3.4.4	Chỉ số đánh giá	71
3.4.5	Các mô hình so sánh	71
3.4.6	Thiết lập thực nghiệm	72
3.4.7	Thực nghiệm chi tiết trên ACTG 388 (Longitudinal Cohort) .	72
3.4.7.1	Hiệu năng tổng thể trên ACTG 388	72
3.4.7.2	So sánh hiệu năng với các mô hình cơ sở	74
3.4.7.3	Kiểm định ý nghĩa thống kê bằng DeLong	75
3.4.7.4	Đặc điểm quần thể nghiên cứu theo phân tầng CD4 .	76
3.4.7.5	Phân tích loại bỏ thành phần (Ablation Study) . . .	76
3.4.7.6	Hiệu năng dự báo trong bối cảnh rủi ro cạnh tranh .	77
3.4.7.7	Phân tích SHAP theo thời gian (Temporal SHAP Analysis)	79
3.4.7.8	Xác thực lâm sàng (Clinical Validation)	80
3.4.7.9	Phân tích khả năng mở rộng (Scalability Analysis) .	82
3.4.8	Kết quả thực nghiệm trên bộ dữ liệu ACTG 175	82
3.4.8.1	So sánh hiệu năng tổng quát giữa các mô hình . . .	82
3.4.8.2	Phân tích AUC phụ thuộc thời gian	83
3.4.8.3	Phân tích sống sót Kaplan–Meier theo nhóm CD4 nền	84
3.4.8.4	Phân tích dự báo theo phác đồ điều trị trên các phân nhóm bệnh nhân	86
3.4.8.5	Phân tích hiệu chỉnh xác suất (Calibration) và ý nghĩa thực tiễn	87
3.4.8.6	Phân tích khuyến nghị điều trị (Treatment Recommendation Analysis)	88
3.4.8.7	Phân tích tầm quan trọng của các đặc trưng (Feature Importance)	90
3.4.8.8	Kiểm định ý nghĩa thống kê giữa các mô hình . . .	91
3.5	Thảo luận về các tham số của mô hình	91
3.6	Kết luận chương	92

4	ƯỚC LƯỢNG NHÂN QUẢ VÀ DỰ BÁO KẾT CỤC TRONG GIÁM SÁT VÀ ĐIỀU TRỊ HIV	94
4.1	Mô tả bài toán	94
4.2	Phân tích lý thuyết và thiết kế mô hình	96
4.3	Xây dựng đồ thị nhân quả	97
4.4	Phương pháp đề xuất	100
4.4.1	Mô hình đề xuất tổng thể	100

4.4.2	Học biểu diễn nhân quả sâu với kiến trúc TITAN	101
4.4.3	Ước lượng hiệu quả can thiệp cá thể hóa với DRLearner . . .	103
4.4.4	Điều phối thích ứng bằng học tăng cường cho ước lượng phản thực	105
4.4.5	Ước lượng bất định và thích ứng với sự thay đổi phân phối . .	107
4.4.5.1	Quy trình huấn luyện mô hình	109
4.5	Thực nghiệm và kết quả	109
4.5.1	Nghiên cứu mô phỏng (Simulation Study)	109
4.5.2	Đánh giá thực nghiệm trên các bộ dữ liệu HIV	114
4.5.2.1	Mô tả dữ liệu	114
4.5.2.2	Tiền xử lý dữ liệu	115
4.5.2.3	Các giả định trong suy luận nhân quả	115
4.5.2.4	Xây dựng và lựa chọn đồ thị nhân quả cho hai bộ dữ liệu.	117
4.5.2.5	Chiến lược chia dữ liệu	120
4.5.2.6	Chỉ số đánh giá	120
4.5.2.7	Các mô hình so sánh	120
4.5.2.8	Thiết lập thực nghiệm.	120
4.5.2.9	Các kết quả chính	121
4.5.3	Phân tích chi phí tính toán và khả năng mở rộng	124
4.6	Kết luận chương	125
	KẾT LUẬN VÀ HƯỚNG PHÁT TRIỂN	127

DANH MỤC CÁC TỪ VIẾT TẮT

Tiếng Anh	Viết tắt	Diễn giải
Human Immunodeficiency Virus	HIV	Virus gây suy giảm miễn dịch ở người
Acquired Immunodeficiency Syndrome	AIDS	Hội chứng suy giảm miễn dịch mắc phải
Antiretroviral Therapy	ART	Điều trị kháng vi-rút HIV
Antiretroviral	ARV	Thuốc kháng vi-rút HIV
Pre-Exposure Prophylaxis	PrEP	Điều trị dự phòng trước phơi nhiễm
Outpatient Clinic	OPC	Phòng khám ngoại trú điều trị HIV
Cluster of Differentiation 4	CD4	Tế bào lympho T CD4+
Artificial Intelligence	AI	Trí tuệ nhân tạo
Recurrent Neural Network	RNN	Mạng nơ-ron hồi quy
Long Short-Term Memory	LSTM	Mạng bộ nhớ dài-ngắn hạn
Gated Recurrent Unit	GRU	Đơn vị hồi quy có cổng
Convolutional Neural Network	CNN	Mạng nơ-ron tích chập
Graph Neural Network	GNN	Mạng nơ-ron đồ thị
Graph Convolutional Network	GCN	Mạng tích chập đồ thị
Graph Attention Network	GAT	Mạng chú ý trên đồ thị
Spatio-Temporal Graph Neural Network	STGNN	Mạng đồ thị không gian-thời gian
Message Passing Neural Network	MPNN	Mạng lan truyền thông điệp
Differentiable Pooling	DiffPool	Phép gộp đồ thị khả vi
Self-Attention Graph Pooling	SAGPool	Phép gộp đồ thị dựa trên tự chú ý
Top-K Pooling	TopKPooling	Phép gộp chọn Top-K nút
Self-Attention	Self-Attention	Cơ chế tự chú ý

Tiếng Anh	Viết tắt	Diễn giải
Multi-head Attention	Multi-head Attention	Cơ chế chú ý đa đầu
Selective State Space Model	SSSM	Mô hình không gian trạng thái có chọn lọc
MAMBA	MAMBA	Mô hình không gian trạng thái chọn lọc
Time-aware Transformer with Augmented Memory	TITAN	Transformer có bộ nhớ theo thời gian
Survival Analysis	SA	Phân tích sống sót
Cox Proportional Hazards	Cox PH	Mô hình nguy cơ tỷ lệ Cox
Inverse Probability of Censoring Weighting	IPCW	Trọng số nghịch xác suất kiểm duyệt
Concordance Index	C-index	Chỉ số tương hợp
Receiver Operating Characteristic – Area Under Curve	ROC-AUC	Diện tích dưới đường cong ROC
Directed Acyclic Graph	DAG	Đồ thị có hướng không chu trình
Average Treatment Effect	ATE	Hiệu ứng điều trị trung bình
Conditional Average Treatment Effect	CATE	Hiệu ứng điều trị có điều kiện
Individual Treatment Effect	ITE	Hiệu ứng điều trị cá thể
Inverse Probability Weighting	IPW	Trọng số nghịch xác suất
Double Machine Learning	DML	Học máy kép
Doubly Robust Learner	DR Learner	Bộ ước lượng nhân quả bền vững kép
Precision in Estimation of Heterogeneous Effect	PEHE	Sai số ước lượng hiệu ứng dị biệt
Mean Squared Error	MSE	Sai số bình phương trung bình
Mean Absolute Error	MAE	Sai số tuyệt đối trung bình
Root Mean Squared Error	RMSE	Căn bậc hai của MSE
Coefficient of Determination	R^2	Hệ số xác định
Expected Calibration Error	ECE	Sai số hiệu chỉnh kỳ vọng
Markov Decision Process	MDP	Quá trình ra quyết định Markov

DANH SÁCH HÌNH VẼ

1	Khung nghiên cứu của luận án.	6
1.1	Gộp nút	17
1.2	Loại bỏ nút	17
1.3	Đồ thị minh họa quan hệ giữa biến can thiệp X và kết cục Y dưới tác động của biến gây nhiễu Z , biến trung gian M và biến công cụ W . . .	28
2.1	mô hình dự báo số ca dương tính HIV mới ở cấp quận/huyện từ dữ liệu giám sát ca bệnh. Mô hình học biểu diễn ca bệnh trên đồ thị toàn thành phố bằng mạng nơ-ron đồ thị dựa trên chú ý, trích xuất các đồ thị con theo khu vực, thực hiện gộp đồ thị và tổng hợp biểu diễn mức khu vực để dự báo số ca mới.	40
2.2	Mô hình đề xuất: Loss huấn luyện và MSE trên tập kiểm định	50
2.3	MPNN-LSTM: Loss huấn luyện và MSE trên tập kiểm định	51
2.4	Mô hình đề xuất: Giá trị R^2 trên tập kiểm định	51
2.5	MPNN-LSTM: Giá trị R^2 trên tập kiểm định	51
2.6	So sánh hiệu năng dự báo tại sáu quận đại diện	52
3.1	Tổng quan mô hình đề xuất cho phân tích sống sót	59
3.2	AUC theo thời gian trên bộ dữ liệu ACTG 388. Các đường biểu diễn so sánh giữa Cox PH, DeepSurv, DeepHit, SurvTRACE, LSTM-Surv, GRU-Surv và mô hình đề xuất. Mô hình đề xuất đạt hiệu năng cao nhất và ổn định nhất theo thời gian.	78
3.3	Tầm quan trọng đặc trưng trung bình theo thời gian của mô hình đề xuất trên bộ dữ liệu ACTG 388. Các cột thể hiện mức đóng góp của từng đặc trưng trong dự báo sống sót.	80
3.4	AUC phụ thuộc thời gian trên bộ dữ liệu ACTG 175.	84
3.5	Đường cong sống sót Kaplan–Meier so sánh giữa nhóm bệnh nhân có $CD4$ ban đầu > 350 tế bào/ μL và $CD4 \leq 350$ tế bào/ μL trong bộ dữ liệu ACTG 175.	85
3.6	Kết quả dự báo theo từng phác đồ điều trị trên các phân nhóm bệnh nhân. Thanh sai số biểu thị khoảng tin cậy 95% (ACTG175).	87
3.7	Phân bố các phác đồ điều trị được mô hình khuyến nghị trên bộ dữ liệu ACTG 175.	89
3.8	So sánh kết quả điều trị giữa nhóm tuân thủ và không tuân thủ khuyến nghị.	90
4.1	Kiến trúc tổng thể của mô hình ước lượng nhân quả đề xuất.	101
4.2	Đồ thị nhân quả chuẩn (ground truth) cho dữ liệu mô phỏng	111
4.3	Đồ thị nhân quả ước lượng từ dữ liệu mô phỏng	112
4.4	Đồ thị nhân quả do chuyên gia xác định từ bộ dữ liệu 1	117

4.5	Đồ thị nhân quả suy diễn từ dữ liệu của bộ dữ liệu 1.	118
4.6	Đồ thị nhân quả do chuyên gia xác định từ bộ dữ liệu 2	119
4.7	Đồ thị nhân quả suy diễn từ dữ liệu của bộ dữ liệu 2.	119

DANH SÁCH BẢNG

2	Ký hiệu sử dụng trong mô hình dự báo ca nhiễm (Chương 3)	xvi
3	Ký hiệu sử dụng trong phân tích sống sót (Chương 4)	xvii
4	Ký hiệu sử dụng trong phân tích nhân quả (Chương 5)	xvii
5	Các cải tiến của mô hình đề xuất so sánh với các nghiên cứu hiện tại .	8
6	Các cải tiến của mô hình đề xuất so sánh với các nghiên cứu hiện tại .	9
7	Các cải tiến của mô hình đề xuất so sánh với các nghiên cứu hiện tại	10
2.1	Các tham số huấn luyện chính	48
2.2	Lựa chọn tổ hợp tối ưu giữa GNN và Graph Pooling	49
2.3	Lựa chọn tham số k cho TopKPooling	49
2.4	So sánh khi thay thế GAT bằng GCN	50
2.5	So sánh mô hình đề xuất với các phương pháp baseline	50
2.6	Kết quả mô hình khi loại bỏ từng đặc trưng dịch tễ của nút	52
2.7	Hiệu quả dự báo theo các khoảng thời gian khác nhau	53
3.1	Tóm tắt các bộ dữ liệu sử dụng trong nghiên cứu	70
3.2	Các chỉ số hiệu năng chính trên ACTG 388	73
3.3	Tác động lâm sàng của mô hình trên ACTG 388	73
3.4	Lợi thế của mô hình theo thời gian trên ACTG 388	73
3.5	So sánh hiệu năng các mô hình phân tích sống sót trên dữ liệu ACTG 388 (dữ liệu longitudinal).	75
3.6	So sánh thống kê cặp giữa các mô hình (giá trị p) trên ACTG 388. . .	76
3.7	Kết quả sống sót theo phân tầng CD4 ban đầu (ACTG 388)	76
3.8	Phân tích loại bỏ: Ảnh hưởng của các thành phần thời gian (ACTG 388)	77
3.9	Giá trị AUC theo thời gian (rủi ro cạnh tranh) trên ACTG 388	78
3.10	Hiệu năng theo từng loại biến cố (rủi ro cạnh tranh) trên ACTG 388 .	79
3.11	Tầm quan trọng đặc trưng toàn cục (trung bình theo thời gian) trên ACTG 388	79
3.12	Quỹ đạo bệnh nhân mẫu: Đáp ứng hoàn toàn (ID: 12884) trên ACTG 388	80
3.13	Mức độ tuân thủ khuyến nghị điều trị (xác thực theo thời gian) trên ACTG 388	81
3.14	So sánh kết quả giữa nhóm có quỹ đạo thuận lợi theo mô hình và nhóm tiêu chuẩn (ACTG 388)	82
3.15	Phân tích khả năng mở rộng của mô hình theo độ dài chuỗi trên ACTG 388	82
3.16	So sánh hiệu năng các mô hình phân tích sống sót trên ACTG175. . .	83
3.17	Giá trị AUC phụ thuộc thời gian trên ACTG175.	83
3.18	Kết quả Kaplan–Meier theo phân tầng CD4 nền (ACTG175).	85
3.19	Ước lượng kết quả dự báo theo phác đồ trên các phân nhóm (ACTG175).	86
3.20	Phân tích Calibration trên ACTG175.	88

3.21	Phân bố khuyến nghị phác đồ điều trị (ACTG175)	88
3.22	Kết quả điều trị theo mức độ tuân thủ khuyến nghị (ACTG175)	89
3.23	Mức độ phù hợp khuyến nghị theo nhóm bệnh nhân (ACTG175) . . .	89
3.24	Tầm quan trọng đặc trưng toàn cục cho dự báo sống sót (ACTG175) .	90
3.25	Tầm quan trọng đặc trưng theo phác đồ điều trị (ACTG175)	91
3.26	So sánh ý nghĩa thống kê giữa các mô hình	91
4.1	So sánh hiệu năng các mô hình suy luận nhân quả trên dữ liệu mô phỏng	112
4.2	Phân tích loại bỏ thành phần của mô hình trên dữ liệu mô phỏng . . .	113
4.3	So sánh hiệu năng dự báo trên dữ liệu mô phỏng	113
4.4	Phân tích phân phối hiệu quả can thiệp trên dữ liệu mô phỏng	113
4.5	Các siêu tham số tối ưu được lựa chọn bởi Optuna.	121
4.6	So sánh mô hình đề xuất với các phương pháp suy luận nhân quả hiện có	122
4.7	Hiệu năng các mô hình tổ hợp sử dụng TITAN, DRLearner, MLP và RL trên bộ dữ liệu 1 và bộ dữ liệu 2	122
4.8	So sánh ATE giữa các mô hình	123
4.9	Ablation study: Ảnh hưởng hiệu năng khi loại bỏ từng thành phần hệ thống trên hai bộ dữ liệu	124
4.10	Monte Carlo Dropout: Bất định dự đoán trên hai bộ dữ liệu	124
4.11	So sánh thời gian huấn luyện (phút)	125
4.12	Hiệu năng suy luận của các phương pháp trên bộ dữ liệu 1	125

TÓM TẮT LUẬN ÁN

Dịch HIV vẫn là một thách thức quan trọng trong y tế công cộng, do đây là một bệnh truyền nhiễm mạn tính cần quản lý lâu dài và có khả năng duy trì chuỗi lây truyền nếu không được kiểm soát hiệu quả. Điều này đòi hỏi hoạt động giám sát liên tục, quản lý điều trị dài hạn và tối ưu hóa các chiến lược can thiệp phù hợp. Trong thực tiễn, hệ thống giám sát và điều trị HIV tạo ra khối lượng lớn dữ liệu dịch tễ, lâm sàng và điều trị với đặc tính đa chiều, biến đổi theo thời gian và chứa nhiều mối quan hệ phụ thuộc phức tạp giữa các yếu tố nguy cơ, tiến triển bệnh và đáp ứng điều trị. Điều này đặt ra nhu cầu phát triển các phương pháp phân tích tiên tiến nhằm hỗ trợ ra quyết định ở nhiều cấp độ khác nhau.

Xuất phát từ yêu cầu đó, luận án tập trung nghiên cứu và phát triển các phương pháp học sâu nhằm hỗ trợ dự báo và ra quyết định trong giám sát và điều trị HIV thông qua ba hướng nghiên cứu. Ở cấp độ quần thể, luận án đề xuất một mô hình dự báo ngắn hạn dựa trên đồ thị không gian–thời gian ở mức từng ca bệnh, cho phép khai thác trực tiếp các mối liên hệ dịch tễ học tiềm năng giữa các trường hợp HIV. Khác với các phương pháp dự báo cục bộ truyền thống, mô hình kết hợp thông tin lan truyền ở quy mô toàn thành phố với đặc trưng của từng khu vực nhằm cải thiện độ chính xác dự báo ngắn hạn cho từng quận/huyện. Ở cấp độ cá thể, luận án phát triển một mô hình phân tích sống sót dựa trên mô hình không gian trạng thái có chọn lọc nhằm hỗ trợ tiên lượng cá thể hóa các biến cố điều trị quan trọng. Phương pháp này cho phép mô hình hóa tường minh các động lực tiến triển bệnh tiềm ẩn theo thời gian, qua đó cải thiện khả năng nắm bắt các phụ thuộc dài hạn trong dữ liệu điều trị HIV. Ở cấp độ hỗ trợ quyết định can thiệp, luận án đề xuất một khung ước lượng hiệu quả can thiệp kết hợp học biểu diễn sâu, suy luận nhân quả và cơ chế điều phối thích ứng nhằm nâng cao độ chính xác và tính ổn định trong đánh giá hiệu quả can thiệp từ dữ liệu quan sát. Cách tiếp cận này cho phép tự động điều chỉnh chiến lược ước lượng theo đặc điểm riêng của từng mẫu dữ liệu nhằm nâng cao độ chính xác và tính ổn định trong đánh giá hiệu quả can thiệp từ dữ liệu quan sát.

Các phương pháp đề xuất được đánh giá trên các bộ dữ liệu HIV thực tế, cho thấy hiệu năng cạnh tranh và tính nhất quán trên các bài toán dự báo dịch tễ, phân tích sống sót và suy luận nhân quả. Kết quả nghiên cứu khẳng định tiềm năng ứng dụng của các phương pháp học sâu trong hỗ trợ ra quyết định dựa trên dữ liệu, góp phần nâng cao hiệu quả giám sát dịch tễ, quản lý điều trị và hoạch định chiến lược phòng, chống HIV.

ABTRACT

HIV remains a significant public health challenge, as it is a chronic infectious disease requiring long-term management and can sustain ongoing transmission if not effectively controlled. This necessitates continuous surveillance, long-term treatment management, and the optimization of appropriate intervention strategies. In practice, HIV surveillance and treatment systems generate large volumes of epidemiological, clinical, and treatment-related data characterized by multidimensional, temporal dynamics, and complex dependencies among risk factors, disease progression, and treatment response. These challenges create a strong need for advanced analytical methods to support decision-making at multiple levels..

Motivated by this need, this dissertation focuses on the development of deep learning-based approaches to support forecasting and decision-making in HIV surveillance and treatment through three complementary research directions. At the population level, the dissertation proposes a short-term forecasting model based on a case-level spatio-temporal graph representation, enabling the direct modeling of potential epidemiological relationships among HIV cases. Unlike conventional localized forecasting approaches, the proposed model integrates citywide transmission patterns with district-specific characteristics to improve short-term forecasting accuracy at the district level. At the individual level, the dissertation develops a survival analysis framework based on a selective state-space model to support personalized prognosis of critical treatment-related outcomes. The proposed approach explicitly models latent disease progression dynamics over time, thereby improving the ability to capture long-term dependencies in longitudinal HIV treatment data. At the intervention decision-support level, the dissertation proposes an intervention effect estimation framework that combines deep representation learning, causal inference, and adaptive coordination mechanisms to improve the accuracy and robustness of intervention effect estimation from observational data. This approach dynamically adjusts estimation strategies according to the characteristics of individual data instances, thereby enhancing reliability in decision support for intervention planning.

The proposed methods are evaluated using real-world HIV datasets and demonstrate competitive and consistent performance across epidemiological forecasting, survival analysis, and causal inference tasks. The findings highlight the potential of deep learning-based approaches for data-driven decision support, contributing to improved epidemiological surveillance, treatment management, and strategic HIV prevention and control planning.

Các ký hiệu toán học chính trong luận án

Ký hiệu sử dụng trong mô hình dự báo ca nhiễm (Chương 3)

G	Đồ thị HIV toàn đô thị
V	Tập các nút trong G
E	Tập các cạnh trong G
X	Đặc trưng nút của đồ thị G
G_r	Đồ thị con tương ứng với khu vực r
V_r	Tập các nút trong G_r
E_r	Tập các cạnh trong G_r
X_r	Đặc trưng nút của G_r
n	Tổng số khu vực trong đô thị
L	Tổng số ca nhiễm HIV trong đô thị
L_r	Số ca nhiễm HIV tại khu vực r
x_i	Vector đặc trưng của nút i
f_{ij}	Vector đặc trưng của cạnh giữa nút i và j
h_i	Biểu diễn embedding của nút i
ρ_{ij}	Mức độ tương quan không gian giữa hai ca bệnh i và j
τ_{ij}	Mức độ tương quan thời gian giữa hai ca bệnh i và j
ξ_{ij}	Độ tương đồng dịch tễ giữa hai ca bệnh i và j
κ_{ij}	Điểm tương quan tổng hợp giữa hai ca bệnh i và j
β	Ngưỡng tạo cạnh
CT_i	Cơ sở xét nghiệm khẳng định của ca bệnh i
lon_i	Kinh độ của CT_i
lat_i	Vĩ độ của CT_i

Ký hiệu sử dụng trong phân tích sống sót (Chương 4)

t	Biến thời gian
T	Biến ngẫu nhiên biểu diễn thời gian xảy ra biến cố
δ	Biến chỉ báo biến cố (1 nếu xảy ra biến cố, 0 nếu bị kiểm duyệt)
x_i	Vector đặc trưng của mẫu thứ i
h_i	Biểu diễn tiềm ẩn của mẫu i
z_i	Trạng thái ẩn được học từ mô hình Mamba
$\lambda(t)$	Hàm nguy cơ tại thời điểm t
$\lambda(t \mid x_i)$	Hàm nguy cơ có điều kiện theo đặc trưng x_i
$S(t)$	Hàm sống sót
$S(t \mid x_i)$	Hàm sống sót có điều kiện
$H(t)$	Hàm nguy cơ tích lũy
$p(T_i \mid x_i)$	Phân phối dự đoán thời gian xảy ra biến cố
$\hat{T}_i$	Thời gian xảy ra biến cố được dự đoán
θ	Tập tham số của mô hình
$p(\theta)$	Phân phối tiên nghiệm của tham số
$\mathcal{L}$	Hàm mất mát

Ký hiệu sử dụng trong phân tích nhân quả (Chương 5)

x_i	Vector đặc trưng của mẫu thứ i
t_i	Biến điều trị của mẫu i
y_i	Kết quả quan sát được của mẫu i
$Y_i(1)$	Kết quả tiềm năng khi có can thiệp
$Y_i(0)$	Kết quả tiềm năng khi không có can thiệp
τ_i	Hiệu ứng điều trị cá thể (ITE), $\tau_i = Y_i(1) - Y_i(0)$
τ	Hiệu ứng điều trị trung bình (ATE)
$\hat{\tau}_i$	Giá trị ước lượng của hiệu ứng điều trị cá thể
$\mu_1(x_i)$	Hàm dự đoán kết quả trong nhóm được điều trị
$\mu_0(x_i)$	Hàm dự đoán kết quả trong nhóm không điều trị
$e(x_i)$	Xác suất được điều trị (propensity score), $P(t_i = 1 \mid x_i)$
$\hat{\mu}_1(x_i)$	Giá trị ước lượng của kết quả khi có điều trị
$\hat{\mu}_0(x_i)$	Giá trị ước lượng của kết quả khi không điều trị
$\hat{e}(x_i)$	Giá trị ước lượng của propensity score
z_i	Biểu diễn ẩn học được từ mô hình TITAN
$\hat{\tau}_i^{(T)}$	Hiệu ứng điều trị cá thể ước lượng bởi TITAN
$\hat{\tau}_i^{(DR)}$	Hiệu ứng điều trị cá thể ước lượng bởi DR Learner
π_i	Trọng số tổ hợp được học bởi cơ chế học tăng cường
θ	Tập tham số của mô hình
$\mathcal{L}$	Hàm mất mát

MỞ ĐẦU

1. Bối cảnh nghiên cứu

HIV/AIDS vẫn là một trong những thách thức lớn đối với y tế công cộng trên toàn cầu. Theo UNAIDS, đến cuối năm 2023 có khoảng 39,9 triệu người đang sống chung với HIV, khoảng 1,3 triệu ca nhiễm mới và 630.000 ca tử vong liên quan đến AIDS [1, 2, 3]. Tại Việt Nam, Quyết định số 1246/QĐ-TTg của Thủ tướng Chính phủ về phê duyệt Chiến lược Quốc gia chấm dứt dịch bệnh AIDS vào năm 2030 [4] đặt mục tiêu đạt các chỉ tiêu 95-95-95, gồm 95% người nhiễm HIV biết tình trạng nhiễm của mình, 95% người được chẩn đoán được điều trị bằng thuốc kháng vi rút (ART) và 95% người đang điều trị đạt ức chế tải lượng vi rút. Các mục tiêu này phản ánh toàn bộ chuỗi quản lý HIV từ phát hiện, điều trị đến kiểm soát hiệu quả điều trị.

Trong bối cảnh đó, công tác quản lý HIV đòi hỏi sự gắn kết chặt chẽ giữa giám sát dịch tễ và điều trị. Khác với nhiều bệnh truyền nhiễm khác, quản lý HIV không kết thúc ở việc phát hiện ca bệnh mà tiếp tục xuyên suốt quá trình kết nối điều trị, duy trì điều trị lâu dài, theo dõi hiệu quả điều trị và đánh giá các chương trình can thiệp [5]. Quá trình này tạo ra lượng lớn dữ liệu từ nhiều nguồn khác nhau như dữ liệu giám sát dịch tễ, hồ sơ điều trị, kết quả xét nghiệm và dữ liệu theo dõi lâm sàng. Sự gia tăng về quy mô, tính đa dạng và tính liên tục của dữ liệu đặt ra yêu cầu khai thác hiệu quả nguồn dữ liệu nhằm cung cấp cơ sở khoa học cho các quyết định trong quản lý HIV.

Trong suốt quá trình quản lý HIV, các quyết định chuyên môn chủ yếu xoay quanh ba nội dung cơ bản: dự báo diễn biến của dịch để hỗ trợ giám sát, dự báo tiên lượng của người bệnh để hỗ trợ theo dõi và điều trị, và lựa chọn các can thiệp phù hợp nhằm nâng cao hiệu quả quản lý và điều trị. Ba nội dung này phản ánh ba cấp độ phân tích trong quản lý HIV, bao gồm cấp độ quần thể (dự báo diễn biến dịch), cấp độ người bệnh (dự báo tiên lượng) và cấp độ can thiệp (hỗ trợ lựa chọn quyết định). Từ đó hình thành ba nhóm bài toán phân tích dữ liệu trọng tâm trong quản lý HIV. Các tiến bộ của các mô hình học sâu đã tạo điều kiện cho việc xây dựng các mô hình có khả năng khai thác hiệu quả nguồn dữ liệu lớn, hỗ trợ giải quyết các bài toán này với độ chính xác và khả năng tổng quát hóa ngày càng cao.

Xuất phát từ yêu cầu quản lý dịch HIV/AIDS, luận án đề xuất ba mô hình học sâu tương ứng với ba nhóm bài toán trên. Mô hình thứ nhất thực hiện dự báo ngắn hạn diễn biến dịch HIV nhằm hỗ trợ cảnh báo sớm và chuẩn bị nguồn lực ứng phó. Mô hình thứ hai dự báo nguy cơ và kết cục điều trị của người bệnh trong một khoảng thời gian tiên lượng xác định trước dựa trên dữ liệu điều trị, cung cấp thông tin hỗ trợ đánh giá nguy cơ và điều chỉnh phác đồ điều trị trong các chu kỳ điều trị kế tiếp của người bệnh. Mô hình thứ ba hỗ trợ ra quyết định thông qua ước lượng hiệu quả của các can thiệp trong giám sát và điều trị HIV đối với từng người bệnh, cung cấp cơ sở định lượng giúp lựa chọn hoặc điều chỉnh kịp thời phác đồ điều trị, kế hoạch theo dõi và chỉ định xét nghiệm phù hợp.

2. Các nghiên cứu hiện có và khoảng trống nghiên cứu trong dự báo và hỗ trợ ra quyết định cho HIV

2.1. Các nghiên cứu hiện có và các hạn chế

Các nghiên cứu về dự báo xu hướng dịch và hạn chế

Dự báo xu hướng dịch là một nội dung quan trọng trong giám sát dịch tễ, giúp phát hiện sớm nguy cơ bùng phát, dự báo nhu cầu nguồn lực và hỗ trợ xây dựng các biện pháp can thiệp phù hợp. Trong nghiên cứu HIV, các phương pháp dự báo đã phát triển theo nhiều hướng tiếp cận khác nhau.

Hướng tiếp cận đầu tiên là các mô hình thống kê truyền thống, tiêu biểu như ARIMA [6, 7]. Các mô hình này mô tả xu hướng biến động của số ca nhiễm theo thời gian và phù hợp với các chuỗi số liệu có quy luật tương đối ổn định. Tuy nhiên, chúng thường giả định mối quan hệ tuyến tính và khó phản ánh ảnh hưởng đồng thời của nhiều yếu tố dịch tễ.

Tiếp theo là các mô hình học sâu trên chuỗi thời gian, chẳng hạn như LSTM hoặc GRU [8, 9]. Nhóm mô hình này có khả năng học các quan hệ phi tuyến và phụ thuộc dài hạn trong dữ liệu, nhờ đó cải thiện độ chính xác so với nhiều phương pháp thống kê. Tuy nhiên, phần lớn các nghiên cứu vẫn xem dữ liệu như một chuỗi thời gian đơn thuần và chưa khai thác được mối liên hệ giữa các khu vực hoặc giữa các ca nhiễm.

Gần đây, các mô hình không gian–thời gian dựa trên đồ thị được phát triển nhằm mô hình hóa đồng thời quan hệ địa lý và sự thay đổi theo thời gian [10, 11, 12, 13]. Trong các mô hình này, mỗi nút thường biểu diễn một khu vực hoặc một đơn vị hành chính, còn các cạnh thể hiện mối liên hệ giữa các khu vực. Cách tiếp cận này phù hợp với nhiều bài toán dự báo dịch bệnh và đã cho kết quả khả quan.

Gần đây, cùng với sự phát triển của học sâu trên dữ liệu đồ thị, một số nghiên cứu đã mở rộng bài toán dự báo HIV theo hướng khai thác đồng thời mối quan hệ không gian giữa các khu vực và sự phụ thuộc theo thời gian của diễn biến dịch. Các mô hình này sử dụng cấu trúc đồ thị để biểu diễn mối liên kết giữa các địa phương, kết hợp với cơ chế học biểu diễn theo thời gian nhằm nâng cao khả năng dự báo diễn biến dịch trên phạm vi nhiều vùng hoặc nhiều đơn vị hành chính [14, 15, 16]

Mặc dù vậy, đối với HIV, các hướng tiếp cận trên vẫn còn một số hạn chế. Phần lớn nghiên cứu sử dụng dữ liệu tổng hợp theo khu vực hoặc theo thời gian, trong khi hệ thống giám sát HIV được xây dựng từ các bản ghi ở mức từng ca bệnh. Vì vậy, các quan hệ dịch tễ giữa các ca nhiễm thường chưa được khai thác trực tiếp. Ngoài ra, việc kết hợp đồng thời thông tin dịch tễ học, quan hệ không gian–thời gian và đặc điểm lan truyền của dịch trong cùng một mô hình vẫn còn là một vấn đề chưa được giải quyết đầy đủ. Đây cũng là khoảng trống mà luận án tập trung nghiên cứu.

Các nghiên cứu về phân tích sống sót và hạn chế

Phân tích sống sót (survival analysis) là một hướng nghiên cứu quan trọng trong thống kê y sinh, được sử dụng để mô hình hóa thời gian xảy ra một biến cố trong điều kiện dữ liệu có kiểm duyệt (censoring). Trong lĩnh vực HIV, phương pháp này được áp dụng rộng rãi để dự báo thời gian xảy ra các kết cục như tử vong, thất bại điều trị hoặc gián đoạn điều trị, từ đó hỗ trợ đánh giá tiên lượng và quản lý bệnh nhân.

Dự báo sống sót của người sống chung với HIV là một bài toán quan trọng trong giám sát và điều trị HIV/AIDS. Nhiều nghiên cứu đã sử dụng dữ liệu theo dõi dọc như số lượng tế bào CD4 và tải lượng virus để tiên lượng nguy cơ tử vong hoặc thời gian sống còn của người bệnh thông qua các mô hình phân tích sống còn và mô hình kết hợp giữa dữ liệu theo dõi dọc và dữ liệu thời gian xảy ra biến cố. Các nghiên cứu này cho thấy việc khai thác thông tin theo thời gian từ các chỉ số lâm sàng giúp cải thiện khả năng dự báo và hỗ trợ ra quyết định điều trị. [17, 18, 19]

Các thuật toán về phân tích sống sót có thể được chia thành ba hướng tiếp cận chính.

Hướng thứ nhất là các mô hình thống kê truyền thống, tiêu biểu là Kaplan–Meier và Cox Proportional Hazards [20, 21]. Đây là những phương pháp nền tảng trong phân tích sống sót và vẫn được sử dụng rộng rãi trong nghiên cứu lâm sàng. Kaplan–Meier cho phép ước lượng xác suất sống sót theo thời gian, trong khi mô hình Cox đánh giá ảnh hưởng của các yếu tố nguy cơ đến thời gian xảy ra biến cố. Tuy nhiên, các mô hình này thường dựa trên những giả định thống kê, đặc biệt là giả định nguy cơ tỷ lệ, nên khả năng mô hình hóa các quan hệ phi tuyến và dữ liệu có sự thay đổi theo thời gian còn hạn chế.

Hướng thứ hai là các mô hình phân tích sống sót dựa trên học sâu. Những mô hình như DeepSurv [22] và DeepHit [23] sử dụng mạng nơ-ron để học các quan hệ phi tuyến giữa đặc trưng của bệnh nhân và nguy cơ xảy ra biến cố. So với các phương pháp truyền thống, nhóm mô hình này có khả năng khai thác dữ liệu có cấu trúc phức tạp và thường đạt hiệu năng dự báo cao hơn khi số lượng đặc trưng lớn.

Gần đây, các mô hình dựa trên cơ chế chú ý (attention) tiếp tục mở rộng khả năng của phân tích sống sót. SurvTRACE [24] sử dụng Transformer để học các phụ thuộc dài hạn trong dữ liệu. Trong lĩnh vực HIV, một số nghiên cứu như ViT-Survive [25] và Group-Wise Self-Attention (GWSA) [26] cũng khai thác cơ chế chú ý nhằm học biểu diễn từ dữ liệu theo dõi điều trị và cải thiện khả năng tiên lượng.

Mặc dù đạt được nhiều kết quả tích cực, các hướng tiếp cận hiện nay vẫn còn một số hạn chế. Các mô hình thống kê truyền thống gặp khó khăn khi xử lý dữ liệu phi tuyến và dữ liệu theo dõi dài hạn. Trong khi đó, các mô hình học sâu và Transformer chủ yếu tập trung vào việc học biểu diễn từ dữ liệu quan sát, nhưng chưa mô hình hóa rõ ràng trạng thái bệnh tiềm ẩn và quá trình tiến triển liên tục của bệnh theo thời gian. Bên cạnh đó, phần lớn các nghiên cứu vẫn chưa quan tâm đầy đủ đến việc lượng hóa mức độ bất định của dự báo, mặc dù đây là yếu tố có ý nghĩa quan trọng trong hỗ trợ quyết định lâm sàng.

Trong thực tế, các chỉ số như tải lượng virus, số lượng tế bào CD4 hoặc thông tin điều trị chỉ phản ánh tình trạng của bệnh nhân tại từng thời điểm quan sát. Nguy cơ xảy ra biến cố còn phụ thuộc vào quá trình tiến triển của bệnh giữa các lần thăm khám, tức là những trạng thái không thể quan sát trực tiếp. Vì vậy, cần có các mô hình có khả năng biểu diễn trạng thái bệnh tiềm ẩn, mô tả động lực tiến triển theo thời gian và đồng thời cung cấp các dự báo có độ tin cậy cao hơn. Đây cũng là khoảng trống mà luận án tập trung giải quyết.

Các nghiên cứu về hỗ trợ ra quyết định can thiệp và hạn chế

Bên cạnh giám sát dịch tễ và tiên lượng điều trị, đánh giá hiệu quả của các biện pháp can thiệp cũng là một nội dung quan trọng trong quản lý HIV. Trong thực tế, nhiều quyết định như lựa chọn phác đồ điều trị, xác định thời điểm khởi trị hoặc đánh giá hiệu quả của các chương trình can thiệp phải được thực hiện trên dữ liệu quan sát thay vì các thử nghiệm ngẫu nhiên có đối chứng. Vì vậy, suy luận nhân quả đã trở thành một hướng nghiên cứu quan trọng nhằm ước lượng tác động của các biện pháp can thiệp từ dữ liệu thực tế.

Các nghiên cứu về suy luận nhân quả có thể chia thành ba nhóm chính.

Nhóm thứ nhất là các phương pháp thống kê truyền thống, như Marginal Structural Models và Inverse Probability Weighting (IPW) [27, 28]. Các phương pháp này được xây dựng trên cơ sở lý thuyết nhân quả và được sử dụng rộng rãi trong nghiên cứu dịch tễ học cũng như y học lâm sàng. Tuy nhiên, hiệu quả của chúng phụ thuộc nhiều vào các giả định về mô hình và chất lượng của các biến gây nhiễu được quan sát.

Nhóm thứ hai là các phương pháp học máy nhân quả, tiêu biểu như TARNet [29], DragonNet [30] và DR Learner [31]. Các phương pháp này kết hợp học biểu diễn với các kỹ thuật ước lượng hiệu ứng điều trị, giúp mô hình hóa tốt hơn các quan hệ phi tuyến và hỗ trợ ước lượng hiệu ứng điều trị ở mức cá thể.

Gần đây, các mô hình học sâu nhân quả tiếp tục được phát triển nhằm khai thác các biểu diễn dữ liệu phức tạp và nâng cao độ chính xác của ước lượng. Các mô hình này sử dụng mạng nơ-ron để học đặc trưng từ dữ liệu quan sát trước khi thực hiện suy luận nhân quả, qua đó cải thiện khả năng xử lý dữ liệu có số chiều lớn hoặc có cấu trúc phức tạp.

Hỗ trợ lựa chọn và đánh giá hiệu quả can thiệp điều trị là một bài toán quan trọng trong quản lý người sống chung với HIV. Do nhiều câu hỏi lâm sàng không thể đánh giá bằng thử nghiệm ngẫu nhiên có đối chứng, các nghiên cứu đã sử dụng dữ liệu quan sát kết hợp với các phương pháp suy luận nhân quả để ước lượng hiệu quả của các chiến lược điều trị. Các ứng dụng tiêu biểu bao gồm đánh giá ảnh hưởng của điều trị Zidovudine đối với thời gian sống còn, ước lượng hiệu quả của liệu pháp kháng vi-rút (HAART) từ dữ liệu theo dõi dọc và xác định thời điểm tối ưu bắt đầu điều trị ART. Những nghiên cứu này cho thấy suy luận nhân quả là một công cụ quan trọng trong hỗ trợ ra quyết định điều trị HIV từ dữ liệu quan sát [32, 33, 34].

Mặc dù đạt được nhiều kết quả tích cực, các phương pháp hiện nay vẫn còn một số hạn chế khi áp dụng cho dữ liệu HIV. Dữ liệu điều trị HIV thường có mức độ không đồng nhất cao giữa các bệnh nhân, trong khi hiệu quả của cùng một biện pháp can thiệp có thể khác nhau giữa các nhóm đối tượng. Nhiều mô hình còn phụ thuộc vào chất lượng của biểu diễn đặc trưng hoặc sử dụng một cơ chế kết hợp cố định giữa học biểu diễn và ước lượng nhân quả, nên khả năng thích ứng với từng trường hợp cụ thể còn hạn chế. Ngoài ra, sự mất cân bằng giữa các nhóm điều trị trong dữ liệu quan sát cũng ảnh hưởng đến độ ổn định của kết quả ước lượng. Những vấn đề này cho thấy vẫn cần phát triển các phương pháp có khả năng xử lý tốt hơn tính không đồng nhất của dữ liệu và nâng cao độ tin cậy của các ước lượng nhân quả trong điều kiện dữ liệu

thực tế.

2.2. Khoảng trống nghiên cứu

Từ các nghiên cứu đã công bố có thể thấy rằng, ba bài toán dự báo dịch tễ, phân tích sống sót và suy luận nhân quả đã đạt được nhiều kết quả quan trọng. Tuy nhiên, mỗi bài toán vẫn còn những hạn chế cần tiếp tục nghiên cứu.

Đối với bài toán dự báo dịch tễ, phần lớn các nghiên cứu vẫn dựa trên dữ liệu tổng hợp theo thời gian hoặc theo khu vực. Việc khai thác trực tiếp các mối quan hệ dịch tễ giữa các ca bệnh và sự lan truyền trong không gian–thời gian còn chưa đầy đủ, làm hạn chế khả năng phản ánh đặc điểm lây truyền của HIV.

Đối với bài toán phân tích sống sót, mặc dù các mô hình học sâu đã cải thiện đáng kể khả năng dự báo, việc mô hình hóa trạng thái bệnh tiềm ẩn, quá trình tiến triển liên tục theo thời gian và lượng hóa bất định trong dự báo vẫn chưa được giải quyết đầy đủ.

Đối với bài toán suy luận nhân quả, các phương pháp hiện có đã cải thiện khả năng ước lượng hiệu ứng điều trị từ dữ liệu quan sát. Tuy nhiên, việc xử lý dữ liệu không đồng nhất, sự mất cân bằng giữa các nhóm điều trị và cơ chế kết hợp linh hoạt giữa học biểu diễn với ước lượng nhân quả vẫn còn là những vấn đề cần tiếp tục nghiên cứu.

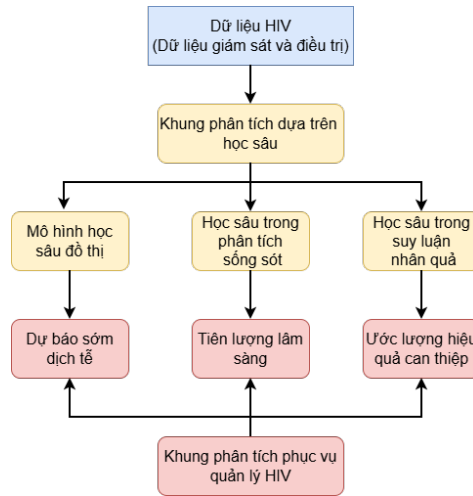
Từ các khoảng trống trên, luận án tập trung phát triển các mô hình học sâu cho ba bài toán gồm dự báo ngắn hạn trong giám sát dịch tễ HIV, phân tích sống sót trong điều trị và suy luận nhân quả hỗ trợ đánh giá can thiệp. Mỗi mô hình được xây dựng nhằm giải quyết một số hạn chế của các nghiên cứu hiện có và nâng cao khả năng hỗ trợ ra quyết định trong quản lý HIV.

3. Khung nghiên cứu của luận án

Luận án này được xây dựng xoay quanh mục tiêu phát triển các khung phương pháp học sâu nhằm phục vụ dự báo ngắn hạn và hỗ trợ ra quyết định trong hệ thống giám sát và điều trị HIV. Luận án đặt nghiên cứu trong toàn bộ chuỗi quản lý bệnh HIV, một bệnh mạn tính có đặc trưng tiến triển dài hạn và yêu cầu theo dõi, điều chỉnh can thiệp liên tục. Trong thực tế, quản lý HIV bao gồm nhiều hoạt động liên kết với nhau như giám sát dịch tễ, phát hiện sớm thông qua xét nghiệm, khởi đầu và duy trì điều trị, cũng như theo dõi đáp ứng điều trị theo thời gian. Do đó, các hệ thống phân tích dữ liệu trong bối cảnh HIV cần khai thác cả dữ liệu giám sát dịch tễ và dữ liệu điều trị, nhằm cung cấp cơ sở định lượng cho các quyết định ở nhiều cấp độ của hệ thống y tế.

Trên cơ sở đó, luận án triển khai ba hướng nghiên cứu hỗ trợ cho hoạt động giám sát, điều trị HIV bằng học sâu. Ở cấp quần thể, nghiên cứu thứ nhất tập trung vào dự báo xu hướng dịch HIV dựa trên dữ liệu giám sát nhằm hỗ trợ phát hiện sớm nguy cơ và tăng cường giám sát dịch. Ở cấp cá thể, nghiên cứu thứ hai áp dụng các phương pháp phân tích sống sót dựa trên học sâu để mô hình hóa tiến triển điều trị và tiên lượng nguy cơ gián đoạn điều trị của bệnh nhân. Cuối cùng, ở cấp quyết định, nghiên cứu thứ ba sử dụng các phương pháp suy luận nhân quả nhằm dự đoán tác động của

các chiến lược xét nghiệm và điều trị đối với kết quả điều trị trong điều kiện dữ liệu quan sát. hình 1 minh họa khung nghiên cứu của luận án.



Hình 1: Khung nghiên cứu của luận án.

Ba hướng nghiên cứu này tạo thành một khung phân tích đa cấp bao gồm dự báo (prediction), tiên lượng (prognosis) và can thiệp (intervention), góp phần hỗ trợ quản lý và ra quyết định trong hệ thống giám sát và điều trị HIV.

4. Mục tiêu, đối tượng và phạm vi nghiên cứu của luận án

4.1. Mục tiêu nghiên cứu của luận án

Luận án tập trung phát triển các mô hình học sâu hiệu quả trong dự báo ngắn hạn và ra quyết định điều trị HIV, gồm ba mục tiêu nghiên cứu chính như sau:

1. Mục tiêu 1: Nghiên cứu và đề xuất mô hình học sâu cho bài toán dự báo số ca nhiễm HIV mới trong ngắn hạn dựa trên dữ liệu giám sát theo từng ca nhiễm. Mô hình cần khai thác các mối liên hệ dịch tễ học theo không gian và thời gian nhằm phản ánh đặc điểm lan truyền của dịch HIV trong bối cảnh đô thị và nâng cao độ chính xác của dự báo.

2. Mục tiêu 2: Nghiên cứu và cải tiến phương pháp phân tích sống sót dựa trên học sâu nhằm dự đoán thời gian xảy ra các biến cố y tế như thất bại điều trị hoặc tử vong. Mô hình hướng tới khả năng nắm bắt các mối quan hệ tiềm ẩn và phụ thuộc dài hạn trong dữ liệu theo thời gian, qua đó nâng cao độ chính xác và khả năng ứng dụng trong quản lý điều trị HIV.

3. Mục tiêu 3: Nghiên cứu và cải tiến các phương pháp dự báo kết cục dưới các kịch bản can thiệp dựa trên dữ liệu quan sát, thông qua các kỹ thuật suy luận nhân quả, nhằm đánh giá hiệu quả của các can thiệp xét nghiệm và điều trị HIV, từ đó hỗ trợ ra quyết định trong các hệ thống quản lý và điều trị HIV.

4.2. Đối tượng nghiên cứu

Đối tượng nghiên cứu của luận án bao gồm các phương pháp và mô hình phân tích dữ liệu phục vụ dự báo và hỗ trợ ra quyết định trong quản lý và điều trị HIV, cụ thể như sau:

1. Các mô hình học sâu và các kỹ thuật xử lý dữ liệu liên quan tới dự báo sớm, phân tích sống sót và ước lượng nhân quả trong điều trị HIV.
2. Các kỹ thuật xây dựng đồ thị lây nhiễm dịch bệnh và đồ thị nhân quả trên dữ liệu giám sát và điều trị HIV
3. Các chỉ số ảnh hưởng tới hiệu năng cũng như khả năng ra quyết định trong việc dự báo sớm và ra quyết định điều trị HIV.

4.3. Phạm vi nghiên cứu

Phạm vi nghiên cứu của luận án được giới hạn trong các nội dung sau:

1. Dự báo ngắn hạn số ca nhiễm HIV mới dựa trên dữ liệu giám sát dịch tễ, với phạm vi dữ liệu được thu thập từ hệ thống giám sát HIV của một đô thị lớn.
2. Dự báo nguy cơ dựa trên phân tích sống sót được thực hiện trên dữ liệu điều trị HIV nhằm dự đoán thời gian xảy ra các biến cố y tế quan trọng trong quá trình điều trị.
3. Dự báo kết cục dưới can thiệp bằng suy luận nhân quả được triển khai trên dữ liệu quan sát liên quan đến điều trị và can thiệp HIV nhằm đánh giá hiệu quả của các biện pháp can thiệp.
4. Luận án tập trung nghiên cứu và phát triển các mô hình học sâu phục vụ dự báo và hỗ trợ ra quyết định, không đi sâu vào các nghiên cứu về sinh học phân tử hoặc cơ chế bệnh học của HIV.

5. Phương pháp nghiên cứu

Nghiên cứu lý thuyết: Thu thập, tổng hợp và phân tích các nghiên cứu có liên quan đến các mô hình trí tuệ nhân tạo trong dự báo và ra quyết định trong lĩnh vực dịch bệnh. Từ đó phân tích các ưu điểm, hạn chế và khoảng trống nghiên cứu trong các hướng nghiên cứu hiện tại, để đề xuất các mô hình cải thiện hiệu năng cũng như khả năng ra quyết định trong các mô hình hiệu quả hơn.

Nghiên cứu thực nghiệm: Các thuật toán đề xuất đã được thử nghiệm đánh giá và so sánh với các kỹ thuật hiện tại nhằm kiểm chứng hiệu quả cũng như cải tiến của cách tiếp cận trong luận án. Đối với mô hình dự báo dịch bệnh, mô hình được kiểm chứng trên bộ dữ liệu ẩn danh, hồi cứu từ hệ thống giám sát ca nhiễm HIV của thành phố Hồ Chí Minh từ tháng 1/2008 tới tháng 12/2009. Trong khi đó, mô hình phân tích sống sót được triển khai trên các bộ dữ liệu công khai về thử nghiệm lâm sàng nổi tiếng là bộ dữ liệu ACTG 175 [35] và ACTG 388 [36]. ACTG 175 được sử dụng như một bộ dữ liệu dạng tĩnh, trong đó mỗi mẫu được biểu diễn bởi các đặc trưng lâm sàng và điều trị tại một thời điểm. Ngược lại, ACTG 388 là dữ liệu dạng dọc (longitudinal), bao gồm các quan sát lặp theo thời gian như tải lượng virus và CD4, cho phép khai thác động học theo thời gian bằng các mô hình học sâu xử lý chuỗi. Cuối cùng, mô hình ước lượng nhân quả được triển khai trên hai bộ dữ liệu công khai, bao gồm dữ liệu khảo sát dân số và sức khỏe Demographic and Health Surveys (DHS) của Ethiopia (EDHS) có tích hợp kết quả xét nghiệm HIV (biomarker) [37], và bộ dữ liệu thử nghiệm lâm sàng mở rộng của ACTG 175 [38]. Dữ liệu EDHS cung cấp thông tin ở mức cộng đồng với tính đại diện dân số cao, trong khi bộ dữ liệu mở rộng của ACTG 175 phản

ánh dữ liệu thử nghiệm lâm sàng có kiểm soát. Sự kết hợp này cho phép kiểm chứng khả năng ước lượng nhân quả trên cả hai bối cảnh, từ đó nâng cao độ tin cậy và khả năng khái quát của mô hình.

6. Những đóng góp chính trong luận án

1. Mô hình dự báo ngắn hạn HIV dựa trên đồ thị ca bệnh không gian–thời gian

Nghiên cứu đề xuất mô hình dự báo ngắn hạn số ca nhiễm HIV mới theo từng khu vực dựa trên đồ thị ca bệnh không gian–thời gian. Khác với các phương pháp hiện nay chủ yếu biểu diễn dữ liệu ở cấp khu vực hoặc chuỗi số ca, mô hình đề xuất chuyển sang biểu diễn ở cấp ca bệnh, cho phép học trực tiếp động lực lan truyền của dịch trên quy mô toàn thành phố. Các đóng góp chính của nghiên cứu gồm:

a) Đề xuất biểu diễn đồ thị ca bệnh thống nhất ở quy mô toàn thành phố.

Nghiên cứu đề xuất biểu diễn dữ liệu giám sát HIV dưới dạng một đồ thị thống nhất, trong đó mỗi nút biểu diễn một ca bệnh và các cạnh được xây dựng từ mối tương quan không gian, thời gian và dịch tễ. Cách biểu diễn này bảo toàn thông tin ở cấp ca bệnh và phản ánh trực tiếp cấu trúc lan truyền của dịch, thay vì chỉ khai thác số ca tổng hợp theo từng khu vực.

b) Đề xuất khung học lan truyền từ toàn thành phố xuống khu vực.

Nghiên cứu đề xuất một khung học trên đồ thị toàn thành phố bằng mạng nơ-ron đồ thị có cơ chế chú ý để học biểu diễn lan truyền toàn cục, sau đó chuyển tri thức này xuống các đồ thị con theo từng quận, huyện nhằm thực hiện dự báo. Cách tiếp cận này kết hợp thông tin lan truyền liên khu vực với đặc trưng cục bộ của từng khu vực, góp phần nâng cao hiệu quả dự báo ngắn hạn.

Bảng 5 chỉ ra các đóng góp của mô hình đề xuất với các nghiên cứu hiện tại.

Bảng 5: Các cải tiến của mô hình đề xuất so sánh với các nghiên cứu hiện tại

Nhóm mô hình	Đại diện	Không gian	Thời gian	Yếu tố dịch tễ	Mức ca bệnh	Tích hợp xu hướng lan truyền đa khu vực vào dự báo
Thống kê	ARIMA	không	có	không	không	không
Học sâu chuỗi thời gian	LSTM, GRU, CNN	không	có	không	không	không
Mô hình không gian–thời gian (ST-GNN)	GCN-LSTM, MPNN, MSGNN	có	có	hạn chế	không	không
Mô hình đề xuất		có	có	có	có	có

Các nghiên cứu về dự báo lây nhiễm HIV đã được thể hiện trong các công trình [CT2] [CT3] [CT6].

2. Đề xuất Mô hình phân tích sống sót cá nhân hóa dựa trên học trạng thái bệnh tiềm ẩn.

Nghiên cứu đề xuất một mô hình phân tích sống sót cho dữ liệu dọc trong điều trị

HIV nhằm dự báo nguy cơ xảy ra biến cố và hỗ trợ quyết định điều trị. Khác với các mô hình chỉ khai thác chuỗi quan sát hoặc tối ưu một mục tiêu dự báo, phương pháp đề xuất kết hợp học tri thức quần thể, mô hình hóa trạng thái bệnh tiềm ẩn và học đa nhiệm trong một khung thống nhất. Các đóng góp chính gồm:

a) Đề xuất cơ chế học tri thức quần thể thông qua bộ nhớ nguyên mẫu.

Nghiên cứu đề xuất mô-đun bộ nhớ học các mẫu lâm sàng từ toàn bộ quần thể bệnh nhân và sử dụng các nguyên mẫu này để tăng cường biểu diễn của từng bệnh nhân. Cơ chế này giúp mô hình tận dụng tri thức quần thể, cải thiện khả năng tổng quát hóa đối với bệnh nhân có dữ liệu theo dõi ngắn hoặc không đầy đủ.

b) Đề xuất cơ chế học trạng thái bệnh tiềm ẩn cho dữ liệu dọc.

Nghiên cứu đề xuất mô hình hóa tiến triển bệnh thông qua trạng thái bệnh tiềm ẩn (latent disease state) được học từ chuỗi dữ liệu theo thời gian bằng mô hình không gian trạng thái có chọn lọc (Mamba). Thay vì học trực tiếp từ các quan sát lâm sàng, mô hình xây dựng một biểu diễn động phản ánh quá trình tiến triển bệnh và sử dụng biểu diễn này làm cơ sở cho các nhiệm vụ dự báo tiếp theo.

c) Đề xuất kiến trúc dự báo đa nhiệm cho phân tích sống sót và hỗ trợ điều trị.

Nghiên cứu đề xuất kiến trúc dự báo đa nhiệm gồm hai nhánh: dự báo sống sót và dự báo kết cục điều trị, cùng khai thác trạng thái bệnh tiềm ẩn đã học từ dữ liệu dọc. Trong đó, nhánh dự báo sống sót thực hiện ước lượng nguy cơ và đường cong sống sót cá thể hóa, trong khi nhánh dự báo kết cục điều trị theo tiếp cận Bayes dự báo kết cục của từng phương án điều trị và lượng hóa độ bất định của dự báo, qua đó hỗ trợ ra quyết định điều trị cá thể hóa.

Bảng 6 minh họa so sánh đóng góp của mô hình đề xuất đối với các nghiên cứu hiện tại về phân tích sống sót.

Bảng 6: Các cải tiến của mô hình đề xuất so sánh với các nghiên cứu hiện tại

Nhóm mô hình	Đại diện	Phi tuyến	Theo thời gian	Trạng thái tiềm ẩn	Động lực liên tục	Kết hợp điều trị
Thống kê	Cox PH	không	có	không	không	không
Học sâu	DeepSurv, DeepHit	có	có	không	hạn chế	không
Phân tích sống sót dựa trên cơ chế chú ý	SurvTRACE	có	có	hạn chế	hạn chế	không
Mô hình đề xuất		có	có	có	có	có

Các nghiên cứu học sâu trong phân tích sống sót trong HIV đã được thể hiện trong các công trình [CT4] [CT5] [CT8].

3. Mô hình ước lượng nhân quả kết hợp học biểu diễn sâu và ước lượng nhân quả bằng học tăng cường

Nghiên cứu đề xuất một khung học thống nhất nhằm ước lượng hiệu quả can thiệp và dự báo kết cục điều trị từ dữ liệu quan sát. Khác với các phương pháp kết hợp cố định giữa học sâu và suy luận nhân quả, phương pháp đề xuất tổ chức quá trình suy luận theo ba thành phần hỗ trợ, bao gồm mô hình hóa cấu trúc nhân quả, học biểu diễn

nhân quả và điều phối thích ứng theo từng cá thể. Các đóng góp chính gồm:

a) Đề xuất khung mô hình hóa nhân quả kết hợp tri thức chuyên gia và dữ liệu.

Nghiên cứu đề xuất xây dựng đồ thị nhân quả từ cả tri thức chuyên gia và dữ liệu quan sát, đồng thời sử dụng phân tích độ ổn định để lựa chọn cấu trúc phù hợp phục vụ suy luận nhân quả. Cách tiếp cận này giúp lựa chọn tập biến đầu vào theo nguyên lý nhân quả thay vì chỉ dựa trên tương quan thống kê.

b) Đề xuất khung học biểu diễn kết hợp với ước lượng hiệu quả can thiệp.

Nghiên cứu đề xuất một kiến trúc hai nhánh, trong đó một nhánh học biểu diễn sâu của dữ liệu, còn nhánh kia thực hiện ước lượng hiệu quả can thiệp theo nguyên lý doubly robust. Cấu trúc này cho phép khai thác đồng thời khả năng học biểu diễn phi tuyến và tính bền vững của suy luận nhân quả trong dữ liệu quan sát.

c) Đề xuất cơ chế điều phối thích ứng theo từng cá thể.

Nghiên cứu đề xuất cơ chế điều phối dựa trên học tăng cường nhằm lựa chọn động cách kết hợp giữa hai nhánh suy luận cho từng cá thể thay vì sử dụng một chiến lược cố định. Cách tiếp cận này giúp mô hình thích ứng với tính dị biệt của dữ liệu và nâng cao độ ổn định của ước lượng hiệu quả can thiệp.

Bảng 7 minh họa so sánh đóng góp của mô hình đề xuất đối với các nghiên cứu hiện tại về suy luận nhân quả.

Bảng 7: Các cải tiến của mô hình đề xuất so với các nghiên cứu hiện tại

Nhóm mô hình	Đại diện	Phi tuyến	Xử lý không đồng nhất	Biểu diễn sâu	Tối ưu kết hợp biểu diễn sâu và ước lượng nhân quả
Thống kê truyền thống	IPW, g-formula, Cox-based	không	hạn chế	không	không
Học máy nhân quả	Causal Forest, DML, ORF, X-Learner	có	có	không	không
học sâu nhân quả	TARNet, DragonNet, CEVAE	có	có	có	không
Mô hình nhân quả dựa trên Transformer	CAT	có	có	có	không
Mô hình đề xuất		có	có	có	có

Các nghiên cứu về ước lượng nhân quả HIV đã được thể hiện trong công trình [CT1]. .

7. Cấu trúc luận án

Luận án gồm phần mở đầu, bốn chương nội dung chính, kết luận, danh mục công trình công bố và tài liệu tham khảo.

Phần mở đầu trình bày bối cảnh nghiên cứu, động lực nghiên cứu, các thách thức, bài toán nghiên cứu, khoảng trống nghiên cứu, mục tiêu, phạm vi nghiên cứu, phương pháp nghiên cứu, khung nghiên cứu cũng như các đóng góp chính của luận án.

Chương 1 trình bày tổng quan các nghiên cứu liên quan đến dự báo ngắn hạn, phân tích sống sót và phân tích nhân quả, đồng thời giới thiệu các cơ sở lý thuyết và phương

pháp nền tảng phục vụ cho các nghiên cứu được đề xuất trong luận án.

Chương 2 tập trung vào bài toán dự báo ngắn hạn dịch HIV, đề xuất mô hình học sâu dựa trên đồ thị không gian–thời gian nhằm khai thác các mối quan hệ dịch tễ học giữa các ca nhiễm để hỗ trợ giám sát dịch tễ.

Chương 3 tập trung vào bài toán phân tích sống sót trong điều trị HIV, đề xuất mô hình học sâu cho tiên lượng nguy cơ và thời gian xảy ra các sự kiện điều trị, phục vụ hỗ trợ quản lý điều trị cá nhân hóa.

Chương 4 tập trung vào bài toán ước lượng nhân quả từ dữ liệu quan sát thực tế, đề xuất phương pháp hỗ trợ đánh giá hiệu quả can thiệp và hỗ trợ ra quyết định trong quản lý HIV.

Phần kết luận tổng hợp các kết quả chính của luận án, phân tích các hạn chế còn tồn tại và đề xuất các hướng nghiên cứu tiếp theo.

Chương 1. CƠ SỞ LÝ THUYẾT

Chương này trình bày các nền tảng lý thuyết cần thiết cho các bài toán nghiên cứu của luận án, bao gồm dự báo dịch bệnh ngắn hạn, học sâu trên dữ liệu chuỗi thời gian và cấu trúc đồ thị, phân tích sống sót, cũng như suy luận nhân quả trong dữ liệu quan sát.

1.1. Cơ sở lý thuyết cho dự báo dịch bệnh ngắn hạn

1.1.1. Lý thuyết chung về dự báo dịch bệnh

Dự báo dịch bệnh là một thành phần quan trọng trong hệ thống giám sát y tế công cộng, hỗ trợ các nhà quản lý dự đoán xu hướng diễn biến của dịch bệnh và chủ động xây dựng các chiến lược ứng phó phù hợp. Tùy thuộc vào khoảng thời gian dự báo và mục tiêu sử dụng, dự báo dịch bệnh có thể được chia thành dự báo ngắn hạn, trung hạn và dài hạn. Trong đó, dự báo ngắn hạn tập trung vào việc ước lượng diễn biến của dịch bệnh trong một khoảng thời gian tương đối gần, thường từ vài ngày đến vài tuần hoặc vài tháng tiếp theo, nhằm phục vụ các quyết định quản lý mang tính vận hành như phân bổ nguồn lực, tổ chức xét nghiệm, điều chỉnh hoạt động giám sát hoặc triển khai các biện pháp can thiệp kịp thời.

Khác với dự báo dài hạn vốn quan tâm nhiều hơn đến xu hướng tổng thể hoặc các kịch bản phát triển trong tương lai, dự báo ngắn hạn nhấn mạnh khả năng phản ánh chính xác trạng thái hiện tại và những biến động có thể xảy ra trong tương lai gần. Điều này đặt ra yêu cầu cho mô hình dự báo không chỉ học được xu hướng chung của dữ liệu theo thời gian mà còn phải đủ nhạy để nhận diện các thay đổi cục bộ, biến động bất thường hoặc các tín hiệu cảnh báo sớm có thể ảnh hưởng đến diễn biến dịch bệnh. Trong bối cảnh y tế công cộng, độ chính xác của dự báo ngắn hạn có ý nghĩa đặc biệt quan trọng vì đây là khoảng thời gian trực tiếp phục vụ quá trình ra quyết định và tổ chức đáp ứng thực tế.

Về mặt toán học, bài toán dự báo ngắn hạn có thể được mô tả như một bài toán học ánh xạ từ dữ liệu quan sát lịch sử sang các giá trị trong tương lai. Giả sử chuỗi dữ liệu quan sát theo thời gian được ký hiệu là:

$$X = \{x_1, x_2, \dots, x_T\} \quad (1.1)$$

trong đó x_t biểu diễn trạng thái của hệ thống tại thời điểm t , chẳng hạn như số ca bệnh mới được ghi nhận, các chỉ số dịch tễ tổng hợp hoặc biểu diễn trạng thái của hệ thống tại thời điểm đó. Mục tiêu của dự báo ngắn hạn là học một hàm ánh xạ:

$$\hat{Y}_{T+1:T+H} = f(X) \quad (1.2)$$

trong đó $f(\cdot)$ là mô hình dự báo, H là khoảng thời gian dự báo, và $\hat{Y}_{T+1:T+H}$ là chuỗi giá trị dự báo trong tương lai.

Trong trường hợp đơn giản, bài toán này có thể chỉ dựa trên một chuỗi thời gian đơn biến, nơi mô hình khai thác mối quan hệ giữa các quan sát lịch sử để suy ra xu hướng tương lai. Tuy nhiên, trong các hệ thống dịch bệnh thực tế, diễn biến của dịch

bệnh thường chịu tác động đồng thời bởi nhiều yếu tố khác nhau, khiến cấu trúc bài toán trở nên phức tạp hơn đáng kể.

Một đặc trưng nền tảng của dự báo dịch bệnh là tính phụ thuộc theo thời gian. Số ca bệnh tại thời điểm hiện tại không xuất hiện độc lập mà chịu ảnh hưởng từ trạng thái của hệ thống trong quá khứ. Nếu ký hiệu số ca bệnh tại thời điểm t là y_t , khi đó mối quan hệ tổng quát có thể được mô tả như:

$$y_t = g(y_{t-1}, y_{t-2}, \dots, y_{t-k}) \quad (1.3)$$

trong đó k biểu diễn độ dài lịch sử có ảnh hưởng đến trạng thái hiện tại. Giả định này phản ánh nguyên lý cơ bản của các mô hình chuỗi thời gian, trong đó trạng thái hiện tại được quyết định một phần bởi các quan sát trước đó. Tuy nhiên, trong thực tiễn dịch bệnh, động học lây truyền thường không chỉ phụ thuộc vào diễn biến thời gian đơn thuần mà còn chịu ảnh hưởng đồng thời bởi các yếu tố không gian và dịch tễ [39].

Bên cạnh yếu tố thời gian, sự lan truyền của dịch bệnh thường gắn với yếu tố không gian. Các ca bệnh không phân bố ngẫu nhiên mà có thể tập trung theo các khu vực địa lý, cộng đồng hoặc các mạng lưới tiếp xúc có liên hệ với nhau. Vì vậy, trạng thái dịch bệnh tại một thời điểm không chỉ phản ánh xu hướng lịch sử mà còn chịu ảnh hưởng bởi cấu trúc phân bố không gian của dịch bệnh. Đồng thời, các đặc điểm dịch tễ học cũng đóng vai trò quan trọng trong việc hình thành động học lây truyền. Những khác biệt về nhóm nguy cơ, hành vi phơi nhiễm, đặc điểm dân số hoặc khả năng tiếp cận dịch vụ y tế có thể tạo ra các mô hình lây truyền khác nhau ngay cả trong những bối cảnh có số lượng ca bệnh tương đương.

Do đó, trong bối cảnh thực tế, bài toán dự báo dịch bệnh có thể được mở rộng từ mô hình chuỗi thời gian đơn biến sang một mô hình đa chiều hơn:

$$y_t = g(T_t, S_t, E_t) \quad (1.4)$$

trong đó T_t biểu diễn thông tin phụ thuộc theo thời gian, S_t biểu diễn yếu tố không gian, và E_t biểu diễn các đặc điểm dịch tễ học liên quan đến trạng thái của hệ thống tại thời điểm t . Biểu diễn này cho thấy diễn biến của dịch bệnh không chỉ là kết quả của chuỗi quan sát theo thời gian mà là sự tổng hợp của nhiều thành phần có thể tương tác đồng thời với nhau.

Đối với HIV, mức độ phức tạp của bài toán còn cao hơn do đặc điểm lây truyền khác biệt so với nhiều bệnh truyền nhiễm cấp tính. HIV không tạo ra các cụm bùng phát ngắn hạn rõ ràng như cúm hoặc sốt xuất huyết mà có tiến trình lây truyền âm thầm, kéo dài và chịu ảnh hưởng bởi nhiều yếu tố như hành vi nguy cơ, mạng lưới tiếp xúc, khả năng tiếp cận xét nghiệm, thời điểm phát hiện bệnh và mức độ tiếp cận điều trị. Do đó, số ca được ghi nhận tại một thời điểm không phản ánh trực tiếp trạng thái lây truyền tức thời mà là kết quả tổng hợp của nhiều quá trình tiềm ẩn diễn ra trong quá khứ.

Bên cạnh đó, dữ liệu giám sát HIV thường được ghi nhận ở mức ca bệnh, trong đó mỗi trường hợp mang theo thông tin về thời điểm phát hiện, vị trí ghi nhận và các đặc

điểm dịch tễ liên quan. So với các chuỗi số liệu tổng hợp truyền thống, loại dữ liệu này cung cấp mức độ chi tiết cao hơn về cấu trúc của hệ thống dịch bệnh. Tuy nhiên, chính đặc điểm này cũng làm gia tăng độ phức tạp của bài toán mô hình hóa, bởi mối quan hệ giữa các ca bệnh có thể mang tính tiềm ẩn, không tuyến tính và thay đổi theo thời gian. Khi dữ liệu được tổng hợp ở cấp độ cao như quận hoặc thành phố, nhiều thông tin vì mô phản ánh cấu trúc lây truyền tiềm ẩn giữa các ca bệnh có thể bị mất đi, làm giảm khả năng mô hình hóa chính xác động học của dịch bệnh.

Từ các phân tích trên có thể thấy rằng dự báo dịch bệnh ngắn hạn, đặc biệt trong bối cảnh HIV, không chỉ là bài toán học phụ thuộc theo thời gian mà là bài toán mô hình hóa động học đa chiều với sự tương tác đồng thời giữa yếu tố thời gian, không gian và đặc điểm dịch tễ học. Điều này đặt ra nhu cầu sử dụng các phương pháp mô hình hóa có khả năng biểu diễn đồng thời nhiều loại quan hệ phức tạp trong dữ liệu. Trong bối cảnh đó, cấu trúc đồ thị là một hướng tiếp cận phù hợp nhằm mô hình hóa các mối liên hệ giữa các thành phần trong hệ thống dịch bệnh, nội dung sẽ được trình bày trong mục tiếp theo.

1.1.2. Các mô hình học sâu cho dự báo chuỗi thời gian dịch bệnh

Đối với các bệnh truyền nhiễm có tính chu kỳ như cúm mùa, sốt xuất huyết, COVID-19 hoặc HIV/AIDS, các mạng hồi tiếp đã cho thấy hiệu quả vượt trội trong việc khai thác các quy luật ẩn trong chuỗi dữ liệu thời gian. Nhờ vào khả năng mô hình hóa quan hệ phi tuyến và phụ thuộc dài hạn, học sâu cung cấp một công cụ mạnh mẽ để dự báo chính xác xu hướng và mức độ lan truyền của dịch bệnh trong tương lai.

Dựa trên giả định rằng trạng thái hiện tại của hệ thống phụ thuộc vào các quan sát trong quá khứ, nhiều mô hình học sâu đã được phát triển nhằm khai thác phụ thuộc theo thời gian trong dữ liệu chuỗi. Một trong những kiến trúc nền tảng là mạng nơ-ron hồi quy (RNN) [40], trong đó trạng thái ẩn được cập nhật tuần tự theo từng thời điểm để lưu trữ thông tin lịch sử. Về mặt toán học, cơ chế cập nhật của RNN có thể được biểu diễn như sau:

$$h_t = \phi(W_h h_{t-1} + W_x x_t + b) \quad (1.5)$$

trong đó x_t là dữ liệu đầu vào tại thời điểm t , h_t là trạng thái ẩn biểu diễn thông tin tích lũy từ quá khứ, W_h và W_x là các ma trận trọng số học được, b là hệ số dịch, và $\phi(\cdot)$ là hàm kích hoạt phi tuyến. Mặc dù RNN cho phép mô hình hóa phụ thuộc theo thời gian, kiến trúc này thường gặp khó khăn trong việc học các phụ thuộc dài hạn do hiện tượng tiêu biến hoặc bùng nổ gradient trong quá trình huấn luyện.

Để khắc phục hạn chế này, các biến thể như mạng bộ nhớ ngắn dài hạn (LSTM) [41] và Mạng hồi quy có cổng (GRU) [42] được phát triển nhằm cải thiện khả năng ghi nhớ thông tin dài hạn. Đặc biệt, LSTM bổ sung một cơ chế bộ nhớ thông qua trạng thái ô nhớ cùng các cổng điều khiển bao gồm cổng quên, cổng vào và cổng ra, cho phép mô hình lựa chọn thông tin cần duy trì, cập nhật hoặc loại bỏ trong quá trình truyền trạng thái theo thời gian. Về mặt toán học, trạng thái của LSTM có thể được mô tả tổng quát như sau:

$$(h_t, c_t) = \text{LSTM}(x_t, h_{t-1}, c_{t-1}) \quad (1.6)$$

trong đó h_t là trạng thái ẩn tại thời điểm t , còn c_t là trạng thái bộ nhớ giúp duy trì thông tin dài hạn. Nhờ cơ chế này, LSTM và GRU cho thấy hiệu quả tốt hơn RNN truyền thống trong nhiều bài toán dự báo chuỗi thời gian.

Các nghiên cứu như [8, 43] cho thấy LSTM và GRU là hai kiến trúc phổ biến trong dự báo HIV và COVID-19. Các biến thể như BiLSTM [42], Stacked-LSTM [44], ARIMA-LSTM [9], và LSTM-Markov [45] được sử dụng để cải thiện hiệu quả dự báo bằng cách học chiều hai chiều hoặc kết hợp yếu tố thống kê và xác suất.

Mạng tích chập (CNN) [46] sử dụng các bộ lọc tích chập để trích xuất đặc trưng cục bộ từ chuỗi dữ liệu, thường được triển khai dưới dạng 1D-CNN trong các bài toán chuỗi thời gian. CNN đặc biệt hiệu quả trong việc học các mẫu ngắn hạn và khai thác dữ liệu có nhiều đặc trưng, như minh họa trong mô hình DeepCov cho dự báo COVID-19 [47]. Ngoài ra, các kiến trúc kết hợp như CNN-LSTM cho phép đồng thời khai thác khả năng học đặc trưng cục bộ của CNN và phụ thuộc dài hạn của LSTM, và đã được áp dụng trong dự báo các bệnh truyền nhiễm như HIV và Lao [48, 49].

Gần đây, các mô hình dựa trên Transformer [50] cũng được ứng dụng trong dự báo chuỗi thời gian nhờ khả năng học các phụ thuộc dài hạn thông qua cơ chế tự chú ý (self-attention) mà không cần xử lý tuần tự như RNN hoặc LSTM. Cơ chế tự chú ý tổng quát được biểu diễn như sau:

$$\text{Attention}(Q, K, V) = \text{softmax} \left(\frac{QK^T}{\sqrt{d_k}} \right) V \quad (1.7)$$

trong đó Q , K , và V lần lượt là truy vấn, khóa và giá trị, còn d_k là chiều của vector khóa dùng để chuẩn hóa. Cơ chế này cho phép mô hình học linh hoạt mối quan hệ giữa các thời điểm quan sát mà không phụ thuộc vào xử lý tuần tự. Các nghiên cứu gần đây đã áp dụng Transformer trong dự báo dịch bệnh như cúm mùa và các bệnh truyền nhiễm khác, cho thấy tiềm năng trong mô hình hóa chuỗi thời gian phức tạp [51, 52].

Mặc dù các mô hình trên đã đạt được nhiều kết quả tích cực trong dự báo chuỗi thời gian, phần lớn chúng chủ yếu tập trung vào việc học phụ thuộc theo thời gian từ dữ liệu tuần tự. Trong bối cảnh dịch bệnh, đặc biệt với HIV, nơi dữ liệu không chỉ chịu ảnh hưởng bởi yếu tố thời gian mà còn gắn với phân bố không gian và các đặc điểm dịch tễ học phức tạp, các mô hình này có thể gặp hạn chế trong việc biểu diễn đầy đủ cấu trúc tương tác của hệ thống. Điều này tạo động lực cho việc phát triển các phương pháp học sâu dựa trên cấu trúc đồ thị nhằm mô hình hóa đồng thời nhiều mối quan hệ phức tạp trong dữ liệu dịch bệnh.

1.1.3. Biểu diễn đồ thị trong phân tích dịch bệnh

Dự báo dịch bệnh không chỉ là bài toán học phụ thuộc theo thời gian mà trong nhiều bối cảnh còn liên quan đến các mối quan hệ tương tác phức tạp giữa các thành phần trong hệ thống dịch tễ. Diễn biến của dịch bệnh có thể chịu ảnh hưởng đồng thời bởi

yếu tố thời gian, phân bố không gian và các đặc điểm dịch tễ học, khiến dữ liệu không còn đơn thuần là một chuỗi quan sát tuần tự mà mang bản chất quan hệ giữa nhiều thành phần có liên kết với nhau. Điều này đặc biệt có ý nghĩa đối với các bệnh truyền nhiễm có cấu trúc lây truyền phức tạp như HIV, nơi các mối liên hệ dịch tễ có thể đóng vai trò quan trọng trong động học của dịch bệnh.

Trong những trường hợp như vậy, các phương pháp biểu diễn dữ liệu truyền thống dưới dạng chuỗi hoặc bảng có thể gặp hạn chế trong việc phản ánh đầy đủ cấu trúc tương tác của hệ thống. Biểu diễn đồ thị là một khuôn khổ toán học phù hợp để mô hình hóa các hệ thống có cấu trúc quan hệ như vậy. Trong lý thuyết đồ thị, một hệ thống có thể được biểu diễn dưới dạng: $G = (V, E)$, trong đó V là tập các nút đại diện cho các thành phần trong hệ thống, còn E là tập các cạnh biểu diễn mối liên hệ giữa các nút. Cách biểu diễn này cho phép mô hình hóa không chỉ từng thành phần riêng lẻ mà còn cả cấu trúc kết nối giữa chúng.

Xét một đồ thị $G = (V, E)$ gồm N nút, với $V = \{v_1, v_2, \dots, v_N\}$ là tập nút và E là tập cạnh. Mỗi nút $v_i \in V$ được biểu diễn bởi vector đặc trưng $x_i \in \mathbb{R}^d$, trong đó d là số chiều đặc trưng của nút. Đối với mỗi cạnh nối giữa hai nút v_i và v_j , đặc trưng cạnh được ký hiệu là $z_{ij} \in \mathbb{R}^{d_e}$, trong đó d_e là số chiều của đặc trưng cạnh, mô tả các thuộc tính hoặc mức độ tương tác giữa hai nút.

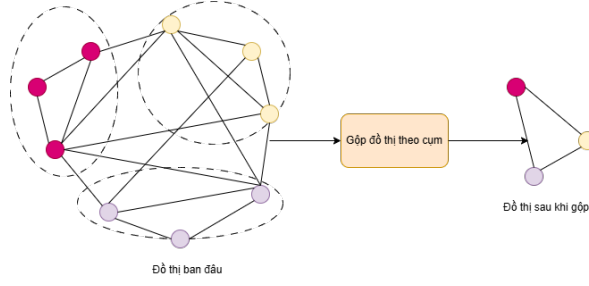
Một cách biểu diễn phổ biến của cấu trúc đồ thị là ma trận kề như sau: $A \in \mathbb{R}^{N \times N}$, trong đó phần tử A_{ij} biểu diễn sự tồn tại hoặc cường độ liên kết giữa nút v_i và v_j . Trong trường hợp đồ thị không trọng số, ma trận kề thường mang giá trị nhị phân để biểu thị có hoặc không có liên kết giữa hai nút, trong khi với đồ thị có trọng số, các phần tử có thể nhận giá trị thực để phản ánh mức độ tương tác giữa các đơn vị.

So với các biểu diễn dữ liệu truyền thống, cấu trúc đồ thị cho phép mô hình hóa trực tiếp các mối quan hệ giữa nhiều đơn vị đồng thời, đồng thời tích hợp linh hoạt thông tin về đặc trưng nút, đặc trưng cạnh và cấu trúc kết nối tổng thể. Nhờ đó, biểu diễn đồ thị trở thành nền tảng phù hợp cho việc phân tích các hệ thống dịch tễ phức tạp. Trên cơ sở này, các phương pháp học sâu trên đồ thị được phát triển nhằm học biểu diễn từ cấu trúc kết nối và lan truyền thông tin giữa các nút, nội dung sẽ được trình bày trong mục tiếp theo.

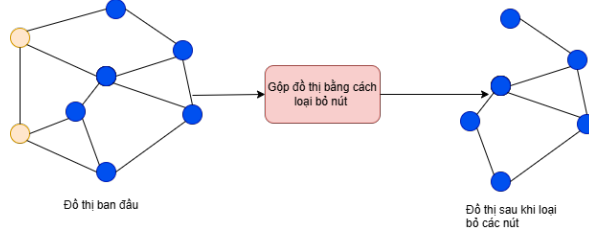
1.1.4. Các phương pháp gộp đồ thị (Graph Pooling)

Gộp đồ thị là một thành phần quan trọng trong các mô hình học sâu trên đồ thị, nhằm giảm số lượng nút và đơn giản hóa cấu trúc đồ thị để trích xuất các biểu diễn tổng quát và có ý nghĩa hơn. Theo phân loại trong [53], các kỹ thuật graph pooling được chia thành hai nhóm chính: gộp nút (node cluster pooling) và loại bỏ nút (node drop pooling).

Gộp nút là phương pháp nhóm các nút có đặc trưng tương đồng thành các cụm. Mỗi cụm sau đó được biểu diễn như một nút mới trong đồ thị rút gọn. Các phương pháp tiêu biểu thuộc nhóm này bao gồm DiffPool [54], Mincut Pool [55], StructPool [56], và MemPool [57]. Quá trình ánh xạ từ nút gốc sang nút cụm được thực hiện bằng một ma trận ánh xạ $S \in \mathbb{R}^{n \times k}$, trong đó n là số nút ban đầu và k là số cụm được tạo ra. Biểu diễn đặc trưng và cấu trúc liên kết của đồ thị mới được tính như sau:



Hình 1.1: Gộp nút



Hình 1.2: Loại bỏ nút

$$X' = S^T X, \quad A' = S^T A S$$

Trong đó, $X \in \mathbb{R}^{n \times d}$ là ma trận đặc trưng của các nút ban đầu, $A \in \mathbb{R}^{n \times n}$ là ma trận liên kết, và X', A' lần lượt là đặc trưng và liên kết của đồ thị sau khi được rút gọn.

Loại bỏ nút là phương pháp lựa chọn các nút quan trọng nhất bằng cách gán một điểm số cho từng nút dựa trên mức độ ảnh hưởng của chúng. Sau đó, chỉ giữ lại những nút có điểm số cao nhất và loại bỏ phần còn lại. Các phương pháp nổi bật trong nhóm này bao gồm TopKPooling [58], SAGPool [59], GSAPool [60], và IPool [61]. Cụ thể, với vector điểm số $s \in \mathbb{R}^n$, ta chọn ra k nút có điểm cao nhất theo chỉ số $\mathcal{I}$, rồi tạo đồ thị con tương ứng như sau:

$$\mathcal{I} = \text{TopK}(s, k), \quad X' = X[\mathcal{I}], \quad A' = A[\mathcal{I}, \mathcal{I}]$$

Trong đó, $\text{TopK}(s, k)$ là phép chọn ra chỉ số của k phần tử lớn nhất trong s , còn X', A' là đặc trưng và liên kết của đồ thị rút gọn chỉ bao gồm các nút được chọn.

Trong bối cảnh dự báo ngắn hạn dịch HIV dựa trên đồ thị, không phải mọi thông tin trong mạng lưới lây truyền đều đóng góp như nhau vào kết quả dự báo ở cấp quận/huyện. Một số mẫu quan hệ dịch tễ có thể phản ánh tín hiệu dự báo quan trọng hơn so với các thông tin ít liên quan hoặc mang tính nhiễu. Vì vậy, quá trình học biểu diễn đồ thị cần có khả năng chọn lọc và nhấn mạnh các tín hiệu dự báo quan trọng để xây dựng biểu diễn cô đọng nhưng vẫn bảo toàn thông tin thiết yếu cho bài toán dự báo.

1.1.5. Học sâu trên cấu trúc đồ thị cho dự báo dịch bệnh

Biểu diễn dữ liệu dưới dạng đồ thị cho phép mô hình hóa các mối liên hệ giữa các thành phần trong hệ thống dịch tễ. Tuy nhiên, bản thân cấu trúc đồ thị mới chỉ thể hiện các kết nối giữa các nút, còn để khai thác được các mối liên hệ này cho bài toán dự

báo cần có các mô hình học phù hợp. Đây là lý do các phương pháp học sâu trên đồ thị được phát triển nhằm học biểu diễn của các nút và quan hệ trong đồ thị [62].

Khác với dữ liệu dạng bảng hoặc chuỗi thời gian thông thường, dữ liệu đồ thị không có cấu trúc cố định theo hàng hoặc lưới đều đặn. Mỗi nút có thể có số lượng nút lân cận khác nhau, đồng thời mức độ quan trọng của các mối liên hệ cũng có thể không giống nhau. Trong bối cảnh giám sát HIV theo ca bệnh, điều này đặc biệt phù hợp vì mỗi ca ghi nhận không nên được xem xét hoàn toàn độc lập mà cần được đặt trong mối liên hệ với các ca khác thông qua thời gian phát hiện, địa bàn ghi nhận và các đặc điểm dịch tễ liên quan.

Một nguyên lý cơ bản của học sâu trên đồ thị là mỗi nút cập nhật biểu diễn của mình bằng cách tổng hợp thông tin từ các nút lân cận, thường được gọi là cơ chế truyền thông điệp. Quá trình này có thể được mô tả tổng quát như:

$$m_i^{(l)} = \text{AGGREGATE}(\{h_j^{(l)} : j \in N(i)\}) \quad (1.8)$$

$$h_i^{(l+1)} = \text{UPDATE}(h_i^{(l)}, m_i^{(l)}) \quad (1.9)$$

trong đó $N(i)$ là tập các nút lân cận của nút i , $m_i^{(l)}$ là thông tin tổng hợp từ các nút lân cận, còn $h_i^{(l+1)}$ là biểu diễn mới của nút sau khi cập nhật. Cơ chế này cho phép mỗi nút không chỉ giữ thông tin riêng mà còn tiếp nhận thêm thông tin từ các nút có liên hệ trong đồ thị.

Một trong những mô hình nền tảng là mạng tích chập đồ thị (GCN) [63], trong đó biểu diễn của mỗi nút được cập nhật bằng cách tổng hợp thông tin từ chính nó và các nút lân cận thông qua cấu trúc liên kết của đồ thị:

$$H^{(l+1)} = \sigma(\tilde{D}^{-1/2} \tilde{A} \tilde{D}^{-1/2} H^{(l)} W^{(l)}) \quad (1.10)$$

trong đó $H^{(l)}$ là biểu diễn nút tại lớp l , $\tilde{A}$ là ma trận kề có bổ sung liên kết tự thân, $\tilde{D}$ là ma trận bậc tương ứng, $W^{(l)}$ là tham số học được, và $\sigma(\cdot)$ là hàm kích hoạt phi tuyến.

Tuy nhiên, trong thực tế, không phải mọi nút lân cận đều có mức độ quan trọng như nhau. Để giải quyết vấn đề này, mạng chú ý đồ thị (Graph Attention Network – GAT) [64] được đề xuất nhằm học mức độ đóng góp khác nhau của từng nút lân cận thông qua cơ chế chú ý:

$$e_{ij} = a(W h_i, W h_j) \quad (1.11)$$

trong đó h_i và h_j lần lượt là vector đặc trưng của nút i và nút lân cận j , W là ma trận biến đổi tuyến tính học được, còn $a(\cdot)$ là hàm attention dùng để tính mức độ liên quan giữa hai nút.

Các hệ số chú ý sau đó được chuẩn hóa bằng hàm softmax theo công thức sau:

$$\alpha_{ij} = \frac{\exp(e_{ij})}{\sum_{k \in N(i)} \exp(e_{ik})} \quad (1.12)$$

trong đó $N(i)$ là tập nút lân cận của nút i , còn α_{ij} biểu diễn mức độ quan trọng tương đối của nút j đối với quá trình cập nhật nút i . Biểu diễn nút sau khi tổng hợp thông tin được cập nhật như sau:

$$h'_i = \sigma \left(\sum_{j \in N(i)} \alpha_{ij} W h_j \right) \quad (1.13)$$

trong đó $\sigma(\cdot)$ là hàm kích hoạt phi tuyến. Cơ chế này cho phép mô hình học thích ứng mức độ ảnh hưởng khác nhau giữa các nút thay vì giả định mọi nút lân cận đóng góp như nhau.

Tuy nhiên, trong bài toán dự báo dịch, học cấu trúc quan hệ trên đồ thị vẫn chưa đủ vì dữ liệu dịch tễ còn thay đổi theo thời gian. Điều này dẫn đến nhóm mô hình đồ thị không gian-thời gian, trong đó mục tiêu là học đồng thời cả cấu trúc quan hệ và động học thời gian. Nếu ký hiệu chuỗi đồ thị quan sát như sau:

$$G_{1:T} = \{G_1, G_2, \dots, G_T\} \quad (1.14)$$

thì bài toán dự báo tổng quát có thể được biểu diễn như sau:

$$\hat{Y}_{T+1:T+H} = f(G_{1:T}) \quad (1.15)$$

trong đó $G_{1:T}$ là chuỗi đồ thị trong quá khứ, H là khoảng thời gian dự báo, và $\hat{Y}_{T+1:T+H}$ là kết quả dự báo trong tương lai.

Bên cạnh cách tiếp cận sử dụng chuỗi đồ thị theo thời gian, một hướng tiếp cận khác là xây dựng một đồ thị hợp nhất, trong đó các quan hệ thời gian, không gian và đặc điểm khác được mã hóa trực tiếp vào cấu trúc đồ thị hoặc thuộc tính cạnh. Trong cách tiếp cận này, bài toán không còn được biểu diễn như một chuỗi các đồ thị rời rạc theo thời gian, mà là học biểu diễn trên một đồ thị duy nhất phản ánh đồng thời nhiều loại quan hệ giữa các nút. So với các mô hình đồ thị không gian-thời gian truyền thống, hướng tiếp cận này cho phép mô hình hóa trực tiếp cấu trúc tương tác đa chiều trong dữ liệu giám sát dịch bệnh.

Các nghiên cứu điển hình về học sâu đồ thị theo không gian - thời gian bao gồm MPNN-LSTM [12], GCN-LSTM [65], GCN-GRU [66], GAT-GRU [67], EvolveGCN [68], GCN-GRU [13], MSGNN [69].

1.2. Cơ sở lý thuyết về phân tích sống sót

1.2.1. Lý thuyết chung về phân tích sống sót

Bên cạnh các bài toán dự báo ở cấp độ quần thể, nhiều bài toán y tế tập trung vào việc tiên lượng kết cục ở cấp độ cá thể, trong đó mỗi quan tâm không chỉ là sự kiện có xảy ra hay không mà còn là thời điểm xảy ra sự kiện đó. Các bài toán như thời gian đến tử vong, tái phát bệnh, thất bại điều trị hoặc gián đoạn điều trị đều thuộc nhóm bài toán này. Khác với các bài toán phân loại hoặc hồi quy thông thường, mục tiêu ở đây không đơn thuần là dự đoán một nhãn hoặc một giá trị số, mà là mô hình hóa quá trình

diễn tiến theo thời gian cho đến khi một sự kiện quan tâm xảy ra.

Phân tích sống sót là nhánh phương pháp thống kê được phát triển để giải quyết các bài toán thời gian đến sự kiện. Giả sử T là biến ngẫu nhiên biểu diễn thời gian cho đến khi xảy ra sự kiện quan tâm. Một trong những đại lượng nền tảng của phân tích sống sót là hàm sống sót, được định nghĩa như sau:

$$S(t) = P(T > t) \quad (1.16)$$

Hàm này biểu diễn xác suất đối tượng nghiên cứu vẫn chưa xảy ra sự kiện tại thời điểm t . Trong bối cảnh y tế, điều này có thể được hiểu là xác suất một đối tượng vẫn duy trì trạng thái chưa gặp biến cố vượt qua thời điểm quan sát.

Liên quan chặt chẽ với hàm sống sót là hàm mật độ xác suất:

$$f(t) = -\frac{dS(t)}{dt} \quad (1.17)$$

đại lượng này phản ánh phân bố xác suất thời gian xảy ra sự kiện. Tuy nhiên, trong nhiều bài toán y học, đại lượng có ý nghĩa thực tiễn hơn là nguy cơ tức thời tại một thời điểm xác định, được mô tả bởi hàm nguy cơ:

$$h(t) = \lim_{\Delta t \rightarrow 0} \frac{P(t \leq T < t + \Delta t \mid T \geq t)}{\Delta t} \quad (1.18)$$

Hàm nguy cơ phản ánh tốc độ xảy ra sự kiện tại thời điểm t , với điều kiện đối tượng vẫn chưa gặp sự kiện trước thời điểm đó. Giữa hàm nguy cơ (hazard function) và hàm sống sót (survival function) tồn tại mối quan hệ toán học sau:

$$S(t) = \exp \left(- \int_0^t h(u) du \right) \quad (1.19)$$

Trong đó, $S(t)$ biểu diễn xác suất đối tượng chưa xảy ra biến cố tại thời điểm t , $h(u)$ là hàm nguy cơ tức thời tại thời điểm u , $\int_0^t h(u) du$ là nguy cơ tích lũy theo thời gian, và $\exp(\cdot)$ là hàm mũ tự nhiên. Mối quan hệ này cho thấy đây là hai cách biểu diễn khác nhau của cùng một quá trình diễn tiến theo thời gian.

Một đặc trưng quan trọng khiến phân tích sống sót khác biệt so với các bài toán học máy thông thường là hiện tượng kiểm duyệt (censoring). Trong dữ liệu thực tế, không phải mọi đối tượng đều trải qua sự kiện trong khoảng thời gian quan sát. Một số trường hợp có thể chưa xảy ra sự kiện khi nghiên cứu kết thúc hoặc không còn tiếp tục được theo dõi, khiến thời gian xảy ra sự kiện thực sự chưa được quan sát đầy đủ. Nếu loại bỏ các trường hợp này, mô hình sẽ mất đi một lượng thông tin quan trọng và có thể dẫn đến sai lệch trong quá trình ước lượng. Để phân biệt giữa các trường hợp đã quan sát được sự kiện và các trường hợp bị kiểm duyệt, một biến chỉ báo thường được định nghĩa như sau:

$$\delta_i = \begin{cases} 1, & \text{nếu sự kiện được quan sát} \\ 0, & \text{nếu bị kiểm duyệt} \end{cases} \quad (1.20)$$

Dữ liệu cho mỗi đối tượng khi đó được biểu diễn dưới dạng: (t_i, δ_i, x_i) , trong đó t_i là thời gian quan sát, δ_i là trạng thái sự kiện, và x_i là tập đặc trưng đầu vào.

Trong phân tích sống sót cổ điển, ước lượng Kaplan–Meier [20] là phương pháp không tham số phổ biến để ước lượng hàm sống sót, được thể hiện như sau:

$$\hat{S}(t) = \prod_{t_i \leq t} \left(1 - \frac{d_i}{n_i}\right) \quad (1.21)$$

trong đó d_i là số sự kiện xảy ra tại thời điểm t_i , còn n_i là số đối tượng vẫn còn trong nhóm nguy cơ ngay trước thời điểm đó. Phương pháp này hữu ích trong mô tả dữ liệu và so sánh các nhóm, nhưng không phù hợp khi cần mô hình hóa ảnh hưởng đồng thời của nhiều đặc trưng đầu vào. Một trong những mô hình nền tảng trong phân tích sống sót là mô hình Cox proportional hazards [21], thể hiện qua công thức sau:

$$h(t|x) = h_0(t) \exp(\beta^T x) \quad (1.22)$$

trong đó $h_0(t)$ là hàm nguy cơ nền, x là vector đặc trưng đầu vào, và β là hệ số hồi quy cần ước lượng. Mô hình này cho phép lượng hóa ảnh hưởng của các yếu tố nguy cơ mà không cần giả định dạng cụ thể của hàm nguy cơ nền, dưới giả định rằng tỷ lệ nguy cơ giữa các đối tượng là không đổi theo thời gian.

Mặc dù mô hình Cox là nền tảng quan trọng trong nghiên cứu y sinh học, mô hình này vẫn tồn tại một số hạn chế trong các bài toán hiện đại. Thứ nhất, mối quan hệ giữa các đặc trưng đầu vào và nguy cơ xảy ra sự kiện có thể mang tính phi tuyến mạnh. Thứ hai, trong nhiều bài toán lâm sàng, đặc trưng đầu vào không phải là các giá trị cố định mà thay đổi theo thời gian, phản ánh diễn tiến liên tục của đối tượng nghiên cứu. Thứ ba, quỹ đạo tiến triển giữa các đối tượng có thể rất không đồng nhất, khiến các giả định thống kê cổ điển trở nên hạn chế.

Trong bối cảnh này, các phương pháp học sâu ngày càng được quan tâm như một hướng tiếp cận mở rộng cho phân tích sống sót, nhờ khả năng học biểu diễn phi tuyến và mô hình hóa dữ liệu phức tạp. Đặc biệt, đối với dữ liệu điều trị dọc, việc học biểu diễn theo thời gian trở thành một thành phần quan trọng trong quá trình tiên lượng. Điều này tạo cơ sở cho việc tích hợp các phương pháp học sâu có khả năng mô hình hóa động học thời gian trong bài toán phân tích sống sót, nội dung sẽ được trình bày trong mục tiếp theo.

1.2.2. Các phương pháp học sâu trong phân tích sống sót

Sự phát triển của học sâu đã mở rộng đáng kể khả năng của phân tích sống sót trong việc mô hình hóa các mối quan hệ phi tuyến và xử lý dữ liệu có cấu trúc phức tạp. Khác với các mô hình thống kê cổ điển vốn phụ thuộc vào các giả định tuyến tính hoặc các dạng hàm nguy cơ cố định, các phương pháp học sâu cho phép học trực tiếp biểu diễn đặc trưng từ dữ liệu đầu vào, từ đó xây dựng các mô hình tiên lượng linh hoạt hơn đối với các bài toán thời gian đến sự kiện.

Một trong những phương pháp tiên phong là DeepSurv [22], trong đó mạng nơ-ron sâu được sử dụng để thay thế thành phần tuyến tính trong mô hình Cox proportional

hazards. Hàm nguy cơ khi đó được biểu diễn như:

$$h(t|x) = h_0(t) \exp(f_\theta(x)) \quad (1.23)$$

trong đó $f_\theta(x)$ là hàm phi tuyến được học bởi mạng nơ-ron sâu. Cách tiếp cận này cho phép mô hình hóa các mối quan hệ phức tạp giữa đặc trưng đầu vào và nguy cơ xảy ra sự kiện, vượt qua hạn chế tuyến tính của Cox truyền thống. Tuy nhiên, DeepSurv vẫn kế thừa giả định proportional hazards, nghĩa là tỷ lệ nguy cơ tương đối giữa các đối tượng được giả định ổn định theo thời gian, đồng thời chủ yếu phù hợp với dữ liệu đặc trưng tĩnh thay vì dữ liệu thay đổi theo thời gian.

Một hướng tiếp cận khác là DeepHit [23], trong đó thay vì chỉ ước lượng hazard function, mô hình học trực tiếp phân bố xác suất xảy ra sự kiện tại từng thời điểm:

$$P(T = t|x) = f_\theta(x, t) \quad (1.24)$$

trong đó f_θ là hàm ánh xạ phi tuyến được học bởi mạng sâu. Cách tiếp cận này cho phép dự báo trực tiếp phân bố thời gian xảy ra sự kiện và linh hoạt hơn trong mô hình hóa survival distribution. Tuy nhiên, việc biểu diễn thời gian dưới dạng rời rạc có thể làm giảm khả năng phản ánh tính liên tục của tiến trình lâm sàng, đặc biệt khi thời gian theo dõi kéo dài hoặc có mật độ quan sát không đồng đều.

Đối với các bài toán phân tích sống sót sử dụng dữ liệu dọc (longitudinal), các nghiên cứu như Dynamic-DeepHit [70] đã tích hợp các mô hình tuần tự để mô hình hóa diễn tiến theo thời gian của các đặc trưng lâm sàng. Quá trình cập nhật trạng thái ẩn tổng quát có thể được mô tả như sau:

$$h_t = f(h_{t-1}, x_t) \quad (1.25)$$

trong đó x_t là đặc trưng quan sát tại thời điểm t , còn h_t là trạng thái ẩn biểu diễn lịch sử tiến triển của đối tượng. Cách tiếp cận này cho phép khai thác thông tin động theo thời gian thay vì chỉ sử dụng các đặc trưng tĩnh. Trong bối cảnh điều trị HIV, đây là hướng tiếp cận phù hợp vì trạng thái lâm sàng thường thay đổi theo quá trình điều trị, phản ánh qua các chỉ số xét nghiệm, thay đổi phác đồ hoặc các yếu tố can thiệp khác. Tuy nhiên, các mô hình này thường gặp khó khăn trong việc nắm bắt các phụ thuộc dài hạn khi chuỗi quan sát kéo dài, đồng thời hiệu quả huấn luyện có thể bị ảnh hưởng bởi đặc tính xử lý tuần tự theo từng bước thời gian.

Gần đây, các mô hình dựa trên Transformer như SurvTRACE [24] được đề xuất nhằm tận dụng cơ chế tự chú ý để học các quan hệ phụ thuộc phức tạp trong dữ liệu sống sót. Cơ chế này cho phép mô hình khai thác linh hoạt các tương tác giữa nhiều đặc trưng đầu vào mà không bị giới hạn bởi các giả định tuyến tính của các mô hình thống kê truyền thống. Trong bối cảnh dữ liệu điều trị HIV, đây là một lợi thế quan trọng vì diễn tiến lâm sàng thường chịu ảnh hưởng bởi nhiều yếu tố tương tác đồng thời như đáp ứng điều trị, thay đổi chỉ số xét nghiệm, lịch sử can thiệp hoặc các biến cố lâm sàng trước đó.

Tuy nhiên, mặc dù các phương pháp học sâu đã cải thiện đáng kể khả năng tiên

lượng trong phân tích sống sót, một số hạn chế vẫn còn tồn tại. Các mô hình mở rộng từ Cox vẫn chịu ràng buộc bởi giả định tỷ lệ nguy cơ không đổi theo thời gian, DeepHit dựa trên biểu diễn thời gian rời rạc, các mô hình tuần tự gặp khó khăn trong việc nắm bắt các phụ thuộc dài hạn khi chuỗi quan sát kéo dài, trong khi các mô hình Transformer tuy cho phép học linh hoạt các quan hệ phức tạp thông qua cơ chế tự chú ý nhưng có chi phí tính toán tăng theo bình phương chiều dài chuỗi là $O(n^2)$, nên có thể tốn kém đối với dữ liệu chuỗi dài. Trong bối cảnh điều trị HIV, các quan sát lâm sàng chỉ phản ánh các thời điểm đo lường rời rạc, trong khi tiến triển bệnh lý thực tế là một quá trình động liên tục với các trạng thái tiềm ẩn thay đổi theo thời gian và có tính khác biệt cao giữa các đối tượng.

1.2.3. Các mô hình không gian trạng thái

Các mô hình không gian trạng thái (SSM, viết tắt của state-space models) [71, 72] cung cấp một khung lý thuyết phù hợp để mô hình hóa tiến triển bệnh, trong đó trạng thái bệnh tiềm ẩn của bệnh nhân được xem như một biến trạng thái động tiến hóa theo thời gian và chi phối các quan sát lâm sàng. Các mô hình không gian trạng thái cung cấp một cách tiếp cận tuần tự để theo dõi sự thay đổi của thông tin lâm sàng theo thời gian. Trạng thái ẩn của chúng biến đổi theo thời gian và có thể phản ánh các quá trình lâm sàng tiềm ẩn [72, 73]. Các nghiên cứu nền tảng của Durbin và Koopman [74], cùng với các mở rộng sau đó của Linderman và cộng sự [75], cho thấy rằng các trạng thái ẩn có thể tóm lược các động lực lâm sàng không quan sát được. Các mô hình không gian trạng thái hiện đại không còn bị giới hạn bởi tính tuyến tính, mà cho phép trạng thái ẩn biến đổi một cách linh hoạt hơn. Điều này khiến các mô hình trở nên phù hợp với dữ liệu lâm sàng, vốn thường biến thiên giữa các cá thể và theo thời gian.

Ý tưởng cốt lõi của SSM là giả định rằng hệ thống được đặc trưng bởi một trạng thái ẩn tại mỗi thời điểm, trong đó trạng thái này được cập nhật liên tục dựa trên trạng thái trước đó và thông tin đầu vào hiện tại. Dạng tổng quát của mô hình được biểu diễn như sau:

$$z_t = f(z_{t-1}, x_t) \quad (1.26)$$

$$y_t = g(z_t) \quad (1.27)$$

trong đó z_t là trạng thái ẩn tại thời điểm t , phản ánh trạng thái bệnh tích lũy chưa quan sát trực tiếp; x_t là thông tin đầu vào tại thời điểm đó, chẳng hạn như đặc trưng lâm sàng, kết quả xét nghiệm hoặc thông tin điều trị; và y_t là đầu ra quan sát được. Hàm chuyển trạng thái $f(\cdot)$ mô tả động lực tiến hóa của trạng thái tiềm ẩn theo thời gian, trong khi hàm $g(\cdot)$ ánh xạ trạng thái này sang không gian quan sát.

Khác với các mô hình chỉ học ánh xạ trực tiếp từ đầu vào sang đầu ra, SSM nhấn mạnh quá trình tiến hóa liên tục của trạng thái tiềm ẩn, cho phép mô hình hóa rõ ràng cơ chế tích lũy thông tin lịch sử. Đặc điểm này đặc biệt phù hợp với dữ liệu lâm sàng dọc, nơi tiến triển bệnh không chỉ phụ thuộc vào quan sát hiện tại mà còn chịu ảnh hưởng bởi trạng thái bệnh tích lũy từ các thời điểm trước đó.

Trong bài toán phân tích sống sót, trạng thái ẩn z_t có thể được xem như một biểu diễn động của nguy cơ xảy ra biến cố tại thời điểm t . Từ đó, xác suất nguy cơ có điều kiện (hazard) có thể được ước lượng như một hàm của trạng thái này:

$$\lambda_i(t) = \phi(z_t) \quad (1.28)$$

trong đó $\phi(\cdot)$ là hàm ánh xạ từ trạng thái tiềm ẩn sang xác suất nguy cơ. Trong thiết lập phân tích sống sót rời rạc theo thời gian, hàm sống sót của cá thể i được xác định như sau:

$$S_i(t) = \prod_{k=1}^t (1 - \lambda_i(k)) \quad (1.29)$$

Cách biểu diễn này thiết lập mối liên hệ trực tiếp giữa động lực bệnh tiềm ẩn và mục tiêu dự báo sống sót, cho phép mô hình học không chỉ mối quan hệ tức thời giữa biến đầu vào và biến cố mà còn cả quỹ đạo tiến triển nguy cơ theo thời gian.

1.3. Bất định trong dự báo y tế

Trong các bài toán dự báo y tế, độ chính xác của mô hình là một tiêu chí quan trọng nhưng chưa đủ để đảm bảo giá trị thực tiễn trong hỗ trợ ra quyết định [76]. Trong môi trường lâm sàng, cùng một giá trị dự báo có thể đi kèm với các mức độ tin cậy rất khác nhau tùy thuộc vào mức độ đầy đủ của dữ liệu, tính đại diện của mẫu huấn luyện hoặc sự hiện diện của các trường hợp bất thường. Do đó, bên cạnh dự báo kết quả, việc biểu diễn mức độ bất định của mô hình cũng đóng vai trò quan trọng nhằm hỗ trợ đánh giá độ tin cậy của khuyến nghị.

Về mặt khái niệm, bất định trong học máy thường được chia thành hai nhóm chính [77]. Loại thứ nhất là bất định do dữ liệu (aleatoric uncertainty), phản ánh nhiễu nội tại của quá trình quan sát như sai số đo lường, dữ liệu thiếu hụt hoặc biến thiên sinh học tự nhiên. Loại thứ hai là bất định do mô hình (epistemic uncertainty), phản ánh mức độ thiếu chắc chắn của mô hình khi dữ liệu huấn luyện chưa đủ đại diện hoặc khi gặp các trường hợp ngoài phân bố đã quan sát. Trong bối cảnh điều trị HIV, hai dạng bất định này đều có ý nghĩa quan trọng do dữ liệu lâm sàng thường không đầy đủ, quỹ đạo điều trị có tính dị thể cao và các hồ sơ hiếm có thể xuất hiện trong thực tế.

Nếu dự báo thông thường chỉ ước lượng $P(y|x)$ thì trong khuôn khổ có xét bất định mô hình, mục tiêu mở rộng thành $P(y|x, D)$

trong đó x là dữ liệu đầu vào, y là kết quả cần dự báo, và D là dữ liệu huấn luyện. Điều này phản ánh rằng đầu ra dự báo không chỉ phụ thuộc vào dữ liệu hiện tại mà còn phụ thuộc vào mức độ tri thức mà mô hình đã học được [78].

Một trong những khuôn khổ tự nhiên để mô hình hóa bất định là suy luận Bayes, trong đó tham số mô hình không được xem là các giá trị cố định mà là các biến ngẫu nhiên có phân bố xác suất [78]. Theo định lý Bayes:

$$P(\theta|D) = \frac{P(D|\theta)P(\theta)}{P(D)} \quad (1.30)$$

trong đó $P(\theta)$ là phân bố tiên nghiệm, $P(D|\theta)$ là likelihood, và $P(\theta|D)$ là phân bố hậu nghiệm. Khi đó, dự báo được tính bằng cách lấy kỳ vọng trên toàn bộ phân bố tham số:

$$P(y|x, D) = \int P(y|x, \theta)P(\theta|D)d\theta \quad (1.31)$$

Cách tiếp cận này cho phép mô hình không chỉ đưa ra giá trị dự báo trung bình mà còn lượng hóa mức độ bất định liên quan đến tham số mô hình. Trong các mô hình học sâu, các cách tiếp cận học sâu theo khung bayes đã được phát triển nhằm xấp xỉ phân bố hậu nghiệm của tham số để hỗ trợ ước lượng bất định [79].

Trong bài toán tiên lượng điều trị HIV, biểu diễn bất định đặc biệt quan trọng vì dự báo thường được sử dụng để hỗ trợ phân tầng nguy cơ và ra quyết định can thiệp. Một mô hình chỉ đưa ra xác suất dự báo mà không phản ánh mức độ chắc chắn có thể dẫn đến sự tự tin quá mức trong các trường hợp khó hoặc dữ liệu bất thường. Do đó, ước lượng bất định không chỉ giúp nâng cao độ tin cậy của các mô hình tiên lượng mà còn cung cấp thông tin hữu ích cho các bài toán hỗ trợ quyết định rộng hơn, bao gồm cả các ứng dụng liên quan đến suy luận nhân quả và tối ưu hóa can thiệp.

1.4. Cơ sở lý thuyết về ước lượng nhân quả

1.4.1. Lý thuyết chung về ước lượng nhân quả

Trong các bài toán y tế công cộng, đặc biệt trong bối cảnh dịch bệnh, dự báo nguy cơ hoặc tiên lượng kết cục chủ yếu phản ánh điều gì có thể xảy ra dựa trên dữ liệu quan sát. Tuy nhiên, trong nhiều tình huống hỗ trợ ra quyết định, mục tiêu quan trọng hơn là đánh giá tác động của một can thiệp cụ thể, chẳng hạn hiệu quả của chiến lược điều trị, chương trình hỗ trợ tuân thủ hoặc các biện pháp can thiệp cộng đồng. Đây là bài toán cốt lõi của suy luận nhân quả.

Một trong những khuôn khổ nền tảng của suy luận nhân quả là khung kết cục tiềm năng (potential outcomes framework) [80]. Với mỗi đối tượng nghiên cứu, giả sử $T \in \{0, 1\}$ là biến can thiệp, trong đó $T = 1$ biểu thị có can thiệp và $T = 0$ biểu thị không can thiệp, còn Y là kết quả quan tâm. Khi đó, mỗi đối tượng được giả định có hai kết quả tiềm năng: $Y(1)$, $Y(0)$, trong đó $Y(1)$ là kết quả nếu đối tượng nhận can thiệp, còn $Y(0)$ là kết quả nếu không nhận can thiệp. Hiệu ứng điều trị cá thể được định nghĩa là: $\tau = Y(1) - Y(0)$, trong đó τ biểu diễn chênh lệch kết quả giữa hai trạng thái can thiệp đối với một đối tượng cụ thể.

Tuy nhiên, trong thực tế, với mỗi đối tượng chỉ có thể quan sát một trạng thái can thiệp, nên chỉ một trong hai kết quả tiềm năng được ghi nhận. Đây là vấn đề phản thực cơ bản trong suy luận nhân quả. Do đó, mục tiêu thường chuyển sang ước lượng hiệu ứng điều trị trung bình trên quần thể (Average Treatment Effect – ATE): $ATE = E[Y(1) - Y(0)]$, trong đó $E[\cdot]$ là toán tử kỳ vọng trên quần thể nghiên cứu. Trong bối cảnh HIV, đại lượng này có thể được hiểu là mức thay đổi trung bình của một kết cục quan tâm nếu toàn bộ đối tượng nhận can thiệp so với khi không nhận can thiệp.

Việc ước lượng tác động nhân quả từ dữ liệu quan sát đòi hỏi một số giả định nền tảng. Thứ nhất là giả định tính nhất quán (consistency), trong đó kết quả quan sát được phải tương ứng với kết quả tiềm năng dưới trạng thái can thiệp thực tế $Y = Y(T)$. Giả định này đảm bảo rằng định nghĩa can thiệp là rõ ràng và nhất quán giữa các đối tượng.

Thứ hai là giả định không có nhiễu chưa quan sát (unconfoundedness): $(Y(1), Y(0)) \perp T \mid X$ trong đó X là tập đặc trưng nền đã quan sát, còn ký hiệu $\perp$ biểu thị tính độc lập có điều kiện. Giả định này hàm ý rằng sau khi điều kiện hóa trên các biến nền phù hợp, việc phân bổ can thiệp không còn phụ thuộc vào các kết quả tiềm năng.

Thứ ba là giả định tính chồng lấn (overlap hay positivity): $0 < P(T = 1|X) < 1$. Giả định này yêu cầu rằng với mỗi tổ hợp đặc trưng quan sát được X , xác suất một đối tượng nhận can thiệp $T = 1$ phải lớn hơn 0 và nhỏ hơn 1. Nói cách khác, trong mỗi nhóm có đặc điểm tương tự nhau, phải tồn tại cả đối tượng nhận can thiệp và không nhận can thiệp. Điều này bảo đảm rằng có thể so sánh kết quả giữa hai nhóm trên cùng một miền đặc trưng, thay vì phải ngoại suy hiệu quả can thiệp cho những nhóm chỉ xuất hiện ở một trạng thái điều trị duy nhất. Điều kiện này đảm bảo rằng với mỗi cấu hình đặc trưng nền, cả hai khả năng có và không can thiệp đều có thể được quan sát.

Trong bối cảnh điều trị HIV, các giả định này đặc biệt quan trọng vì phân bổ can thiệp thường không ngẫu nhiên mà phụ thuộc vào mức độ bệnh, lịch sử điều trị, khả năng tuân thủ hoặc các yếu tố lâm sàng khác. Điều này tạo ra hiện tượng nhiễu (confounding), khi cùng một yếu tố ảnh hưởng đồng thời đến cả khả năng nhận can thiệp và kết quả điều trị.

Một công cụ phổ biến để xử lý nhiễu là điểm xu hướng (propensity score), được định nghĩa: $e(X) = P(T=1|X)$, trong đó $e(X)$ biểu diễn xác suất nhận can thiệp dựa trên đặc trưng nền X . Điểm xu hướng được sử dụng để cân bằng khác biệt giữa các nhóm quan sát. Bên cạnh đó, mô hình hóa kết quả (outcome modeling) trực tiếp học kỳ vọng kết quả có điều kiện:

$$\mu_1(X) = E[Y|X, T = 1] \quad (1.32)$$

$$\mu_0(X) = E[Y|X, T = 0] \quad (1.33)$$

trong đó $\mu_1(X)$ và $\mu_0(X)$ lần lượt là kỳ vọng kết quả khi có và không có can thiệp. Từ đó suy ra hiệu ứng điều trị cá thể như sau:

$$ITE(X) = \mu_1(X) - \mu_0(X) \quad (1.34)$$

trong đó $ITE(X)$ biểu diễn hiệu ứng điều trị cá thể có điều kiện theo đặc trưng nền X .

Các phương pháp hiện đại thường kết hợp cả mô hình phân bổ điều trị và mô hình kết quả thông qua các chiến lược ước lượng bền vững kép (doubly robust) [31] nhằm tăng độ ổn định của ước lượng, đặc biệt phù hợp với dữ liệu HIV quan sát có tính không đồng nhất cao và cơ chế điều trị phức tạp.

Mặc dù các phương pháp nhân quả cổ điển có nền tảng lý thuyết mạnh, chúng thường gặp hạn chế khi dữ liệu phức tạp, phi tuyến mạnh hoặc chứa nhiều tương tác phức tạp như trong dữ liệu y tế hiện đại. Điều này tạo động lực cho việc kết hợp suy luận nhân quả với các phương pháp học sâu nhằm xây dựng các mô hình ước lượng linh hoạt hơn.

1.4.2. Cấu trúc đồ thị nhân quả và cơ chế nhận dạng tác động can thiệp

Mô hình hóa quan hệ nhân quả là nền tảng cốt lõi trong phân tích can thiệp và ra quyết định dựa trên dữ liệu quan sát. Một trong những công cụ mạnh mẽ để biểu diễn và suy luận các quan hệ này là đồ thị nhân quả có hướng và không chu trình (Directed Acyclic Graph – DAG). DAG cho phép trực quan hóa cấu trúc sinh dữ liệu, làm rõ mối quan hệ giữa các biến ngẫu nhiên và xác định các yếu tố cần kiểm soát nhằm nhận dạng hiệu ứng nhân quả.

Trong DAG, mỗi nút (node) tương ứng với một biến ngẫu nhiên và mỗi cạnh có hướng (directed edge) biểu diễn một quan hệ nhân quả trực tiếp. Xét tập biến $V = \{X_1, \dots, X_p, T, Y\}$, trong đó T là biến can thiệp (treatment) và Y là kết cục (outcome). Một cạnh $T \rightarrow Y$ biểu thị ảnh hưởng trực tiếp của can thiệp lên kết quả. Ngược lại, nếu tồn tại biến Z sao cho $Z \rightarrow T$ và $Z \rightarrow Y$, thì Z được gọi là biến gây nhiễu (confounder), vì nó tạo ra các liên hệ giả giữa T và Y thông qua đường backdoor $T \leftarrow Z \rightarrow Y$ nếu không được điều chỉnh.

Bên cạnh đó, DAG còn phân biệt các vai trò khác của biến trong hệ thống. Biến trung gian M nằm trên đường truyền $T \rightarrow M \rightarrow Y$ và phản ánh cơ chế tác động gián tiếp của can thiệp lên kết quả; do đó, việc điều chỉnh M có thể làm sai lệch ước lượng tổng tác động. Trong khi đó, biến công cụ W ảnh hưởng đến T nhưng không tác động trực tiếp đến Y , cung cấp một chiến lược nhận dạng thay thế trong trường hợp tồn tại gây nhiễu không quan sát.

Ý nghĩa nhân quả của DAG được thiết lập thông qua mô hình phương trình cấu trúc (Structural Causal Model – SCM). Cụ thể, mỗi biến được sinh ra theo:

$$X_j = f_j(\text{Pa}_G(X_j), U_j), \quad j = 1, \dots, p, \quad (1.35)$$

$$T = f_T(\text{Pa}_G(T), U_T), \quad (1.36)$$

$$Y = f_Y(\text{Pa}_G(Y), U_Y), \quad (1.37)$$

trong đó $\text{Pa}_G(Z)$ là tập các nút cha của biến Z trong đồ thị G , và U là tập nhiễu ngoại sinh. Khi các nhiễu này độc lập, phân phối chung được phân rã theo quy tắc Markov:

$$p(v) = \prod_{Z \in V} p(z \mid \text{Pa}_G(z)). \quad (1.38)$$

Trong bối cảnh dữ liệu quan sát, sai lệch ước lượng chủ yếu phát sinh từ các đường backdoor, tức các đường đi từ T đến Y bắt đầu bằng một cạnh hướng vào T . Để loại bỏ ảnh hưởng của nhiễu, cần lựa chọn một tập biến điều chỉnh sao cho chặn được toàn bộ các đường backdoor này.

Theo lý thuyết nhân quả của Pearl [81], nếu tồn tại một tập biến S không chứa hậu

duệ của T và chặn tất cả các đường backdoor từ T đến Y , thì phân phối hậu can thiệp có thể được nhận dạng từ dữ liệu quan sát. Khi đó, phân phối $p(y \mid do(T = t))$ được biểu diễn thông qua việc điều chỉnh theo S , cụ thể:

$$p(y \mid do(T = t)) = \sum_s p(y \mid T = t, S = s) p(S = s). \quad (1.39)$$

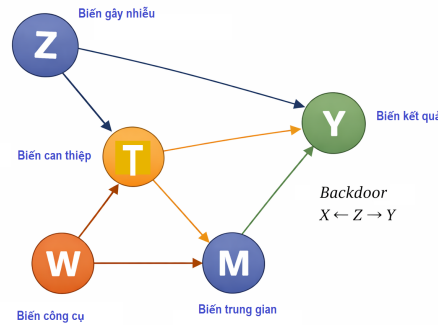
Biểu thức trên cho thấy phân phối hậu can thiệp có thể được ước lượng bằng cách lấy trung bình có trọng số của các phân phối có điều kiện theo S . Từ đó, kỳ vọng hậu can thiệp được xác định bởi:

$$\mathbb{E}[Y \mid do(T = t)] = \sum_s \mathbb{E}[Y \mid T = t, S = s] p(S = s), \quad (1.40)$$

và hiệu quả can thiệp trung bình (ATE) được xác định như sau:

$$\tau = \mathbb{E}_S [\mathbb{E}[Y \mid T = 1, S] - \mathbb{E}[Y \mid T = 0, S]]. \quad (1.41)$$

Việc lựa chọn biến điều chỉnh cần đặc biệt thận trọng. Điều chỉnh biến trung gian có thể loại bỏ một phần tác động thực của T lên Y , trong khi điều chỉnh biến hội tụ (collider), tức là biến nhận đồng thời ảnh hưởng từ T và một yếu tố khác (ví dụ $T \rightarrow C \leftarrow U$), có thể mở các đường đi không nhân quả và gây thiên lệch trong ước lượng. Do đó, DAG không chỉ hỗ trợ nhận dạng hiệu ứng nhân quả mà còn đóng vai trò quan trọng trong việc xác định chiến lược điều chỉnh phù hợp. Trong trường hợp tồn tại gây nhiễu không quan sát, biến công cụ có thể được sử dụng như một hướng tiếp cận bổ sung để tăng cường khả năng nhận dạng.



Hình 1.3: Đồ thị minh họa quan hệ giữa biến can thiệp X và kết cục Y dưới tác động của biến gây nhiễu Z , biến trung gian M và biến công cụ W .

1.4.3. Học máy nhân quả trong dữ liệu quan sát y tế

Mặc dù các phương pháp suy luận nhân quả cổ điển như propensity score matching [82], inverse probability weighting [83] và doubly robust estimation [31] đã cung cấp nền tảng quan trọng cho việc ước lượng tác động can thiệp từ dữ liệu quan sát, các phương pháp này vẫn gặp khó khăn khi áp dụng cho dữ liệu y tế có cấu trúc phức tạp. Trong thực tế, dữ liệu lâm sàng thường có số lượng đặc trưng lớn, quan hệ phi tuyến, tương tác đa chiều và sự khác biệt đáng kể giữa các cá thể. Khi đó, các mô hình thống

kê truyền thống có thể chưa đủ linh hoạt để mô hình hóa đầy đủ quan hệ giữa đặc điểm nền, điều trị và kết cục.

Trong lĩnh vực HIV, vấn đề này càng rõ rệt do kết cục điều trị có thể chịu ảnh hưởng đồng thời bởi tình trạng miễn dịch, lịch sử điều trị, hành vi tuân thủ, yếu tố xã hội và đặc điểm dịch tễ. Các yếu tố này thường không tác động độc lập mà tương tác với nhau theo cách phức tạp. Vì vậy, việc kết hợp suy luận nhân quả với học máy và học sâu trở thành một hướng tiếp cận quan trọng nhằm cải thiện khả năng học biểu diễn dữ liệu và ước lượng hiệu ứng điều trị trong các bối cảnh dữ liệu phức tạp.

Một ý tưởng trung tâm của học máy nhân quả là học biểu diễn dữ liệu trước khi thực hiện ước lượng nhân quả. Thay vì ước lượng hiệu ứng can thiệp trực tiếp từ không gian đặc trưng gốc X , mô hình học một ánh xạ sang không gian biểu diễn tiềm ẩn: $z = \phi(X)$, trong đó X là tập đặc trưng đầu vào, $\phi(\cdot)$ là hàm học biểu diễn và z là biểu diễn tiềm ẩn. Biểu diễn này được kỳ vọng giữ lại các thông tin quan trọng liên quan đến điều trị và kết cục, đồng thời làm cho bài toán ước lượng nhân quả trở nên ổn định hơn trong những dữ liệu có nhiều quan hệ phi tuyến hoặc nhiễu.

Một chủ đề quan trọng khác là tính dị thể của hiệu ứng điều trị. Trong nhiều bài toán y tế, cùng một can thiệp có thể tạo ra tác động khác nhau giữa các nhóm bệnh nhân. Nếu chỉ xem xét hiệu ứng điều trị trung bình $ATE = \mathbb{E}[Y(1) - Y(0)]$, các khác biệt quan trọng giữa các cá thể hoặc nhóm bệnh nhân có thể bị che khuất. Do đó, học máy nhân quả thường quan tâm đến các đại lượng cá thể hóa hơn như hiệu ứng điều trị cá thể:

$$ITE(x) = \mathbb{E}[Y(1) - Y(0) \mid X = x] \quad (1.42)$$

và hiệu ứng điều trị có điều kiện:

$$CATE(x) = \mathbb{E}[Y(1) - Y(0) \mid X \in \text{subgroup}] \quad (1.43)$$

Các đại lượng này giúp đánh giá tác động của can thiệp trên từng hồ sơ hoặc nhóm bệnh nhân cụ thể, từ đó hỗ trợ tốt hơn cho ra quyết định cá thể hóa.

Để xử lý các bài toán nhân quả phức tạp, nhiều mô hình học máy nhân quả hiện đại đã được phát triển. Một số phương pháp, như TARNet [29] và DragonNet [30], sử dụng mạng học sâu để học biểu diễn tiềm ẩn trước khi ước lượng hiệu ứng điều trị. Causality-Aware Transformer (CAT) [84] tích hợp mô-đun nhân quả vào kiến trúc transformer nhằm điều chỉnh trọng số attention dựa trên cấu trúc quan hệ nguyên nhân giữa các đặc trưng đầu vào. Điều này giúp mô hình học được đại diện bối cảnh giàu thông tin và nâng cao khả năng giải thích của dự báo nhân quả. Một số hướng khác, như T-learner, S-learner và X-learner, xây dựng bài toán ước lượng nhân quả thông qua các mô hình học máy riêng cho điều trị, kết cục hoặc hiệu ứng điều trị [85].

Nhìn chung, học máy nhân quả là sự kết hợp giữa nền tảng suy luận nhân quả và khả năng mô hình hóa linh hoạt của học máy hiện đại. Hướng tiếp cận này đặc biệt phù hợp với các bài toán y tế có dữ liệu phi tuyến, không đồng nhất và chứa hiệu ứng can thiệp khác nhau giữa các cá thể. Đây là cơ sở quan trọng cho các nghiên cứu sử dụng học máy và học sâu để hỗ trợ đánh giá hiệu quả can thiệp từ dữ liệu quan sát

thực tế.

1.5. Các chỉ số đánh giá

1.5.1. Các chỉ số đánh giá trong dự báo

Sai số bình phương trung bình (Mean Squared Error - MSE). Sai số bình phương trung bình đo lường độ lệch trung bình giữa giá trị dự đoán và giá trị thực tế. Nó được định nghĩa như sau:

$$\text{MSE} = \frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2 \quad (1.44)$$

Trong đó, y_i là giá trị thực tế tại thời điểm i , $\hat{y}_i$ là giá trị mô hình dự đoán, và N là số lượng quan sát. MSE càng nhỏ chứng tỏ dự báo càng gần với giá trị thực tế, nhưng dễ bị ảnh hưởng bởi các ngoại lệ (outliers) do lấy bình phương sai số.

Hệ số xác định (R^2 - Coefficient of Determination). Chỉ số R^2 đánh giá mức độ mô hình dự báo giải thích được phương sai trong dữ liệu thực tế, so với mô hình dự đoán trung bình. Nó được định nghĩa như:

$$R^2 = 1 - \frac{\sum_{i=1}^N (y_i - \hat{y}_i)^2}{\sum_{i=1}^N (y_i - \bar{y})^2} \quad (1.45)$$

Trong đó, $\bar{y}$ là giá trị trung bình của các giá trị thực tế. Chỉ số R^2 nằm trong khoảng $(-\infty, 1]$, với giá trị gần 1 cho thấy mô hình dự báo hiệu quả; ngược lại, giá trị gần 0 hoặc âm cho thấy mô hình không giải thích được biến động trong dữ liệu.

1.5.2. Các chỉ số đánh giá mô hình sống sót

C-index (Concordance Index). C-index đo lường khả năng mô hình xếp đúng thứ tự thời gian sống giữa các cặp cá thể. Một mô hình có giá trị C-index càng cao càng thể hiện khả năng phân biệt tốt giữa những người sống lâu hơn và những người xảy ra biến cố sớm hơn. Giá trị được tính theo công thức như sau.

$$C = \frac{\sum_{i,j} \mathbb{I}(T_i < T_j) \cdot \mathbb{I}(\hat{f}_i > \hat{f}_j)}{\sum_{i,j} \mathbb{I}(T_i < T_j)} \quad (1.46)$$

Trong đó, $\hat{f}_i$ là nguy cơ ước lượng của cá thể i , $\mathbb{I}$ là hàm chỉ báo, và T_i là thời gian sống thực tế. C-index nằm trong khoảng $[0.5, 1]$, với 0.5 tương đương dự đoán ngẫu nhiên.

Brier Score. Brier Score đo sai số bình phương giữa xác suất sống sót dự đoán và kết cục thực tế tại thời điểm t . Để xử lý dữ liệu kiểm duyệt, phương pháp trọng số nghịch đảo xác suất kiểm duyệt (IPCW) được áp dụng. Công thức như sau.

$$\text{Brier}(t) = \frac{1}{N} \sum_{i=1}^N w_i(t) \cdot (S(t|\mathbf{x}_i) - \mathbb{I}(T_i > t))^2 \quad (1.47)$$

Trong đó $w_i(t)$ là trọng số IPCW, $S(t|\mathbf{x}_i)$ là xác suất sống sót được mô hình ước lượng, và $\mathbb{I}(T_i > t)$ là trạng thái sống thực tại thời điểm t .

Calibration (Hiệu chỉnh xác suất). Calibration kiểm tra mức độ phù hợp giữa xác suất sống sót mô hình dự đoán và xác suất sống thực tế. Một chỉ số thông dụng là Expected Calibration Error (ECE), chia dữ liệu thành M nhóm theo xác suất dự đoán, và đo sai lệch giữa tỷ lệ sống thực tế và xác suất trung bình của mô hình như sau.

$$\text{ECE} = \sum_{m=1}^M \frac{|B_m|}{N} \cdot |\text{acc}(B_m) - \text{conf}(B_m)| \quad (1.48)$$

Trong đó, B_m là nhóm thứ m chứa các cá thể có xác suất dự đoán thuộc khoảng tương ứng; $\text{acc}(B_m)$ là tỷ lệ sống thực tế trong nhóm và $\text{conf}(B_m)$ là trung bình xác suất dự đoán.

AUC theo thời gian (Time-dependent ROC-AUC). AUC theo thời gian đánh giá khả năng mô hình phân biệt giữa các cá thể xảy ra biến cố trước và sau một thời điểm t cụ thể. Giá trị AUC càng cao thể hiện khả năng phân biệt tốt hơn tại thời điểm đó như sau.

$$\text{AUC}(t) = P(\hat{S}(t|\mathbf{x}_i) < \hat{S}(t|\mathbf{x}_j) \mid T_i \leq t, T_j > t) \quad (1.49)$$

F1-score cho bài toán phân lớp thời gian sống. Trong các thiết lập khi thời gian sống được rời rạc hoá thành các lớp (ví dụ như thời gian sống sót > 12 tháng hay không), F1-score được dùng để đánh giá độ chính xác cân bằng giữa độ nhạy và độ chính xác như sau.

$$\text{F1-score} = 2 \cdot \frac{\text{precision} \cdot \text{recall}}{\text{precision} + \text{recall}} \quad (1.50)$$

1.5.3. Các chỉ số đánh giá trong ước lượng nhân quả

Để đánh giá toàn diện khung phương pháp đề xuất, nghiên cứu sử dụng ba nhóm chỉ số chính: (i) độ chính xác ước lượng hiệu ứng nhân quả, (ii) độ tin cậy và hiệu chuẩn bất định, (iii) hiệu năng dự báo kết quả đầu ra.

Giả sử với mỗi cá thể i ta có $\tau_i = Y_i(1) - Y_i(0)$ là hiệu ứng điều trị thực (Individual Treatment Effect – ITE), và $\hat{\tau}_i$ là hiệu ứng điều trị được mô hình ước lượng. Ký hiệu n là số lượng cá thể trong tập dữ liệu.

1. PEHE (Precision in Estimation of Heterogeneous Effect).

Chỉ số này đo độ chính xác của mô hình trong việc ước lượng hiệu ứng điều trị dị biệt giữa các cá thể. Giá trị PEHE càng nhỏ chứng tỏ mô hình cá thể hóa tác động nhân quả càng tốt. công thức như sau:

$$\text{PEHE} = \sqrt{\frac{1}{n} \sum_{i=1}^n (\hat{\tau}_i - \tau_i)^2} \quad (1.51)$$

Trong đó: n là số cá thể trong mẫu, $\hat{\tau}_i$ là hiệu ứng điều trị ước lượng của cá thể

i , τ_i : hiệu ứng điều trị thực, $(\hat{\tau}_i - \tau_i)^2$ là sai số bình phương của cá thể i . 2. *Average Treatment Effect (ATE)*. Các chỉ số này đánh giá độ chính xác của mô hình ở mức độ quần thể và kiểm tra thiên lệch hệ thống trong ước lượng. Công thức như sau:

$$ATE = \frac{1}{n} \sum_{i=1}^n \tau_i \quad (1.52)$$

$$\widehat{ATE} = \frac{1}{n} \sum_{i=1}^n \hat{\tau}_i \quad (1.53)$$

a) Sai số ATE được tính toán như sau

$$ATE \text{ Error} = \left| \widehat{ATE} - ATE \right| \quad (1.54)$$

b) Độ lệch ATE (Bias) được tính toán như sau:

$$ATE \text{ Bias} = \widehat{ATE} - ATE \quad (1.55)$$

Trong đó ATE là hiệu ứng điều trị trung bình thực, $\widehat{ATE}$ là hiệu ứng điều trị trung bình ước lượng.

Các chỉ số này đánh giá độ chính xác của mô hình ở mức độ quần thể và kiểm tra thiên lệch hệ thống trong ước lượng.

3. *Hệ số xác định của ITE (R^2)* Chỉ số này đo tỷ lệ phương sai của hiệu ứng thực được mô hình giải thích. có công thức như sau:

$$R^2 = 1 - \frac{\sum_{i=1}^n (\hat{\tau}_i - \tau_i)^2}{\sum_{i=1}^n (\tau_i - \bar{\tau})^2} \quad (1.56)$$

trong đó:

$$\bar{\tau} = \frac{1}{n} \sum_{i=1}^n \tau_i \quad (1.57)$$

Ở đây $\bar{\tau}$ là trung bình hiệu ứng điều trị thực, Tử số là tổng sai số bình phương và mẫu số: tổng phương sai của hiệu ứng thực.

4. *ITE-MAE (Mean Absolute Error)* Chỉ số này đo sai số tuyệt đối trung bình của hiệu ứng cá thể và ít nhạy cảm với ngoại lệ hơn PEHE. Công thức như sau:

$$ITE\text{-}MAE = \frac{1}{n} \sum_{i=1}^n |\hat{\tau}_i - \tau_i| \quad (1.58)$$

Trong đó $|\cdot|$ là giá trị tuyệt đối.

5. *Độ lệch chuẩn của ITE* Chỉ số này phản ánh mức độ phân tán của hiệu ứng điều

trị ước lượng. Công thức như sau:

$$\text{ITE Std} = \sqrt{\frac{1}{n} \sum_{i=1}^n (\hat{\tau}_i - \bar{\hat{\tau}})^2} \quad (1.59)$$

trong đó:

$$\bar{\hat{\tau}} = \frac{1}{n} \sum_{i=1}^n \hat{\tau}_i \quad (1.60)$$

6. *Tỷ lệ bao phủ (Coverage Rate)* Coverage đo tỷ lệ hiệu ứng thực nằm trong khoảng tin cậy dự báo.

$$\text{Coverage} = \frac{1}{n} \sum_{i=1}^n \mathbf{1}\{\tau_i \in \widehat{CI}_i^{95\%}\} \quad (1.61)$$

Trong đó $\widehat{CI}_i^{95\%}$ là khoảng tin cậy 95% của cá thể i , $\mathbf{1}\{\cdot\}$ là hàm chỉ báo.

7. *Calibration Error* Chỉ số này đánh giá mức độ quá tự tin hoặc quá thận trọng của mô hình trong ước lượng bất định. có công thức như sau:

$$\text{Calibration Error} = |\text{Coverage} - 0.95| \quad (1.62)$$

8. *Accuracy*. Chỉ số này đo tỷ lệ dự báo đúng tổng thể. Có công thức như sau:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (1.63)$$

Trong đó, TP (True Positive), TN (True Negative), FP (False Positive) và FN (False Negative) lần lượt là số trường hợp dự báo đúng biến cố, dự báo đúng không biến cố, dự báo nhầm biến cố và bỏ sót biến cố.

Chương 2. DỰ BÁO SỐ CA DƯƠNG TÍNH HIV DỰA TRÊN ĐỒ THỊ CA NHIỄM

Chương này trình bày phương pháp dự báo ngắn hạn số ca nhiễm HIV được đề xuất trong nghiên cứu. Xuất phát từ yêu cầu thực tiễn trong giám sát dịch tễ, dự báo ở cấp khu vực trong thành phố (quận, huyện) có ý nghĩa quan trọng trong hỗ trợ phát hiện sớm các khu vực nguy cơ và phân bổ nguồn lực can thiệp phù hợp. Để đáp ứng yêu cầu này, nghiên cứu tiếp cận dữ liệu giám sát HIV ở cấp ca nhiễm, xây dựng đồ thị ca nhiễm toàn thành phố nhằm mô hình hóa các mối quan hệ không gian, thời gian và dịch tễ học trong quá trình lan truyền của dịch. Trên cơ sở đó, mô hình học động lực lan truyền ở quy mô toàn thành phố và sử dụng thông tin này để dự báo số ca nhiễm tại từng quận, huyện.

2.1. Mô tả bài toán

Xét một thành phố trong đó các ca dương tính HIV được ghi nhận liên tục thông qua hệ thống giám sát. Mỗi ca dương tính được đặc trưng bởi thời điểm phát hiện, thông tin vị trí không gian và một tập các thuộc tính dịch tễ đã được mã hóa. Bài toán đặt ra là dự báo ngắn hạn số ca dương tính HIV mới tại các khu vực không gian liên kề (ví dụ như quận hoặc huyện) trong cùng một thành phố, đồng thời phản ánh được tính liên thông không gian và xu hướng lan truyền tiềm ẩn của dịch ở quy mô toàn thành phố.

Để mô hình hóa các mối tương quan dịch tễ giữa các ca bệnh, trước hết nghiên cứu xây dựng một đồ thị các ca dương tính ở quy mô toàn thành phố tại thời điểm t , ký hiệu là $G^{(t)} = (V^{(t)}, E^{(t)})$. Trong đó, mỗi nút $v \in V^{(t)}$ đại diện cho một ca dương tính đã được ghi nhận, và mỗi cạnh $e \in E^{(t)}$ phản ánh mức độ tương quan tiềm ẩn giữa hai ca bệnh dựa trên các yếu tố không gian, thời gian và đặc trưng dịch tễ.

Khác với các cách tiếp cận dựa trên chuỗi thời gian tổng hợp theo khu vực, phương pháp này xuất phát từ giả định rằng động lực dịch tễ được hình thành từ các mối tương quan giữa các ca dương tính đã được ghi nhận, thay vì chỉ từ sự biến động tổng số ca theo thời gian tại từng khu vực riêng lẻ. Do đó, đồ thị cấp thành phố $G^{(t)}$ được sử dụng để học các biểu diễn không-thời gian phản ánh động lực dịch tễ tổng thể.

Từ đồ thị toàn thành phố, một đồ thị con tương ứng với từng khu vực d được trích xuất, ký hiệu là $G_d^{(t)}$. Đồ thị con này bao gồm các nút thuộc khu vực d và các cạnh liên quan, qua đó phản ánh cấu trúc tương quan cục bộ trong bối cảnh toàn cục.

Bài toán dự báo có thể được biểu diễn dưới dạng học một hàm ánh xạ từ đồ thị con theo thời gian của từng khu vực sang số ca dương tính dự báo trong tương lai, được mô tả như sau:

$$\hat{y}_{d,t+h} = f(G_d^{(t)}) \quad (2.1)$$

trong đó $f(\cdot)$ là hàm dự báo được học bởi mô hình đề xuất; h là khoảng thời gian dự báo trong tương lai (ví dụ 15 hoặc 30 ngày); và $\hat{y}_{d,t+h}$ biểu thị số ca dương tính

HIV mới được dự báo tại khu vực d vào thời điểm $t + h$.

2.2. Phân tích bài toán và cơ sở thiết kế mô hình

Dữ liệu giám sát HIV được thu thập ở mức từng ca bệnh, bao gồm thông tin về thời gian phát hiện, vị trí địa lý và các đặc trưng dịch tễ của mỗi trường hợp. Khác với các phương pháp dự báo truyền thống chỉ sử dụng chuỗi số ca theo thời gian, dữ liệu này còn chứa các mối liên hệ tiềm ẩn giữa các ca bệnh. Những ca được phát hiện gần nhau về không gian, xuất hiện trong khoảng thời gian tương đồng hoặc có các đặc điểm dịch tễ giống nhau có khả năng phản ánh cùng một quá trình lan truyền của dịch. Vì vậy, nếu chỉ tổng hợp dữ liệu theo quận, huyện trước khi xây dựng mô hình thì nhiều thông tin về mối liên hệ giữa các ca bệnh sẽ bị mất.

Xuất phát từ đặc điểm trên, nghiên cứu lựa chọn biểu diễn dữ liệu giám sát dưới dạng một đồ thị ca bệnh ở quy mô toàn thành phố. Trong đồ thị này, mỗi nút tương ứng với một ca dương tính HIV và các cạnh phản ánh mức độ liên hệ giữa các ca bệnh. Việc xây dựng cạnh dựa trên ba giả định mô hình hóa. Thứ nhất, các ca bệnh có vị trí gần nhau được giả định có mức độ liên hệ theo không gian cao hơn. Thứ hai, các ca được phát hiện trong khoảng thời gian gần nhau được giả định có khả năng thuộc cùng một giai đoạn lan truyền của dịch. Thứ ba, các ca có đặc điểm dịch tễ tương đồng, chẳng hạn cùng đường lây truyền, được giả định có mức độ liên hệ cao hơn trong mạng lây nhiễm. Các giả định này không nhằm xác định trực tiếp quan hệ lây truyền giữa hai cá thể mà được sử dụng để xây dựng cấu trúc đồ thị, tạo điều kiện cho mô hình học các mối tương quan dịch tễ tiềm ẩn từ dữ liệu giám sát.

Sau khi đồ thị được xây dựng, bài toán đặt ra là học được biểu diễn của từng ca bệnh sao cho vừa phản ánh được đặc điểm của chính ca bệnh đó, vừa khai thác được thông tin từ các ca bệnh có liên quan. Mạng nơ-ron đồ thị là lựa chọn phù hợp vì có khả năng lan truyền và tổng hợp thông tin trên cấu trúc đồ thị. Trong nghiên cứu này, luận án sử dụng cơ chế chú ý để mô hình tự động học mức độ ảnh hưởng của các nút lân cận thay vì giả định mọi liên kết đều có vai trò như nhau. Nhờ đó, các mối quan hệ có ý nghĩa hơn đối với quá trình lan truyền của dịch sẽ được gán trọng số cao hơn trong quá trình học biểu diễn.

Biểu diễn được học ở quy mô toàn thành phố sau đó được sử dụng để xây dựng các đồ thị con tương ứng với từng quận, huyện. Do các nút trong các đồ thị con đã trải qua quá trình lan truyền thông tin trên toàn bộ đồ thị, biểu diễn của chúng đồng thời chứa thông tin cục bộ của khu vực và thông tin toàn cục của mạng lây nhiễm. Vì vậy, dự báo số ca dương tính HIV tại từng khu vực không chỉ dựa trên dữ liệu của riêng khu vực đó mà còn tận dụng được mối liên hệ với các khu vực khác trong thành phố.

Để giảm độ phức tạp của đồ thị và tập trung vào các nút có ảnh hưởng lớn đến nhiệm vụ dự báo, nghiên cứu sử dụng các phương pháp graph pooling theo hướng chọn lọc nút như TopKPooling hoặc SAGPool. Các phương pháp này giữ lại những nút được mô hình đánh giá là quan trọng và loại bỏ các nút có mức đóng góp thấp, qua đó giảm nhiễu và tạo ra biểu diễn mức đồ thị gọn hơn nhưng vẫn bảo toàn phần lớn thông tin cần thiết cho dự báo.

Từ các phân tích trên, mô hình được thiết kế theo hướng học các mối tương quan dịch tễ ở cấp ca bệnh trên đồ thị toàn thành phố, sau đó chuyển hóa tri thức này thành biểu diễn của từng khu vực để dự báo số ca dương tính HIV mới trong ngắn hạn. Cách tiếp cận này vừa khai thác được thông tin ở quy mô toàn thành phố, vừa bảo tồn được đặc trưng của từng khu vực, tạo cơ sở cho kiến trúc mô hình được trình bày trong mục tiếp theo.

2.3. Biểu diễn đồ thị ca dương tính theo không gian –thời gian cấp thành phố

2.3.1. Khung biểu diễn và động cơ mô hình hóa.

Thay vì mô hình hóa các khu vực hành chính như những thực thể độc lập (Tương tự như các mô hình dự báo dịch bệnh theo không gian-thời gian hiện nay), nghiên cứu này xem các ca dương tính HIV là đơn vị cơ bản của quá trình lan truyền dịch và xây dựng một đồ thị ca dương tính ở cấp thành phố. Cách biểu diễn này nhằm phản ánh trực tiếp các mối tương quan không gian-thời gian tiềm ẩn giữa các ca dương tính được ghi nhận trong hệ thống giám sát, thay vì chỉ khai thác sự biến thiên của số ca tổng hợp theo từng khu vực.

Động cơ của cách tiếp cận này xuất phát từ giả định rằng các đặc điểm dịch tễ quan sát được trong dữ liệu giám sát HIV được hình thành từ sự tương tác gián tiếp và chồng lấn về không gian, thời gian và đường lây truyền giữa các ca dương tính, chứ không chỉ từ các chuỗi số ca theo khu vực. Do đó, việc mô hình hóa ở cấp ca dương tính cho phép bảo toàn tính dị biệt giữa các bản ghi và tạo nền tảng để học các cấu trúc lan truyền tiềm ẩn ở quy mô toàn thành phố.

Định nghĩa đồ thị và đặc trưng nút. Đồ thị ca dương tính được ký hiệu là $G = (V, E)$, trong đó $V = \{v_1, v_2, \dots, v_L\}$ là tập các nút đại diện cho các ca dương tính HIV được ghi nhận trong khoảng thời gian từ t_1 đến t_2 trên toàn thành phố, và $L = |V|$ là tổng số ca dương tính tương ứng.

Mỗi nút $v_i \in V$ được gán một vector đặc trưng $\mathbf{x}_i \in \mathbb{R}^{d_x}$, thuộc ma trận đặc trưng nút $\mathbf{X} \in \mathbb{R}^{L \times d_x}$. Vector đặc trưng này bao gồm các yếu tố dịch tễ đã được mã hóa số, cụ thể gồm sáu nhóm thông tin: nhóm tuổi, giới tính, nghề nghiệp, nhóm nguy cơ, đường lây truyền và khu vực ghi nhận ca dương tính. Trong đó, thông tin khu vực được sử dụng như một thuộc tính của nút, không phải là một nút riêng trong đồ thị. Cách biểu diễn này cho phép mô hình học được sự khác biệt giữa các ca dương tính, đồng thời duy trì khả năng tổng hợp thông tin theo khu vực ở các bước xử lý tiếp theo.

2.3.2. Xây dựng cạnh dựa trên các yếu tố không gian, thời gian và dịch tễ.

Tập cạnh $E \subseteq V \times V$ biểu diễn các mối tương quan đồng thời về không gian, thời gian và đặc điểm dịch tễ giữa các ca dương tính HIV trong phạm vi một thành phố. Mỗi cạnh $e_{ij} \in E$ được gán một vector đặc trưng $\mathbf{f}_{ij} \in \mathbb{R}^{d_f}$, thuộc ma trận đặc trưng cạnh $\mathbf{F} \in \mathbb{R}^{m \times d_f}$, trong đó m là số lượng cạnh trong đồ thị và d_f là số chiều đặc trưng cạnh.

Các đặc trưng cạnh được sử dụng bao gồm: (i) khoảng cách theo kinh độ giữa các phòng xét nghiệm khẳng định HIV, (ii) khoảng cách theo vĩ độ giữa các phòng xét nghiệm khẳng định HIV, và (iii) khoảng cách thời gian giữa thời điểm chẩn đoán

dương tính của hai ca. Việc sử dụng vị trí phòng xét nghiệm khẳng định HIV thay cho vị trí cư trú cá nhân phản ánh đúng bản chất của dữ liệu giám sát dịch tễ, đồng thời đảm bảo tuân thủ các yêu cầu về bảo mật và ẩn danh dữ liệu. Cách tiếp cận này phù hợp với thực tế thu thập dữ liệu HIV, trong đó thông tin địa lý chi tiết ở cấp cá nhân thường không sẵn có hoặc không được phép sử dụng cho mục đích nghiên cứu.

2.3.3. Khái niệm tiềm năng lan truyền không gian–thời gian.

Không phải mọi cặp ca dương tính được ghi nhận trong cùng một thành phố đều có mối liên hệ dịch tễ trực tiếp hoặc gián tiếp. Khả năng hai ca dương tính thuộc cùng một chuỗi lan truyền phụ thuộc đồng thời vào mức độ gần nhau về không gian, thời điểm chẩn đoán và đặc điểm đường lây truyền. Dựa trên nguyên lý này, chương này đưa ra khái niệm *tiềm năng lan truyền không gian–thời gian* nhằm lượng hóa mức độ liên thông dịch tễ tiềm ẩn giữa hai ca dương tính trong dữ liệu giám sát.

Khái niệm này không nhằm suy diễn quan hệ lây nhiễm trực tiếp giữa các cá nhân, mà đóng vai trò như một thước đo tổng hợp để xác định những cặp ca dương tính có mức độ tương quan đủ mạnh để được mô hình hóa thông qua một cạnh trong đồ thị.

Hàm tương quan tổng hợp. Tiềm năng lan truyền không gian–thời gian giữa hai ca dương tính v_i và v_j được lượng hóa thông qua hệ số tương quan

$$\rho_{ij} = \xi_{ij} \times \tau_{ij} \times \kappa_{ij}, \quad \rho_{ij} \in (0, 1]. \quad (2.2)$$

Ba thành phần trong biểu thức trên lần lượt phản ánh tương quan không gian (ξ_{ij}), tương quan thời gian (τ_{ij}) và tương quan về đường lây truyền (κ_{ij}). Việc sử dụng phép nhân phản ánh giả định rằng mối liên hệ dịch tễ giữa hai ca dương tính chỉ có ý nghĩa khi cả ba khía cạnh này đồng thời được thỏa mãn; nếu một trong các thành phần có giá trị thấp, hệ số tương quan tổng thể sẽ giảm mạnh, qua đó loại bỏ các liên kết yếu hoặc kém ý nghĩa.

Thành phần tương quan không gian. Thành phần tương quan không gian ξ_{ij} được xác định dựa trên khoảng cách địa lý giữa các phòng xét nghiệm khẳng định HIV nơi hai ca dương tính được ghi nhận:

$$\xi_{ij} = \begin{cases} 1, & \text{nếu } v_i \text{ và } v_j \text{ được ghi nhận trong cùng khu vực,} \\ 1 - \frac{d_{ij}}{\max(d_{ij})}, & \text{ngược lại,} \end{cases} \quad (2.3)$$

trong đó d_{ij} là khoảng cách Euclid giữa hai phòng xét nghiệm, được tính từ tọa độ kinh độ và vĩ độ:

$$d_{ij} = \sqrt{(\text{long}_{CT_i} - \text{long}_{CT_j})^2 + (\text{lat}_{CT_i} - \text{lat}_{CT_j})^2}. \quad (2.4)$$

Cách chuẩn hóa tuyến tính này đảm bảo $\xi_{ij} \in (0, 1]$ và duy trì mối quan hệ đơn điệu giữa khoảng cách địa lý và mức độ tương quan, đồng thời tránh khuếch đại quá mức ảnh hưởng của khoảng cách như trong các hàm suy giảm mũ.

Thành phần tương quan thời gian. Thành phần tương quan thời gian τ_{ij} phản ánh

mức độ gần nhau về thời điểm chẩn đoán dương tính của hai ca:

$$\tau_{ij} = \begin{cases} 1, & \text{nếu } r_i = r_j, \\ 1 - \frac{|r_i - r_j|}{\max(r_{ij})}, & \text{ngược lại,} \end{cases} \quad (2.5)$$

trong đó r_i và r_j lần lượt là thời điểm chẩn đoán dương tính của hai ca dương tính. Thành phần này phản ánh giả định rằng các ca được phát hiện gần nhau về thời gian có nhiều khả năng chịu ảnh hưởng từ các bối cảnh lan truyền tương đồng.

Thành phần tương quan đường lây truyền. Thành phần tương quan về đường lây truyền κ_{ij} được xác định như sau:

$$\kappa_{ij} = \begin{cases} 1, & \text{nếu } v_i \text{ và } v_j \text{ có cùng đường lây truyền,} \\ \frac{L_i + L_j}{L}, & \text{ngược lại,} \end{cases} \quad (2.6)$$

trong đó L_i và L_j là số lượng ca dương tính có cùng đường lây truyền với v_i và v_j , còn L là tổng số ca dương tính trong toàn bộ tập dữ liệu. Cách xây dựng này cho phép phản ánh xác suất nền của các đường lây truyền phổ biến trong dữ liệu giám sát, giúp mô hình thích ứng với sự thay đổi cấu trúc dịch tễ giữa các giai đoạn khác nhau.

Ngưỡng tạo cạnh và kiểm soát mật độ đồ thị Một cạnh $e_{ij} \in E$ chỉ được tạo nếu hệ số tương quan thỏa mãn điều kiện

$$\rho_{ij} > \beta, \quad (2.7)$$

trong đó $\beta \in (0, 1)$ là siêu tham số điều khiển mức độ chặt chẽ của kết nối trong đồ thị. Trong đồ thị được xây dựng, ngưỡng β được sử dụng để quyết định việc hình thành cạnh giữa hai nút, mỗi nút tương ứng với một ca dương tính HIV được ghi nhận. Một cạnh chỉ được tạo ra khi hệ số tương quan, được xác định từ khoảng cách không gian, thời điểm chẩn đoán và sự tương đồng về đường lây truyền, vượt qua giá trị ngưỡng đã định.

Việc áp dụng ngưỡng β nhằm loại bỏ các liên kết giữa những ca dương tính không đủ gần nhau về mặt dịch tễ, từ đó giảm nhiễu trong đồ thị. Đồng thời, cơ chế này giúp kiểm soát mật độ kết nối, tránh việc mỗi nút liên thông với quá nhiều nút khác, khiến thông tin lan truyền trở nên không chọn lọc và làm suy giảm khả năng học các cấu trúc lan truyền có ý nghĩa. Nhờ đó, đồ thị duy trì được khả năng học và phản ánh tốt hơn động lực lan truyền HIV tiềm ẩn trong dữ liệu giám sát. Giá trị $\beta=0.5$ được lựa chọn sau quá trình thử nghiệm thực nghiệm nhằm đạt được cấu trúc đồ thị đủ thừa để đảm bảo hiệu quả tính toán nhưng vẫn duy trì khả năng biểu diễn các tương quan dịch tễ cần thiết cho dự báo.

2.3.4. Nhúng đặc trưng nút và cạnh về cùng không gian biểu diễn

Trong đồ thị ca bệnh, đặc trưng nút x_i (mô tả thuộc tính dịch tễ đã mã hóa của từng ca bệnh) và đặc trưng cạnh f_{ij} (mô tả liên thông không-thời gian và/hoặc dịch tễ giữa hai ca bệnh) có thể khác nhau về kiểu dữ liệu (rời rạc/liên tục) và số chiều. Để tích

hợp đồng thời hai nguồn thông tin này trong các khối chú ý trên đồ thị, nghiên cứu sử dụng hàm MLP để ánh xạ chúng về cùng không gian nhưng có chiều d_{em} . Cụ thể:

$$\text{FeatureEMBED}(v_i) = \text{MLP}(x_i, d_{em}), \quad (2.8)$$

$$\text{FeatureEMBED}(e_{ij}) = \text{MLP}(f_{ij}, d_{em}). \quad (2.9)$$

Trong đó, $x_i \in \mathbb{R}^{d_x}$ là đặc trưng dịch tễ của nút v_i , còn $f_{ij} \in \mathbb{R}^{d_f}$ là đặc trưng không-thời gian (và/hoặc dịch tễ) của cạnh e_{ij} . Bước nhúng này đảm bảo tính nhất quán hình học khi tích hợp thông tin nút và cạnh trong các khối GNN/chú ý ở các bước sau, đồng thời giúp mô hình học được các biểu diễn liên tục phù hợp cho bài toán hồi quy dự báo.

2.4. Mô hình dự báo

2.4.1. Tổng quan mô hình đề xuất

Cho đồ thị ca bệnh cấp đô thị $G = (V, E)$, trong đó V là tập các ca dương tính HIV, E là tập các cạnh biểu diễn mối liên hệ không gian-thời gian giữa các ca bệnh, $X^{(0)} \in \mathbb{R}^{N \times d_x}$ là ma trận đặc trưng nút và $F^{(0)} \in \mathbb{R}^{|E| \times d_e}$ là ma trận đặc trưng cạnh. Mục tiêu của mô hình là học một ánh xạ:

$$f : (G, X^{(0)}, F^{(0)}) \rightarrow \hat{\mathbf{y}}, \quad (2.10)$$

trong đó $\hat{\mathbf{y}} = (\hat{y}_1, \hat{y}_2, \dots, \hat{y}_n) \in \mathbb{R}^n$ là số ca dương tính HIV mới được dự báo cho n khu vực (quận, huyện) trong kỳ dự báo tiếp theo.

Kiến trúc đề xuất gồm hai giai đoạn học biểu diễn và một tầng dự báo. Trước hết, mạng nơ-ron đồ thị dựa trên chú ý được sử dụng để học biểu diễn các ca bệnh trên đồ thị toàn thành phố như sau:

$$H = \text{AttGNN}(G, X^{(0)}, F^{(0)}), \quad H \in \mathbb{R}^{N \times d_{em}}, \quad (2.11)$$

trong đó mỗi hàng của ma trận H là biểu diễn của một ca bệnh sau khi đã tổng hợp thông tin từ các ca bệnh lân cận thông qua cơ chế chú ý. AttGNN là mạng đồ thị dựa trên cơ chế chú ý.

Tiếp theo, với mỗi khu vực $r \in \{1, \dots, n\}$, một đồ thị con $G_r = (V_r, E_r)$ được trích xuất từ đồ thị toàn thành phố và ma trận biểu diễn nút tương ứng được xác định bởi $H_r = H[V_r]$. Đồ thị con sau đó được rút gọn bằng phương pháp gộp đồ thị theo hướng sàng lọc nút (*node drop pooling*):

$$(G'_r, H'_r) = \text{Pool}(G_r, H_r), \quad (2.12)$$

nhằm giữ lại các ca bệnh mang nhiều thông tin nhất cho dự báo. Trên đồ thị đã được rút gọn, một lớp mạng nơ-ron đồ thị dựa trên chú ý tiếp tục được áp dụng để học biểu diễn ở cấp khu vực như sau:

$$H''_r = \text{AttGNN}(G'_r, H'_r). \quad (2.13)$$

Tiếp theo, phép tổng hợp toàn cục (*Global Mean Pooling*) được sử dụng để ánh xạ mỗi đồ thị con thành một vector biểu diễn mức khu vực như sau:

$$z_r = \text{Mean}(H_r''), \quad z_r \in \mathbb{R}^{d_{em}}. \quad (2.14)$$

Các vector biểu diễn của tất cả các khu vực được xếp theo hàng (*row-wise stacking*) để tạo thành ma trận biểu diễn khu vực như sau:

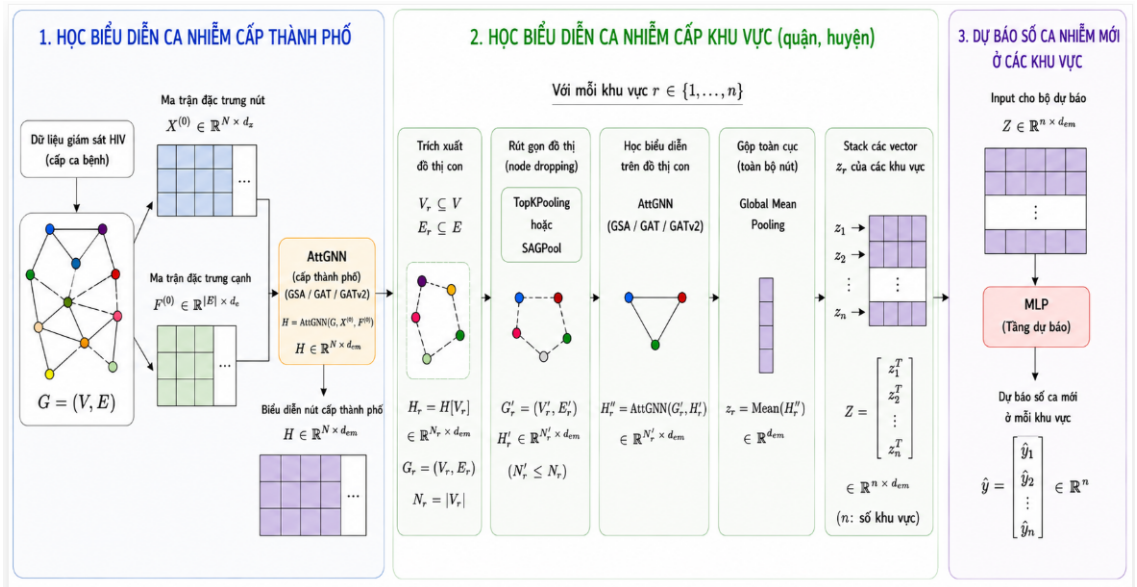
$$Z = \begin{bmatrix} z_1 \\ z_2 \\ \vdots \\ z_n \end{bmatrix} \in \mathbb{R}^{n \times d_{em}}, \quad (2.15)$$

trong đó mỗi hàng của ma trận Z là vector biểu diễn của một khu vực và mỗi cột tương ứng với một chiều của không gian biểu diễn d_{em} . Do đó, Z là ma trận đặc trưng mức khu vực, tổng hợp thông tin của toàn bộ các khu vực trong đô thị và được sử dụng làm đầu vào của tầng dự báo.

Cuối cùng, mạng MLP thực hiện bài toán hồi quy để dự báo đồng thời số ca dương tính HIV mới của tất cả các khu vực

$$\hat{y} = \text{MLP}(Z). \quad (2.16)$$

Kiến trúc tổng quát của mô hình được minh họa trong Hình 2.1. Kiến trúc này cho phép chuyển hóa tri thức lan truyền dịch tế được học ở quy mô toàn thành phố thành biểu diễn đặc trưng của từng khu vực, từ đó hỗ trợ dự báo đồng thời số ca dương tính HIV mới trong ngắn hạn cho tất cả các khu vực một cách thống nhất và nhất quán.



Hình 2.1: mô hình dự báo số ca dương tính HIV mới ở cấp quận/huyện từ dữ liệu giám sát ca bệnh. Mô hình học biểu diễn ca bệnh trên đồ thị toàn thành phố bằng mạng nơ-ron đồ thị dựa trên chú ý, trích xuất các đồ thị con theo khu vực, thực hiện gộp đồ thị và tổng hợp biểu diễn mức khu vực để dự báo số ca mới.

2.4.2. Mạng nơ-ron đồ thị dựa trên chú ý cho đồ thị đô thị

Nghiên cứu khảo sát ba kiến trúc GNN dựa trên chú ý cho đồ thị đô thị: (i) Graph with Self-Attention (GSA), (ii) Graph Attention Network (GAT), và (iii) Graph Attention Network v2 (GATv2). Điểm chung của các kiến trúc này là cho phép mô hình tự học trọng số ảnh hưởng giữa các ca bệnh lân cận dựa trên dữ liệu, thay vì gán trọng số cố định. Nhờ đó, mô hình có thể thích nghi với cấu trúc liên thông thay đổi theo giai đoạn dịch và theo bối cảnh đô thị.

1. *Đồ thị kết hợp cơ chế tự chú ý (Graph with Self-Attention - GSA).* Cơ chế tự chú ý trong kiến trúc Transformer mô hình hóa các mối quan hệ phụ thuộc giữa các phần tử trong một chuỗi, bằng cách cho phép mỗi phần tử tự động gán trọng số cho các phần tử còn lại dựa trên biểu diễn đặc trưng của chúng. cơ chế tự chú ý cho phép nắm bắt các mối quan hệ dài hạn và không tuần tự một cách trực tiếp thông qua phép so khớp trong không gian đặc trưng.

Trong bài toán dự báo lan truyền HIV trên dữ liệu giám sát, mỗi nút trong đồ thị đại diện cho một ca dương tính được ghi nhận, còn các cạnh phản ánh mức độ gần nhau về không gian, thời gian và đặc điểm dịch tễ. Tuy nhiên, sự tồn tại của một cạnh không đồng nghĩa với việc hai ca dương tính có vai trò tương đương trong động lực lan truyền dịch. Trên thực tế, mức độ ảnh hưởng tiềm ẩn của mỗi ca dương tính đối với các ca khác có thể khác nhau, tùy thuộc vào bối cảnh không gian, thời gian và các yếu tố dịch tễ liên quan.

Việc áp dụng cơ chế tự chú ý trên đồ thị cho phép mô hình học một cách mềm (soft) mức độ ảnh hưởng lan truyền tiềm ẩn giữa các ca dương tính, thông qua việc gán các trọng số chú ý khác nhau cho từng cặp nút được kết nối. Các trọng số này phản ánh mức độ đóng góp tương đối của thông tin từ các ca dương tính lân cận trong quá trình cập nhật biểu diễn của ca đang xét, thay vì giả định sự đóng góp đồng đều giữa các nút.

Do đó, GSA đặc biệt phù hợp với bài toán dự báo HIV ngắn hạn, khi mục tiêu là mô hình hóa sự không đồng nhất trong mức độ ảnh hưởng lan truyền giữa các ca dương tính, đồng thời khai thác hiệu quả cấu trúc đồ thị không gian–thời gian được xây dựng từ dữ liệu giám sát thực tế. Cách tiếp cận này cho phép biểu diễn nút phản ánh động lực lan truyền dịch ở quy mô toàn thành phố, thay vì chỉ phản ánh các đặc trưng tĩnh của từng ca dương tính riêng lẻ.

Cơ chế tự chú ý nhận đầu vào $D = [d_1, \dots, d_n] \in \mathbb{R}^{n \times d_{em}}$, và tạo ma trận trọng số chú ý:

$$SCORE_{att} = \text{softmax} \left(\frac{QK^T}{\sqrt{d_k}} \right), \quad (2.17)$$

trong đó $Q = DW_q^T$, $K = DW_k^T$, với $W_q, W_k \in \mathbb{R}^{d_k \times d_{em}}$ là tham số học được. Đầu ra của attention:

$$Y = SCORE_{att} \cdot V, \quad (2.18)$$

với $V = DW_v^T$, $W_v \in \mathbb{R}^{d_v \times d_{em}}$.

Multi-head attention kết hợp C head để nắm bắt nhiều kiểu quan hệ:

$$\text{Multihead}(Q, K, V) = \text{CONCAT}(Y^1, \dots, Y^C)W_o, \quad (2.19)$$

trong đó $W_o \in \mathbb{R}^{C \cdot y \times y}$.

2. *Tích hợp đặc trưng cạnh vào trọng số chú ý trên đồ thị.* Để phản ánh tương tác không gian–thời gian giữa hai ca bệnh, đặc trưng cạnh được đưa trực tiếp vào phép tính chú ý. Trọng số chú ý giữa hai nút v_i và v_j tại head c được định nghĩa:

$$\alpha_{c,ij} = \frac{(q_{c,i}, k_{c,j} + f_{c,ij})}{\sum_{u \in \mathcal{N}(i)} (q_{c,i}, k_{c,u} + f_{c,iu})}, \quad (2.20)$$

với

$$q_{c,i} = W_{c,q}h_i + b_{c,q}, \quad k_{c,j} = W_{c,k}h_j + b_{c,k}, \quad f_{c,ij} = W_{c,e}f_{ij} + b_{c,e},$$

và

$$(q, k) = \exp\left(\frac{q^\top k}{\sqrt{d_k}}\right).$$

Ở đây $h_i, h_j \in \mathbb{R}^{d_{em}}$ là biểu diễn nút (sau nhúng/hoặc sau các lớp GNN trước đó), $f_{ij} \in \mathbb{R}^{d_{em}}$ là biểu diễn cạnh sau nhúng, $W_{c,q}, W_{c,k} \in \mathbb{R}^{d_k \times d_{em}}$, $W_{c,e} \in \mathbb{R}^{d_k \times d_{em}}$ là tham số học được.

Sau khi có $\alpha_{c,ij}$, biểu diễn nút v_i được cập nhật:

$$h'_i = \parallel_{c=1}^C \left[\sum_{j \in \mathcal{N}(i)} \alpha_{c,ij} (v_{c,j} + e_{c,ij}) \right], \quad (2.21)$$

trong đó

$$v_{c,j} = W_{c,v}h_j + b_{c,v}. \quad (2.22)$$

Việc đưa $f_{c,ij}$ vào thành phần $(k_{c,j} + f_{c,ij})$ cho phép trọng số chú ý không chỉ phụ thuộc vào “nội dung” của ca bệnh lân cận v_j , mà còn phụ thuộc vào mức liên thông không gian–thời gian của cặp (i, j) . Điều này phù hợp với bối cảnh giám sát HIV, trong đó quan hệ giữa hai ca bệnh cần được diễn giải theo khoảng cách không gian–thời gian thay vì chỉ dựa trên tương đồng dịch tễ của từng ca riêng rẽ.

3. *Graph Attention Network (GAT) có mở rộng theo đặc trưng cạnh.* Trong GAT, trọng số chú ý giữa nút đích v_i và nút lân cận v_j (bao gồm cả i) được tính dựa trên đặc trưng nút và đặc trưng cạnh. Ký hiệu $h_i = W_i x_i$, $h_j = W_j x_j$, và $h_{ij} = W_{ij} e_{ij}$. Khi đó hệ số chú ý được tính như sau.

$$\alpha_{ij} = \frac{\exp(\text{LeakyReLU}(W_{a,i}^\top h_i + W_{a,j}^\top h_j + W_{a,ij}^\top h_{ij}))}{\sum_{k \in \mathcal{N}(i) \cup \{i\}} \exp(\text{LeakyReLU}(W_{a,i}^\top h_i + W_{a,k}^\top h_k + W_{a,ik}^\top h_{ik}))}, \quad (2.23)$$

và với multi-head attention, biểu diễn cập nhật:

$$h'_i = \frac{1}{C} \sum_{c=1}^C \sum_{j \in \mathcal{N}(i) \cup \{i\}} \alpha_{ij}^{(c)} h_j. \quad (2.24)$$

4. Graph Attention Network v2 (GATv2)

GATv2 tăng tính linh hoạt của attention bằng cách thay đổi thứ tự phép biến đổi bằng cách áp dụng phi tuyến trước khi chiếu bởi trọng số attention với công thức như sau.

$$\alpha_{ij} = \frac{\exp(\mathcal{W}_a^\top \text{LeakyReLU}(h_i + h_j + h_{ij}))}{\sum_{k \in \mathcal{N}(i) \cup \{i\}} \exp(\mathcal{W}_a^\top \text{LeakyReLU}(h_i + h_k + h_{ik}))}. \quad (2.25)$$

GATv2 cho phép hàm chú ý “thích nghi” hơn theo ngữ cảnh lân cận, giảm tính “tĩnh” của attention trong GAT cổ điển, do đó phù hợp khi cấu trúc liên thông giữa ca bệnh có thể thay đổi theo giai đoạn dịch và theo bối cảnh đô thị.

2.4.3. Trích xuất đồ thị con theo khu vực từ đồ thị ca dương tính cấp thành phố

Mục tiêu dự báo của nghiên cứu là số ca dương tính HIV mới ở cấp khu vực (ví dụ như quận, huyện). Tuy nhiên, các biểu diễn ca dương tính đã được học trước đó nằm trên đồ thị cấp thành phố, phản ánh đầy đủ các tương tác liên khu vực. Do đó, để chuyển tri thức từ cấp thành phố sang cấp khu vực mà không làm mất thông tin lan truyền toàn cục, nghiên cứu tiến hành trích xuất các đồ thị con tương ứng với từng khu vực từ đồ thị đô thị ban đầu.

Cụ thể, với mỗi khu vực $r \in \{1, 2, \dots, n\}$, một đồ thị con

$$G_r = (V_r, E_r)$$

được xây dựng, trong đó $V_r = \{v_{r1}, \dots, v_{rL_r}\}$ là tập các ca dương tính HIV được ghi nhận trong khu vực r , và $L_r = |V_r|$. Các biểu diễn nút trong G_r được kế thừa trực tiếp từ embedding đã được học trên đồ thị cấp thành phố. Nhờ đó, mỗi ca dương tính trong đồ thị con vẫn mang thông tin về các tương tác liên khu vực, thay vì chỉ phản ánh cấu trúc nội khu vực.

Tập cạnh E_r của đồ thị con được tái thiết lập theo cùng quy tắc tạo cạnh đồ thị đã trình bày ở phần trên, tức là chỉ tạo cạnh giữa hai ca dương tính nếu hệ số tương quan dịch tễ $\rho_{ij} > \beta$. Cách làm này đảm bảo tính nhất quán giữa đồ thị cấp thành phố và các đồ thị con theo khu vực, đồng thời kiểm soát mật độ kết nối trong từng đồ thị con, giúp duy trì khả năng học của mô hình ở các bước tiếp theo.

2.4.4. Gộp đồ thị (Graph pooling) cho đồ thị con theo khu vực

Bài toán dự báo số ca dương tính HIV mới cho mỗi khu vực có thể được xem là một bài toán hồi quy ở mức đồ thị (*graph-level regression*), trong đó mỗi đồ thị con G_r cần được ánh xạ thành một biểu diễn cố định có ý nghĩa dịch tễ. Phương pháp gộp đồ thị đóng vai trò quan trọng trong việc giảm kích thước đồ thị, loại bỏ nhiễu và tập trung vào các ca dương tính có ảnh hưởng lớn đến động lực lan truyền tại khu vực tương

ứng.

Trong nghiên cứu này, các kỹ thuật gộp đồ thị theo hướng loại bỏ các nút kém ảnh hưởng trong đồ thị (node dropping) được ưu tiên sử dụng thay vì các kỹ thuật gom cụm các nút (node clustering). Lý do là việc gom cụm các ca dương tính có thể làm mất các đặc trưng dịch tễ hiếm hoặc các mối liên thông không gian–thời gian quan trọng, vốn có ý nghĩa lớn trong giám sát HIV. Ngược lại, node dropping cho phép mô hình chọn lọc một tập con các ca dương tính tiêu biểu, trong khi vẫn bảo toàn cấu trúc đồ thị ban đầu.

Quy trình pooling bao gồm ba bước chính: (i) gán điểm quan trọng cho mỗi nút; (ii) lựa chọn một tỷ lệ k các nút có điểm cao nhất; (iii) tái thiết lập đồ thị với tập nút được chọn.

1. Phương pháp TopKPooling. Với đồ thị con $G_r = (X_r, A_r)$, TopKPooling sử dụng một vector chiều học được $p \in \mathbb{R}^{d_x}$ để tính điểm cho mỗi nút:

$$S_r = \frac{X_r p}{\|p\|}.$$

Các nút có điểm cao nhất được lựa chọn:

$$IDX_r = \text{RANK}(S_r, k).$$

Sau đó, điểm được chuẩn hóa bằng hàm sigmoid:

$$\tilde{S}_r = \text{sigmoid}(S_r(IDX_r)), \quad (2.26)$$

và đặc trưng nút được tái trọng số:

$$X'_r = X_r[IDX_r, :] \odot \left(\tilde{S}_r \mathbf{1}_{d_x}^\top \right), \quad (2.27)$$

trong khi ma trận kề được cập nhật tương ứng:

$$A'_r = A_r[IDX_r, IDX_r].$$

2. Phương pháp SAGPool.

Khác với TopKPooling, SAGPool tính điểm nút dựa trên cả đặc trưng nút và cấu trúc đồ thị thông qua một lớp GCN:

$$S_r = \sigma \left(\tilde{D}_r^{-\frac{1}{2}} \tilde{A}_r \tilde{D}_r^{-\frac{1}{2}} X_r W_{\text{SAPool}} \right), \quad (2.28)$$

trong đó $\tilde{A}_r = A_r + I_r$ là ma trận kề có bổ sung self-loop, $\tilde{D}_r$ là ma trận bậc, và $W_{\text{SAPool}} \in \mathbb{R}^{d_x \times 1}$ là tham số học được. Một tỷ lệ $\lfloor k L_r \rfloor$ các nút có điểm cao nhất được giữ lại để xây dựng đồ thị mới G'_r .

2.4.5. Cập nhật biểu diễn và tổng hợp thông tin mức đồ thị con

Sau khi áp dụng node drop pooling, các đồ thị con đã được rút gọn vẫn giữ lại những ca dương tính quan trọng nhất về mặt dịch tễ. Để tiếp tục khai thác cấu trúc liên thông nội khu vực, nghiên cứu áp dụng thêm một lớp mạng đồ thị chú ý cho mỗi đồ thị con.

Bước này cho phép mô hình tái tổng hợp thông tin cục bộ giữa các ca dương tính đã được chọn lọc, từ đó tăng cường khả năng biểu diễn động lực lan truyền trong từng khu vực.

Tiếp theo, để thu được một biểu diễn cố định cho mỗi đồ thị con G_r , nghiên cứu sử dụng phép tổng hợp toàn cục (global mean pooling):

$$\text{output}(G_r) = \frac{1}{L_r} \sum_{i=1}^{L_r} x_i, \quad \forall r \in \{1, 2, \dots, n\}. \quad (2.29)$$

Phép tổng hợp này giúp biểu diễn mức đồ thị phản ánh xu hướng chung của các ca dương tính trong khu vực, đồng thời giảm ảnh hưởng của các biến động cục bộ không ổn định.

2.4.6. Tầng dự báo và đầu ra của mô hình

Các biểu diễn mức đồ thị của tất cả các khu vực được ghép lại thành một ma trận:

$$\text{output}(G) \in \mathbb{R}^{n \times d_{em}}.$$

Ma trận này sau đó được đưa qua một mạng MLP để thực hiện hồi quy và dự báo số ca dương tính HIV mới cho từng khu vực:

$$\text{PredictOutput} = \text{MLP}(\text{output}(G), 1). \quad (2.30)$$

2.4.7. Lựa chọn cấu hình mô hình tối ưu

Để đánh giá tính hiệu quả và lựa chọn cấu hình phù hợp nhất cho bài toán dự báo HIV ngắn hạn, nghiên cứu tiến hành các thí nghiệm so sánh theo hai chiều chính. Thứ nhất là mạng đồ thị dựa trên cơ chế chú ý được sử dụng ở cấp thành phố, bao gồm GSA, GAT và GATv2. Thứ hai là kỹ thuật graph pooling theo hướng node dropping, bao gồm TopKPooling và SAGPool, được áp dụng cho các đồ thị con theo khu vực.

Việc lựa chọn cấu hình cuối cùng dựa trên hiệu năng dự báo trên dữ liệu thực tế, cũng như tính ổn định trong quá trình huấn luyện. Kết quả chi tiết của các thí nghiệm so sánh và phân tích ảnh hưởng của từng thành phần được trình bày trong phần thực nghiệm và kết quả của chương.

2.4.8. Quy trình huấn luyện mô hình

Dựa trên kiến trúc tổng thể được minh họa trong Hình 2.1, quá trình huấn luyện mô hình được thực hiện theo hai giai đoạn chính, bao gồm học biểu diễn trên đồ thị toàn cục và học dự báo trên các đồ thị con theo khu vực.

Ở giai đoạn thứ nhất, mô hình nhận đầu vào là đồ thị ca bệnh cấp đô thị và thực hiện lan truyền thông tin giữa các nút thông qua các lớp mạng nơ-ron đồ thị dựa trên cơ chế chú ý nhằm học biểu diễn đặc trưng toàn cục. Quá trình này cho phép mô hình nắm bắt các tương quan lan truyền tiềm ẩn giữa các ca bệnh trong toàn hệ thống.

Ở giai đoạn thứ hai, các đồ thị con theo từng khu vực được trích xuất từ đồ thị toàn cục và tiếp tục được xử lý thông qua các khối lọc nút, mạng đồ thị chú ý và phép tổng

hợp để thu được biểu diễn ở cấp khu vực. Các biểu diễn này sau đó được đưa vào lớp MLP để sinh ra giá trị dự báo số ca nhiễm tương ứng.

Trong quá trình huấn luyện, đầu ra của mô hình được so sánh với giá trị thực tế để tính toán hàm mất mát. Thuật toán lan truyền ngược được sử dụng để cập nhật đồng thời các tham số của toàn bộ mô hình thông qua bộ tối ưu Adam. Quá trình này được lặp lại qua nhiều epoch cho đến khi hội tụ. Đồng thời, hiệu năng mô hình được theo dõi trên tập xác thực nhằm lựa chọn mô hình tối ưu và hạn chế hiện tượng quá khớp.

2.5. Thực nghiệm và kết quả

Phần này trình bày các thực nghiệm đánh giá hiệu quả của mô hình đề xuất trong bài toán dự báo ngắn hạn số ca dương tính HIV ở cấp quận/huyện trên dữ liệu giám sát dịch tễ thực tế. Bên cạnh việc so sánh hiệu năng dự báo với các phương pháp đối sánh, các thực nghiệm còn được thiết kế nhằm phân tích đóng góp của từng thành phần trong kiến trúc mô hình, đồng thời đánh giá khả năng dự báo của mô hình ở các khoảng thời gian dự báo khác nhau.

2.5.1. Dữ liệu và tiền xử lý dữ liệu

Nghiên cứu sử dụng bộ dữ liệu hồi cứu về giám sát ca nhiễm HIV tại Thành phố Hồ Chí Minh trong giai đoạn từ tháng 01 năm 2008 đến tháng 12 năm 2019. Bộ dữ liệu bao gồm 40.611 bản ghi tương ứng với các ca nhiễm HIV được ghi nhận tại 23 quận, huyện tại thời điểm đó của thành phố. Các biến được sử dụng trong nghiên cứu bao gồm các giá trị số biểu diễn các yếu tố dịch tễ học của ca nhiễm HIV, cụ thể gồm:

1. Giá trị số đại diện cho các nhóm tuổi.
2. Giá trị số đại diện cho giới tính
3. Giá trị số đại diện cho nghề nghiệp.
4. Giá trị số đại diện cho nhóm nguy cơ.
5. Giá trị số đại diện cho đường lây truyền.
6. Giá trị số đại diện cho địa điểm phát hiện ca dương tính (quận/huyện).
7. Giá trị tọa độ kinh độ và vĩ độ của các cơ sở xét nghiệm HIV tương ứng với từng quận.
8. Khoảng thời gian giữa các ca nhiễm được chuẩn hóa theo các khoảng thời gian ba tháng.

Để phục vụ cho việc mô hình hóa và dự báo, bộ dữ liệu được chia thành các khoảng thời gian ba tháng. Trong mỗi khoảng thời gian, một đồ thị HIV toàn thành phố được xây dựng bằng thư viện Torch Geometric. Trong đồ thị này, mỗi ca nhiễm HIV được biểu diễn như một nút (node) của đồ thị.

Các đặc trưng của nút bao gồm các yếu tố dịch tễ học như nhóm tuổi, giới tính, nghề nghiệp, địa điểm phát hiện ca bệnh, nhóm nguy cơ và đường lây truyền. Trong khi đó, đặc trưng của các cạnh giữa hai nút được xác định dựa trên mối quan hệ không – thời gian giữa các ca bệnh, bao gồm khoảng cách địa lý giữa hai cơ sở xét nghiệm (tính từ kinh độ và vĩ độ) và khoảng thời gian giữa hai lần xác nhận dương tính.

Dữ liệu sử dụng trong nghiên cứu là dữ liệu hồi cứu đã được mã hóa hoàn toàn trước khi được sử dụng cho nghiên cứu. Bộ dữ liệu không chứa bất kỳ thông tin định

danh cá nhân trực tiếp nào như họ tên, địa chỉ, số định danh cá nhân hoặc thông tin liên hệ. Tất cả các biến trong bộ dữ liệu được biểu diễn dưới dạng mã số và không tồn tại khóa ánh xạ cho phép truy ngược lại danh tính hoặc thông tin gốc của đối tượng. Dữ liệu chỉ được sử dụng cho mục đích nghiên cứu nhằm bảo đảm các yêu cầu về bảo mật thông tin và quyền riêng tư dữ liệu.

2.5.2. Chiến lược chia dữ liệu

Dữ liệu trong nghiên cứu được thu thập theo thời gian từ tháng 01 năm 2008 đến tháng 12 năm 2019, tương ứng với 144 tháng quan sát. Để xây dựng dữ liệu đầu vào cho mô hình, nghiên cứu sử dụng cửa sổ thời gian trượt dài 3 tháng. Cụ thể, tại mỗi bước, dữ liệu của ba tháng liên tiếp được sử dụng để xây dựng một đồ thị HIV toàn thành phố. Sau đó, cửa sổ được dịch tiếp 1 tháng để tạo đồ thị kế tiếp. Với cách xây dựng này, từ 144 tháng dữ liệu ban đầu, nghiên cứu thu được 141 đồ thị theo thời gian.

Sau khi xây dựng toàn bộ chuỗi đồ thị, dữ liệu được chia thành hai phần: 80% số đồ thị được sử dụng cho huấn luyện và 20% số đồ thị còn lại được sử dụng cho kiểm tra. Cách chia này giúp mô hình học được quy luật lan truyền từ phần lớn dữ liệu lịch sử, đồng thời vẫn giữ lại một tập dữ liệu độc lập để đánh giá khả năng dự báo.

2.5.3. Chỉ số đánh giá

Hiệu năng dự báo của các mô hình được đánh giá thông qua hai chỉ số gồm sai số bình phương trung bình (Mean Squared Error – MSE) và hệ số xác định (R^2). MSE được sử dụng để đánh giá mức độ sai lệch giữa giá trị dự báo và giá trị thực tế, đồng thời nhấn mạnh các sai số lớn thông qua phép bình phương, nên phù hợp để đánh giá độ chính xác của mô hình dự báo. Trong khi đó, R^2 phản ánh tỷ lệ biến thiên của dữ liệu được mô hình giải thích, qua đó đánh giá khả năng nắm bắt xu hướng biến động của số ca HIV theo thời gian. Hai chỉ số này là các thước đo phổ biến và được sử dụng rộng rãi trong đánh giá các mô hình hồi quy và dự báo [86, 87, 88].

2.5.4. Các mô hình so sánh

Các mô hình so sánh. Để đánh giá toàn diện hiệu quả của mô hình đề xuất, nghiên cứu tiến hành so sánh với ba nhóm phương pháp phổ biến trong dự báo dịch tễ. Nhóm thứ nhất bao gồm các mô hình chuỗi thời gian truyền thống, đại diện là ARIMA. Nhóm thứ hai bao gồm các mô hình học sâu cổ điển dựa trên chuỗi thời gian như LSTM [89], BI-LSTM [42], CNN [46], Stacked-LSTM [44] và CNN-LSTM [48]. Nhóm thứ ba bao gồm các mô hình dự báo đồ thị không gian–thời gian đã được công bố trong các nghiên cứu trước đây, bao gồm DCRNN [90], MPNN-LSTM [91], GAT-LSTM [92], GCN-LSTM [93] và CGN-GRU [94].

Các mô hình so sánh đều được huấn luyện và đánh giá trên cùng tập dữ liệu, sử dụng cùng cửa sổ thời gian và chỉ số đánh giá, bảo tính công bằng trong so sánh.

2.5.5. Thiết lập thực nghiệm

Các tham số huấn luyện chính của mô hình được cố định trong toàn bộ quá trình thực nghiệm và được trình bày trong Bảng 2.1. Các tham số huấn luyện được lựa chọn dựa

trên kinh nghiệm thực nghiệm với các mô hình học sâu và được điều chỉnh trong quá trình huấn luyện để đạt được sự hội tụ ổn định. Adam được sử dụng làm phương pháp tối ưu do khả năng cập nhật tham số hiệu quả và được sử dụng phổ biến trong các bài toán học sâu. Learning rate ban đầu được thiết lập ở mức 0.0005 để tránh dao động mạnh trong quá trình tối ưu, kết hợp với cơ chế CosineAnnealingWarmRestarts để điều chỉnh learning rate theo từng giai đoạn huấn luyện. Weight decay được sử dụng để hạn chế hiện tượng overfitting. Batch size được thiết lập ở mức 4 do giới hạn bộ nhớ khi xử lý dữ liệu đồ thị. Số epoch được chọn đủ lớn để mô hình có thời gian hội tụ. Chiều nhúng của nút và cạnh được thiết lập là 512 nhằm đảm bảo khả năng biểu diễn đủ lớn cho các quan hệ phức tạp trong đồ thị.

Bảng 2.1: Các tham số huấn luyện chính

Tham số	Giá trị
Phương pháp tối ưu	Adam
Learning rate ban đầu	0.0005
Weight decay	0.00001
Bộ điều chỉnh learning rate	CosineAnnealingWarmRestarts
T_0	10
Hàm mất mát	Mean Squared Error (MSE)
Số epoch	1500
Batch size	4
Ngưỡng tạo cạnh β	0.5
Chiều nhúng của nút	512
Chiều nhúng của cạnh	512

2.5.6. Kết quả thực nghiệm

Trong bước đầu tiên, nghiên cứu tiến hành khảo sát có hệ thống các tổ hợp khác nhau giữa các kiến trúc GNN và kỹ thuật graph pooling. Cụ thể, ba kiến trúc GNN được xem xét bao gồm Graph with Self-Attention (GSA), Graph Attention Network (GAT) và Graph Attention Network v2 (GATv2). Đối với tầng graph pooling, hai phương pháp node dropping là TopKPooling và SAGPool được sử dụng. Ngoài ra, Graph Isomorphism Network (GIN) được đưa vào như một mô hình đối chứng nhằm đánh giá hiệu quả của các kiến trúc không sử dụng trọng số cạnh và cơ chế chú ý. Biểu thức cập nhật nút trong GIN được sử dụng trong thí nghiệm như sau:

$$h'_i = \text{MLP} \left((1 + \varepsilon)h_i + \sum_{j \in \mathcal{N}_i} h_j \right), \quad (2.31)$$

trong đó ε là một tham số cố định hoặc học được. Kết quả tổng hợp của các thí nghiệm lựa chọn mô hình được trình bày trong Bảng 2.2.

Kết quả cho thấy tổ hợp sử dụng GSA ở tầng đồ thị đô thị, kết hợp với TopKPooling và GAT ở tầng đồ thị con theo khu vực đạt hiệu năng cao nhất với sai số MSE thấp nhất và hệ số R^2 cao nhất. Các tổ hợp sử dụng SAGPool hoặc GIN trong nhiều trường

Bảng 2.2: Lựa chọn tổ hợp tối ưu giữa GNN và Graph Pooling

	GSA		GATv2		GAT		GIN	
	MSE	R^2	MSE	R^2	MSE	R^2	MSE	R^2
TopKPooling + GAT + Global Mean Pooling	21.64	0.558	23.29	0.525	35.45	0.24	41.43	0.112
SAGPool + GAT + Global Mean Pooling	inf	inf	inf	inf	46.02	0.061	inf	inf
Global Mean Pooling	25.90	0.444	26.48	0.443	24.98	0.464	42.81	0.081

Bảng 2.3: Lựa chọn tham số k cho TopKPooling

	GSA		GATv2	
	MSE	R^2	MSE	R^2
$k = 0.1$	25.30	0.484	28.46	0.419
$k = 0.2$	47.46	0.031	23.99	0.511
$k = 0.3$	21.64	0.558	26.93	0.450
$k = 0.4$	47.85	0.031	23.29	0.525
$k = 0.5$	22.31	0.545	24.83	0.493
$k = 0.6$	23.06	0.530	25.86	0.472
$k = 0.7$	23.95	0.511	23.54	0.512
$k = 0.8$	23.49	0.520	26.46	0.460
$k = 0.9$	47.60	0.029	23.89	0.513

hợp dẫn đến quá trình huấn luyện không ổn định, phản ánh sự không phù hợp của các cấu hình này đối với dữ liệu giám sát HIV có cấu trúc phức tạp.

Lựa chọn tham số k trong TopKPooling. Để đánh giá ảnh hưởng của tỷ lệ nút được giữ lại trong TopKPooling, nghiên cứu tiến hành thử nghiệm với các giá trị k từ 0.1 đến 0.9. Kết quả chi tiết được trình bày trong Bảng 2.3.

Kết quả cho thấy hiệu năng dự báo phụ thuộc mạnh vào giá trị k . Các giá trị k quá nhỏ hoặc quá lớn đều làm suy giảm hiệu năng do hoặc mất thông tin quan trọng, hoặc giữ lại quá nhiều nhiễu. Giá trị $k = 0.3$ đối với GSA và $k = 0.4$ đối với GATv2 cho kết quả tốt nhất.

Phân tích ảnh hưởng của kiến trúc GNN sau pooling. Để đánh giá vai trò của cơ chế chú ý trên đồ thị con theo khu vực, nghiên cứu tiến hành thay thế GAT bằng GCN trong tầng sau pooling. Kết quả được trình bày trong Bảng 2.4.

Việc thay thế GAT bằng GCN giúp quá trình huấn luyện ổn định hơn nhưng làm giảm đáng kể hiệu năng dự báo, cho thấy vai trò quan trọng của cơ chế chú ý và khả năng khai thác thông tin cạnh trong việc học động lực lan truyền HIV.

So sánh với các phương pháp dự báo hiện có. Sau khi lựa chọn được cấu hình mô hình tối ưu, nghiên cứu tiến hành so sánh mô hình đề xuất với các phương pháp dự báo chuỗi thời gian và không gian–thời gian đã được sử dụng trong các nghiên cứu trước. Kết quả được trình bày trong Bảng 2.5.

Kết quả cho thấy mô hình đề xuất vượt trội so với tất cả các phương pháp so sánh cả về MSE và R^2 , chứng minh tính phù hợp của cách tiếp cận học đồ thị cao dương tính cấp đô thị trong bài toán dự báo HIV ngắn hạn.

Phân tích quá trình huấn luyện và đánh giá khả năng tổng quát hóa của mô hình.

Để đánh giá quá trình học và khả năng tổng quát hóa của mô hình đề xuất, nghiên

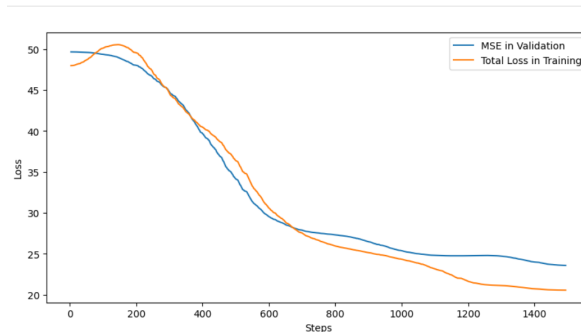
Bảng 2.4: So sánh khi thay thế GAT bằng GCN

	GSA		GATv2		GAT		GIN	
	MSE	R^2	MSE	R^2	MSE	R^2	MSE	R^2
TopKPooling + GCN + Global Mean Pooling	29.88	0.36	28.86	0.38	35.45	0.24	41.43	0.112
SAGPool + GCN + Global Mean Pooling	34.50	0.26	34.24	0.31	42.00	0.10	44.54	0.323

Bảng 2.5: So sánh mô hình đề xuất với các phương pháp baseline

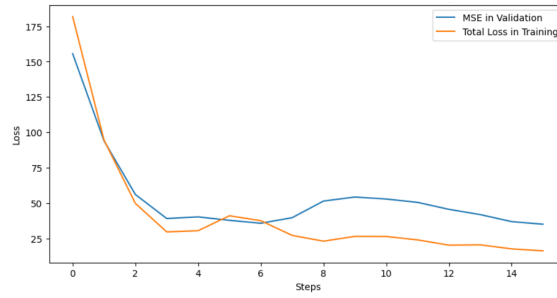
Method	Proposed	MPNN LSTM	GAT LSTM	DC RNN	GCN GRU	GCN LSTM	LSTM	BI LSTM	STACKED LSTM	CNN	CNN LSTM	ARIMA
MSE	21.64	35.11	50.15	52.78	52.32	60.53	48.51	45.99	47.06	46.27	52.22	3944.64
R^2	0.558	0.244	-0.008	-0.136	-0.126	-0.303	0.010	0.010	-0.013	0.003	-0.120	-0.768

cứu tiến hành phân tích động học huấn luyện thông qua các đường cong mất mát (loss), sai số bình phương trung bình (MSE), và hệ số xác định (R^2) trên tập huấn luyện và tập kiểm định. Các kết quả này được so sánh trực tiếp với mô hình đối chứng MPNN-LSTM nhằm làm rõ sự khác biệt về hành vi học và mức độ ổn định của hai phương pháp. Hình 2.2 minh họa sự biến thiên của hàm mất mát huấn luyện và MSE trên tập kiểm định đối với mô hình đề xuất. Có thể quan sát thấy rằng trong khoảng 500 bước huấn luyện đầu tiên, cả tổng loss và MSE đều giảm nhanh, cho thấy mô hình học được các cấu trúc chính trong dữ liệu một cách hiệu quả. Khi quá trình huấn luyện tiếp diễn, loss tiếp tục giảm trong khi MSE trên tập kiểm định dần ổn định, phản ánh khả năng học đặc trưng có tính khái quát tốt cho cả tập huấn luyện và tập kiểm định. Mặc dù tồn tại một số dao động nhỏ của MSE trên tập kiểm định, dấu hiệu overfitting nếu có chỉ ở mức nhẹ và vẫn nằm trong phạm vi chấp nhận được đối với bài toán dự báo dịch tễ thực tế.



Hình 2.2: Mô hình đề xuất: Loss huấn luyện và MSE trên tập kiểm định

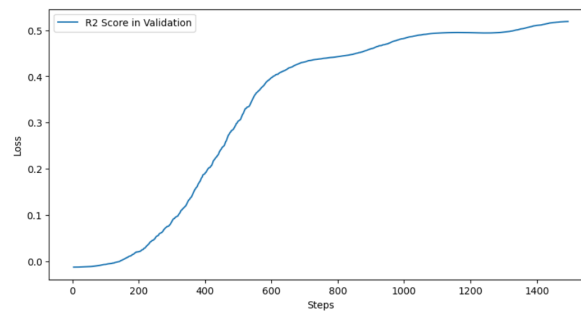
Ngược lại, Hình 3.6 trình bày động học huấn luyện của mô hình MPNN-LSTM. Mặc dù mô hình này cũng học rất nhanh trong giai đoạn đầu, MSE sau đó có xu hướng giảm chậm và sớm đạt trạng thái bão hòa. Điều này cho thấy mô hình có xu hướng ghi nhớ dữ liệu huấn luyện nhưng gặp khó khăn trong việc tổng quát hóa sang dữ liệu chưa quan sát, dẫn đến hiệu năng dự báo kém hơn trên tập kiểm định.



Hình 2.3: MPNN-LSTM: Loss huấn luyện và MSE trên tập kiểm định

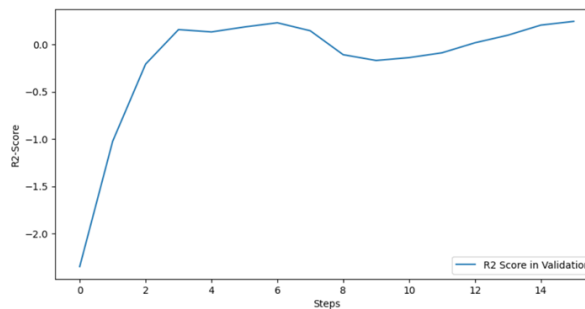
2.5.7. Phân tích hệ số xác định (R^2) trong quá trình kiểm định

Để đánh giá mức độ phù hợp giữa giá trị dự báo và giá trị thực tế, nghiên cứu tiếp tục phân tích sự biến thiên của hệ số xác định R^2 trên tập kiểm định trong suốt quá trình huấn luyện. Hình 2.4 cho thấy R^2 của mô hình đề xuất tăng đều và ổn định ngay từ giai đoạn đầu, nhanh chóng vượt ngưỡng 0.5 và duy trì ở mức cao. Điều này cho thấy mô hình có khả năng giải thích một phần lớn phương sai của dữ liệu kiểm định, phản ánh năng lực tổng quát hóa tốt.



Hình 2.4: Mô hình đề xuất: Giá trị R^2 trên tập kiểm định

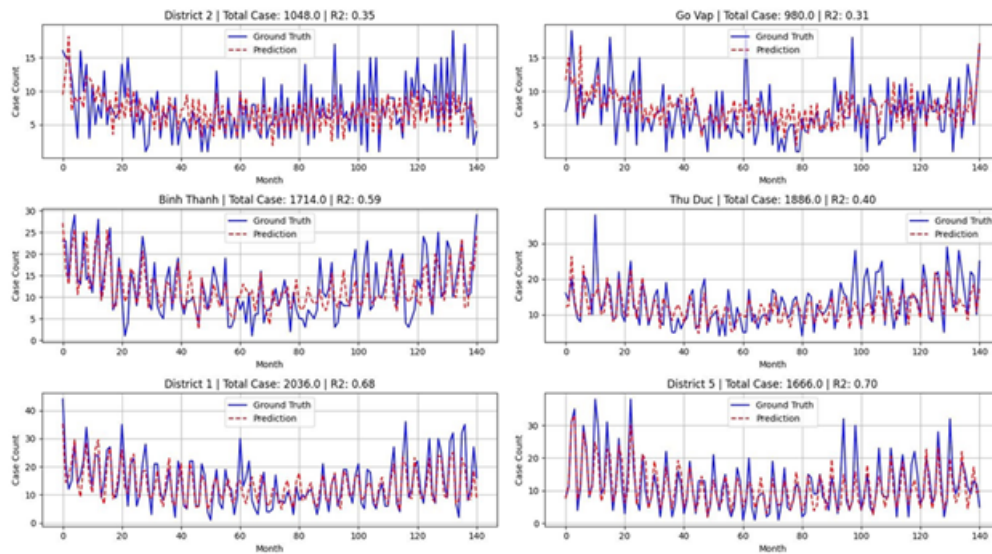
Ngược lại, Hình 2.5 cho thấy R^2 của mô hình MPNN-LSTM dao động mạnh và hiếm khi đạt giá trị cao, phản ánh mức độ phù hợp thấp giữa dự báo và dữ liệu thực tế. Sự khác biệt này khẳng định rằng mô hình đề xuất có khả năng học và duy trì cấu trúc lan truyền HIV ổn định hơn so với mô hình đối chứng.



Hình 2.5: MPNN-LSTM: Giá trị R^2 trên tập kiểm định

Phân tích so sánh hiệu năng giữa các khu vực. Nghiên cứu thực hiện phân tích sâu hiệu năng dự báo tại sáu quận đại diện cho các đặc trưng dữ liệu khác nhau, bao gồm

khu vực trung tâm, khu vực ngoại thành và một trường hợp bất thường. Hình 2.6 trình bày so sánh trực quan giữa giá trị dự báo và thực tế tại các khu vực này.



Hình 2.6: So sánh hiệu năng dự báo tại sáu quận đại diện

Kết quả cho thấy các quận trung tâm (như Quận 1, Quận 5 và Bình Thạnh) có độ chính xác dự báo cao hơn nhờ số lượng ca dương tính lớn và xu hướng thời gian ổn định. Ngược lại, các quận ngoại thành như Quận 2 và Gò Vấp có số ca thấp và biến động yếu, khiến mô hình gặp khó khăn trong việc học các mẫu dịch tễ rõ ràng. Trường hợp Quận Thủ Đức thể hiện hành vi bất thường với biến động mạnh và không tuần hoàn, dẫn đến hiệu năng dự báo giảm nhưng vẫn tốt hơn các khu vực có dữ liệu cực kỳ thưa thớt.

Tổng hợp các kết quả trên cho thấy mô hình có khả năng khai thác hiệu quả thông tin lan truyền ở quy mô toàn thành phố, đồng thời tận dụng các khu vực có dữ liệu phong phú để cải thiện dự báo tại những khu vực hạn chế về dữ liệu. Điều này khẳng định giá trị của cách tiếp cận đồ thị ca bệnh cấp đô thị trong bài toán dự báo HIV ngắn hạn.

Phân tích ảnh hưởng của các yếu tố dịch tễ.

Nhằm đánh giá mức độ đóng góp của từng yếu tố dịch tễ trong đồ thị ca dương tính, nghiên cứu tiến hành thí nghiệm loại bỏ từng đặc trưng nút và huấn luyện lại mô hình trong cùng điều kiện thiết lập. Các kết quả thu được được trình bày trong Bảng 2.6.

Bảng 2.6: Kết quả mô hình khi loại bỏ từng đặc trưng dịch tễ của nút

	Đầy đủ	Bỏ vị trí	Bỏ nhóm tuổi	Bỏ giới	Bỏ nghề	Bỏ đường lây	Bỏ nhóm nguy cơ
MSE	21.64	24.00	21.65	21.96	26.05	22.39	19.52
R^2	0.558	0.510	0.558	0.552	0.459	0.543	0.602

Kết quả cho thấy việc loại bỏ các đặc trưng về vị trí địa lý, nghề nghiệp và đường lây truyền đều làm suy giảm đáng kể hiệu năng dự báo, phản ánh vai trò quan trọng của các yếu tố này trong động lực lan truyền HIV. Ngược lại, đặc trưng nhóm tuổi gần như không ảnh hưởng đến kết quả dự báo, khi các chỉ số MSE và R^2 gần như không

thay đổi so với mô hình đầy đủ.

Đáng chú ý, khi loại bỏ đặc trưng nhóm nguy cơ, hiệu năng mô hình được cải thiện rõ rệt. Điều này cho thấy đặc trưng này có thể chứa nhiều hoặc thông tin chồng chéo với các yếu tố dịch tễ khác trong dữ liệu giám sát, dẫn đến tác động tiêu cực khi được đưa vào mô hình. Kết quả này gợi ý rằng việc lựa chọn đặc trưng trong mô hình dự báo HIV cần được cân nhắc kỹ lưỡng dựa trên chất lượng và tính nhất quán của dữ liệu giám sát.

Đánh giá hiệu quả dự báo theo các khoảng thời gian.

Để đánh giá khả năng dự báo của mô hình trong các khoảng thời gian khác nhau, nghiên cứu tiến hành thí nghiệm với nhiều thiết lập độ dài chuỗi lịch sử và cửa sổ dự báo. Kết quả được trình bày trong Bảng 2.7.

Bảng 2.7: Hiệu quả dự báo theo các khoảng thời gian khác nhau

	15 ngày	1 tháng	2 tháng	3 tháng	1 tháng (1m data)	1 tháng (2m data)
MSE	11.00	19.52	38.83	41.48	22.43	22.64
R²	0.784	0.602	0.369	0.128	0.557	0.553

Kết quả cho thấy mô hình đạt hiệu năng cao nhất trong các kịch bản dự báo ngắn hạn, đặc biệt là trong khoảng từ 15 ngày đến 1 tháng. Khi cửa sổ dự báo được mở rộng lên 2 hoặc 3 tháng, sai số tăng đáng kể và hệ số R^2 giảm mạnh, cho thấy khả năng giải thích biến thiên dữ liệu suy giảm theo thời gian.

Hiện tượng này phản ánh bản chất động của dịch HIV, trong đó các yếu tố hành vi, can thiệp y tế và chính sách phòng chống có thể thay đổi theo thời gian nhưng không được mô hình hóa trực tiếp trong đồ thị đầu vào. Do đó, mô hình phù hợp nhất cho các bài toán dự báo ngắn hạn nhằm hỗ trợ ra quyết định kịp thời trong giám sát và can thiệp y tế cộng đồng.

2.6. Thảo luận về các lựa chọn tham số thiết kế, ưu điểm và hạn chế của mô hình

Mặc dù mô hình đề xuất đạt kết quả dự báo tốt trên bộ dữ liệu nghiên cứu, hiệu năng của mô hình vẫn chịu ảnh hưởng bởi một số lựa chọn thiết kế trong quá trình xây dựng phương pháp.

Thứ nhất, ngưỡng xây dựng đồ thị β quyết định mật độ kết nối giữa các nút trong đồ thị. Khi ngưỡng được lựa chọn thấp, số lượng cạnh tăng lên, giúp khai thác nhiều mối liên hệ hơn giữa các ca bệnh nhưng đồng thời làm đồ thị trở nên dày đặc, gia tăng chi phí tính toán và có thể đưa thêm các kết nối ít có ý nghĩa. Ngược lại, ngưỡng quá cao tạo ra đồ thị thưa hơn, làm giảm khả năng lan truyền thông tin giữa các nút và có thể bỏ sót một số mối quan hệ dịch tễ quan trọng. Trong nghiên cứu này, ngưỡng được lựa chọn dựa trên đặc điểm của bộ dữ liệu và kết quả thực nghiệm nhằm cân bằng giữa khả năng biểu diễn của đồ thị và hiệu quả tính toán.

Thứ hai, cửa sổ thời gian xây dựng đồ thị cũng ảnh hưởng đáng kể đến hiệu quả của mô hình. Cửa sổ thời gian ngắn có thể chưa khai thác đầy đủ thông tin về diễn biến dịch, trong khi cửa sổ quá dài làm số lượng ca bệnh và số cạnh tăng nhanh, khiến

đồ thị trở nên dày đặc và thời gian huấn luyện tăng đáng kể. Kết quả thực nghiệm cho thấy cửa sổ 3 tháng là lựa chọn phù hợp, vừa bảo đảm khả năng biểu diễn thông tin, vừa duy trì hiệu quả tính toán.

Thứ ba, tỷ lệ TopKPooling ảnh hưởng đến lượng thông tin được giữ lại trong quá trình học biểu diễn đồ thị. Nếu tỷ lệ quá nhỏ, nhiều nút quan trọng có thể bị loại bỏ; ngược lại, tỷ lệ quá lớn làm giảm hiệu quả của bước tổng hợp đồ thị. Thực nghiệm với các giá trị từ 0,1 đến 0,9 cho thấy việc lựa chọn tỷ lệ phù hợp giúp mô hình đạt độ chính xác dự báo cao hơn.

Ngoài ra, mô hình được thiết kế cho bài toán dự báo ngắn hạn. Khi kéo dài khoảng thời gian dự báo vượt quá 3 tháng, độ chính xác có xu hướng giảm do nhiều yếu tố mới phát sinh như sự thay đổi hành vi nguy cơ, các biện pháp can thiệp y tế hoặc các yếu tố xã hội chưa được phản ánh trong dữ liệu đầu vào. Vì vậy, mô hình phù hợp hơn với các bài toán hỗ trợ giám sát và cảnh báo sớm trong ngắn hạn.

Nhìn chung, ưu điểm của mô hình là kết hợp được thông tin không gian và thời gian thông qua biểu diễn đồ thị, cho phép khai thác mối liên hệ giữa các ca bệnh để nâng cao hiệu quả dự báo. Tuy nhiên, mô hình vẫn phụ thuộc vào các lựa chọn thiết kế trong quá trình xây dựng đồ thị và hiệu quả dự báo có xu hướng giảm khi áp dụng cho các khoảng thời gian dự báo dài. Trong tương lai, nghiên cứu có thể tiếp tục khảo sát sâu hơn ảnh hưởng của các lựa chọn thiết kế này cũng như phát triển các phương pháp xây dựng đồ thị thích nghi nhằm nâng cao khả năng tổng quát hóa của mô hình.

2.7. kết luận chương

Chương này đã nghiên cứu bài toán dự báo ngắn hạn số ca dương tính HIV từ dữ liệu giám sát ca bệnh. Trên cơ sở phân tích những hạn chế của các phương pháp dự báo chuỗi thời gian truyền thống, luận án đã đề xuất một cách tiếp cận dựa trên học sâu trên đồ thị nhằm khai thác đồng thời các mối quan hệ không gian và thời gian giữa các ca bệnh trong quá trình lan truyền dịch.

Điểm mới đầu tiên của chương là xây dựng đồ thị ca bệnh từ dữ liệu giám sát HIV, trong đó mỗi ca dương tính được biểu diễn bởi một nút và các cạnh phản ánh mối quan hệ không gian – thời gian giữa các ca bệnh. Cách biểu diễn này cho phép mô hình học trực tiếp từ cấu trúc lan truyền của dịch thay vì chỉ sử dụng số liệu tổng hợp theo khu vực như nhiều nghiên cứu trước đây. Trên cơ sở đó, luận án đề xuất mô hình mạng nơ-ron đồ thị dựa trên cơ chế chú ý để học biểu diễn của mạng lây nhiễm ở quy mô toàn thành phố, sau đó tổng hợp thành biểu diễn của từng quận, huyện nhằm dự báo số ca dương tính HIV mới trong ngắn hạn. Kết quả thực nghiệm trên bộ dữ liệu giám sát HIV tại Thành phố Hồ Chí Minh giai đoạn 2008–2019 cho thấy mô hình đề xuất đạt độ chính xác dự báo cao hơn so với các phương pháp dự báo chuỗi thời gian và một số mô hình học sâu được lựa chọn để so sánh.

Mặc dù dự báo số ca dương tính mới có ý nghĩa quan trọng trong giám sát dịch, việc hỗ trợ điều trị HIV không chỉ dừng lại ở dự báo quy mô dịch mà còn đòi hỏi khả năng đánh giá diễn biến của từng người bệnh trong quá trình điều trị. Thời gian xảy ra các biến cố như thất bại điều trị, tiến triển bệnh hoặc tử vong có vai trò quan trọng

trong việc lựa chọn phác đồ, xây dựng kế hoạch theo dõi và đưa ra các can thiệp kịp thời. Đây là bài toán có bản chất khác với dự báo ở cấp quần thể và sẽ được nghiên cứu trong chương tiếp theo thông qua các mô hình phân tích sống sót trên dữ liệu dọc của bệnh nhâ

Chương 3. DỰ BÁO NGUY CƠ BẰNG PHÂN TÍCH SỐNG SỐT TRONG ĐIỀU TRỊ HIV

Chương này trình bày phương pháp phân tích sống sót được đề xuất nhằm dự báo nguy cơ xảy ra biến cố trong điều trị HIV. Trong bối cảnh dữ liệu điều trị HIV bao gồm cả thông tin nền, lịch sử thăm khám theo thời gian và các trường hợp kiểm duyệt, nghiên cứu hướng tới xây dựng một mô hình có khả năng mô hình hóa động lực tiến triển của trạng thái bệnh theo thời gian. Phương pháp được phát triển dựa trên khung mô hình không gian trạng thái, đồng thời tích hợp cơ chế bộ nhớ nhằm học các nguyên mẫu lâm sàng từ quần thể bệnh nhân và chiến lược học đa nhiệm để hỗ trợ ước lượng nguy cơ xảy ra biến cố một cách hiệu quả hơn.

3.1. Mô tả bài toán

Nghiên cứu này tập trung vào bài toán dự báo sống sót cá nhân hóa cho người nhiễm HIV tại thời điểm bắt đầu điều trị kháng vi rút (Antiretroviral Therapy – ART). Trong các nghiên cứu HIV, dữ liệu bệnh nhân thường xuất hiện dưới hai dạng chính: dữ liệu tĩnh tại thời điểm ban đầu và dữ liệu theo chuỗi thời gian trong quá trình theo dõi điều trị.

Trong trường hợp dữ liệu tĩnh, mỗi bệnh nhân i được mô tả bởi một vector đặc trưng ban đầu

$$x_i(0) \in \mathbb{R}^{d_s},$$

bao gồm các thông tin nhân khẩu học, kết quả xét nghiệm ban đầu và các chỉ số lâm sàng được ghi nhận tại thời điểm tham gia nghiên cứu.

Trong trường hợp dữ liệu theo chuỗi thời gian, ngoài các đặc trưng ban đầu, mỗi bệnh nhân còn có một chuỗi các quan sát theo thời gian

$$\{x_{i,t}^{(l)}\}_{t=1}^{T_i},$$

trong đó mỗi $x_{i,t}^{(l)} \in \mathbb{R}^{d_l}$ biểu diễn các thông tin lâm sàng được ghi nhận tại lần khám thứ t , chẳng hạn như số lượng tế bào CD4, tải lượng virus, tình trạng điều trị và các chỉ số sinh học liên quan. Hai dạng dữ liệu này phản ánh hai cấu trúc đầu vào khác nhau, do đó một mô hình dự báo hiệu quả cần có khả năng khai thác thông tin từ cả dữ liệu tĩnh và dữ liệu theo thời gian.

Tập dữ liệu nghiên cứu có thể được biểu diễn dưới dạng

$$\mathcal{D} = \{(X_i, T_i, \delta_i)\}_{i=1}^N,$$

trong đó T_i là thời gian sống sót quan sát được của bệnh nhân i , còn $\delta_i \in \{0, 1\}$ là biến chỉ báo biến cố, với $\delta_i = 1$ biểu thị biến cố xảy ra và $\delta_i = 0$ biểu thị dữ liệu bị kiểm duyệt.

Mục tiêu của bài toán là ước lượng hàm sống sót cá nhân hóa dựa trên toàn bộ

thông tin sẵn có của bệnh nhân, được biểu diễn như sau

$$\hat{S}_i(t) = P\left(T_i > t \mid x_i(0), \{x_{i,s}^{(l)}\}_{s \leq t}\right).$$

Trong khuôn khổ mô hình sống sót rời rạc theo thời gian, xác suất sống sót có thể được tính từ các xác suất nguy cơ (hazard) dự đoán theo công thức

$$\hat{S}_i(t) = \prod_{k=1}^t (1 - \hat{h}_i(k)).$$

Một thách thức quan trọng của bài toán này nằm ở sự khác biệt bản chất giữa hai dạng dữ liệu. Dữ liệu theo chuỗi thời gian phản ánh sự thay đổi động của các chỉ số lâm sàng trong quá trình điều trị, trong khi dữ liệu tĩnh chỉ mô tả trạng thái ban đầu của bệnh nhân. Do đó, mô hình dự báo cần có khả năng tích hợp hiệu quả cả hai dạng thông tin này trong một khung thống nhất, đồng thời đảm bảo khả năng ước lượng rủi ro sống sót một cách chính xác và ổn định trong các bối cảnh lâm sàng khác nhau.

3.2. Phân tích bài toán và cơ sở thiết kế mô hình

Bài toán tiên lượng sống sót trong điều trị HIV đặt ra nhiều thách thức do dữ liệu đầu vào bao gồm nhiều nhóm đặc trưng không đồng nhất như thông tin nhân khẩu học, dịch tễ học, điều trị, xét nghiệm và các chỉ dấu sinh học, đồng thời được thu thập theo dõi dọc trong thời gian dài. Mỗi nhóm đặc trưng phản ánh một khía cạnh khác nhau của trạng thái bệnh và có vai trò khác nhau trong quá trình dự báo. Vì vậy, nghiên cứu sử dụng các bộ mã hóa chuyên biệt để học biểu diễn riêng cho từng nhóm đặc trưng trước khi tích hợp thành một biểu diễn thống nhất.

Đối với dữ liệu theo dõi dọc, mô hình cần phản ánh được quá trình tiến triển của trạng thái bệnh theo thời gian. Các mô hình tuần tự như LSTM và GRU có khả năng khai thác phụ thuộc theo thời gian nhưng thường gặp hạn chế khi chuỗi quan sát dài hoặc dữ liệu không đầy đủ. Trong khi đó, các mô hình dựa trên cơ chế tự chú ý cải thiện khả năng học phụ thuộc dài hạn nhưng chủ yếu tập trung vào mối quan hệ giữa các thời điểm quan sát, thay vì mô hình hóa trực tiếp động lực tiến triển của trạng thái bệnh. Vì vậy, nghiên cứu lựa chọn mô hình không gian trạng thái có chọn lọc (Selective State Space Model - SSSM) dựa trên kiến trúc Mamba [95]. Với cơ chế cập nhật trạng thái thích ứng theo đầu vào, SSSM cho phép lựa chọn các tín hiệu lâm sàng quan trọng và mô hình hóa linh hoạt hơn quá trình tiến triển của bệnh.

Mặc dù SSSM mô hình hóa hiệu quả động lực tiến triển của từng bệnh nhân, việc học chỉ từ chuỗi quan sát của cá thể có thể chưa khai thác được đầy đủ tri thức ở mức quần thể. Trong thực tế, nhiều bệnh nhân có đặc điểm ban đầu khác nhau nhưng vẫn có thể trải qua các diễn biến bệnh tương tự. Vì vậy, nghiên cứu bổ sung mô-đun bộ nhớ tri thức quần thể nhằm lưu trữ các nguyên mẫu lâm sàng được học trong quá trình huấn luyện. Thông qua cơ chế truy hồi, mô hình lựa chọn các nguyên mẫu phù hợp để bổ sung thông tin cho biểu diễn của bệnh nhân hiện tại. Cơ chế này giúp kết hợp tri thức của cá thể với tri thức tích lũy từ quần thể, qua đó nâng cao chất lượng biểu

diễn và khả năng tổng quát hóa của mô hình.

Bên cạnh dự báo thời gian xảy ra biến cố, bài toán còn yêu cầu phân tầng nguy cơ và hỗ trợ đánh giá kết quả điều trị cho từng bệnh nhân. Vì vậy, nghiên cứu lựa chọn thiết kế học đa mục tiêu thay vì tối ưu một mục tiêu duy nhất. Thành phần *survival loss* giúp mô hình học phân bố thời gian xảy ra biến cố; *ranking loss* cải thiện khả năng phân biệt mức độ nguy cơ giữa các bệnh nhân; *treatment-specific outcome loss* hỗ trợ dự báo kết quả điều trị theo từng phương án can thiệp; trong khi điều chuẩn Bayes [79] giúp giảm quá khớp và lượng hóa bất định của dự báo. Việc tối ưu đồng thời các mục tiêu này giúp mô hình đáp ứng đầy đủ hơn yêu cầu của bài toán tiên lượng và hỗ trợ ra quyết định lâm sàng.

Cuối cùng, khả năng diễn giải là một yêu cầu quan trọng đối với các hệ thống hỗ trợ quyết định trong y tế. Do đó, nghiên cứu tích hợp phương pháp SHAP [96] nhằm phân tích mức độ đóng góp của từng đặc trưng vào kết quả dự báo, qua đó nâng cao tính minh bạch và hỗ trợ giải thích các quyết định của mô hình.

Từ các phân tích trên, nghiên cứu xây dựng một kiến trúc thống nhất gồm bộ mã hóa đặc trưng theo nhóm, bộ nhớ tri thức quần thể, mô hình không gian trạng thái có chọn lọc, mô-đun cơ chế học đa mục tiêu và mô-đun giải thích mô hình. Thiết kế này vừa phù hợp với đặc điểm của dữ liệu HIV theo dõi dọc, vừa đáp ứng yêu cầu dự báo sống sót, phân tầng nguy cơ và hỗ trợ ra quyết định điều trị.

3.3. Phương pháp đề xuất

3.3.1. Tổng quan mô hình đề xuất

Hình 3.1 trình bày kiến trúc tổng thể của mô hình đề xuất nhằm dự báo sống sót và hỗ trợ dự báo kết cục điều trị cho người sống chung với HIV từ dữ liệu theo dõi dọc.

Cho tập dữ liệu gồm N bệnh nhân $\mathcal{D} = \{(X_i, T_i, \delta_i)\}_{i=1}^N$, trong đó X_i là chuỗi đặc trưng theo dõi của bệnh nhân, T_i là thời gian theo dõi hoặc thời điểm xảy ra biến cố và δ_i là biến kiểm duyệt. Đối với bệnh nhân thứ i , dữ liệu theo dõi được biểu diễn như sau: $X_i = (x_i^{(1)}, x_i^{(2)}, \dots, x_i^{(L_i)})$,

với L_i là số lần tái khám. Mục tiêu của mô hình là học ánh xạ

$$f : X \longrightarrow (\hat{S}(t), \hat{Y}) \quad (3.1)$$

trong đó $\hat{S}(t)$ là xác suất sống sót dự báo theo thời gian và $\hat{Y}$ là kết cục điều trị dự báo.

Để thực hiện ánh xạ này, mô hình gồm bốn mô-đun chính: bộ mã hóa đặc trưng, mô-đun học tri thức quần thể dựa trên bộ nhớ, mô hình học trạng thái bệnh tiềm ẩn bằng không gian trạng thái có chọn lọc (SSSM) và mô-đun dự báo đa nhiệm. Dòng xử lý tổng quát của mô hình được mô tả như sau

$$X_t \rightarrow \mathbf{h}_t \rightarrow \tilde{\mathbf{h}}_t \rightarrow \mathbf{o}_t \rightarrow (\hat{S}(t), \hat{Y}), \quad (3.2)$$

trong đó $\mathbf{h}_t$ là biểu diễn đặc trưng sau bộ mã hóa, $\tilde{\mathbf{h}}_t$ là biểu diễn sau khi tích hợp tri thức quần thể từ mô-đun bộ nhớ và $\mathbf{o}_t$ là biểu diễn trạng thái bệnh cuối cùng do mô

hình SSSM học được. Biểu diễn này được sử dụng đồng thời cho nhánh dự báo sống sót và nhánh dự báo kết cục điều trị. Sau khi hoàn thành quá trình huấn luyện, phương pháp SHAP được áp dụng để giải thích các dự báo và phân tích mức độ đóng góp của từng đặc trưng theo thời gian.



Hình 3.1: Tổng quan mô hình đề xuất cho phân tích sống sót

3.3.2. Mã hóa đặc trưng lâm sàng

Dữ liệu lâm sàng HIV bao gồm nhiều nhóm thông tin có bản chất khác nhau, chẳng hạn như đặc điểm nhân khẩu học, các chỉ số lâm sàng và thông tin điều trị. Việc sử dụng một bộ mã hóa chung cho toàn bộ dữ liệu có thể làm giảm khả năng biểu diễn do sự khác biệt về ngữ nghĩa giữa các nhóm đặc trưng. Vì vậy, luận án sử dụng cơ chế mã hóa đa nguồn, trong đó mỗi nhóm đặc trưng được học bởi một bộ mã hóa riêng trước khi được tích hợp thành một biểu diễn thống nhất.

Tại mỗi lần tái khám thứ t của bệnh nhân i , tập đặc trưng đầu vào được chia thành ba nhóm gồm đặc trưng nhân khẩu học, đặc trưng lâm sàng và đặc trưng điều trị. Biểu diễn của từng nhóm được học thông qua một mạng Multi-Layer Perceptron (MLP) độc lập

$$\mathbf{h}_{i,t}^m = f_m(\mathbf{x}_{i,t}^m), \quad m \in \{\text{demo, clinical, treat}\}, \quad (3.3)$$

trong đó $\mathbf{x}_{i,t}^m$ và $\mathbf{h}_{i,t}^m$ lần lượt là vector đặc trưng đầu vào và biểu diễn tiềm ẩn của nhóm đặc trưng m , còn $f_m(\cdot)$ là bộ mã hóa MLP tương ứng.

Đối với các đặc trưng lâm sàng được suy diễn hoặc tính toán từ dữ liệu gốc, các đặc trưng này được xem như những biến lâm sàng bổ sung và được xử lý bởi cùng bộ mã hóa lâm sàng. Thiết kế này giúp mô hình dễ dàng mở rộng để tích hợp các chỉ dấu sinh học mới mà không cần thay đổi kiến trúc tổng thể.

Sau khi được mã hóa, các biểu diễn của ba nhóm đặc trưng được nối lại để tạo

thành biểu diễn hợp nhất

$$\mathbf{h}_{i,t}^{\text{concat}} = [\mathbf{h}_{i,t}^{\text{demo}}; \mathbf{h}_{i,t}^{\text{clinical}}; \mathbf{h}_{i,t}^{\text{treat}}], \quad (3.4)$$

trong đó $[\cdot; \cdot]$ ký hiệu phép nối (concatenation) các vector đặc trưng.

Đối với dữ liệu theo dõi dọc, biểu diễn hợp nhất tiếp tục được bổ sung thông tin về thứ tự thời gian thông qua mã hóa vị trí dạng sin-cos (sinusoidal positional encoding) theo [62]. Biểu diễn đầu vào của mô hình được xác định như sau

$$\mathbf{h}_{i,t} = \mathbf{h}_{i,t}^{\text{concat}} + PE(t), \quad (3.5)$$

trong đó $PE(t)$ là vector mã hóa vị trí tại lần tái khám thứ t . Việc bổ sung thông tin vị trí giúp mô hình phân biệt các thời điểm theo dõi khác nhau trong khi vẫn sử dụng chung kiến trúc mã hóa cho toàn bộ chuỗi dữ liệu. Đối với dữ liệu chỉ gồm một lần quan sát ban đầu, thành phần mã hóa vị trí được cố định tại thời điểm gốc. Biểu diễn $\mathbf{h}_{i,t}$ sau bước mã hóa đa nguồn được sử dụng làm đầu vào của mô-đun học tri thức quần thể dựa trên bộ nhớ được trình bày trong mục tiếp theo.

3.3.3. Thành phần học tri thức quần thể dựa trên bộ nhớ

Bộ mã hóa đặc trưng tạo ra biểu diễn tiềm ẩn của bệnh nhân tại mỗi lần tái khám từ các nhóm đặc trưng nhân khẩu học, lâm sàng và điều trị. Biểu diễn này phản ánh trạng thái hiện tại của bệnh nhân nhưng chủ yếu được xây dựng từ thông tin của chính cá thể đó. Trong thực tế, những bệnh nhân có đặc điểm ban đầu khác nhau vẫn có thể trải qua các diễn biến bệnh tương tự. Để khai thác các mẫu tiến triển chung ở mức quần thể, nghiên cứu đề xuất một thành phần học tri thức dựa trên bộ nhớ nguyên mẫu khả học. Thành phần này gồm 3 phần: ngân hàng bộ nhớ nguyên mẫu toàn cục, cơ chế truy xuất bộ nhớ bằng *cross-attention* có xét đến thời gian và cơ chế hợp nhất tri thức quần thể với biểu diễn của bệnh nhân.

Bộ nhớ nguyên mẫu toàn cục. Gọi $\mathbf{h}_{i,t} \in \mathbb{R}^d$ là biểu diễn tiềm ẩn của bệnh nhân i tại lần tái khám thứ t sau bước mã hóa đặc trưng. Mô hình duy trì một ngân hàng bộ nhớ toàn cục gồm K véc tơ nguyên mẫu

$$\mathbf{M} = [\mathbf{m}_1, \mathbf{m}_2, \dots, \mathbf{m}_K]^T \in \mathbb{R}^{K \times d}, \quad (3.6)$$

trong đó mỗi $\mathbf{m}_k \in \mathbb{R}^d$ là một nguyên mẫu tiềm ẩn có thể học. Các nguyên mẫu không tương ứng với một bệnh nhân cụ thể và cũng không được xác định trước từ tri thức chuyên gia. Thay vào đó, chúng được khởi tạo ngẫu nhiên cùng với các tham số khác của mô hình và được cập nhật thông qua lan truyền ngược trong quá trình huấn luyện. Sau nhiều vòng lặp tối ưu, các nguyên mẫu dần biểu diễn những trạng thái hoặc xu hướng tiến triển bệnh thường gặp trong quần thể nghiên cứu.

Trong giai đoạn huấn luyện, ngân hàng bộ nhớ được cập nhật đồng thời với toàn bộ mạng. Khi suy luận, ngân hàng bộ nhớ đã học được được giữ cố định và được sử dụng chung cho tất cả bệnh nhân.

Truy xuất bộ nhớ bằng cơ chế chú ý liên kết có xét đến thời gian. Tại mỗi lần

tái khám, biểu diễn hiện tại của bệnh nhân được sử dụng để truy vấn các nguyên mẫu phù hợp nhất trong ngân hàng bộ nhớ. Mặc dù biểu diễn $\mathbf{h}_{i,t}$ đã chứa thông tin về vị trí của lần tái khám thông qua bước mã hóa vị trí, khoảng cách thực tế giữa các lần theo dõi vẫn mang ý nghĩa lâm sàng quan trọng. Vì vậy, mô hình sử dụng thêm thông tin thời gian theo dõi để xây dựng truy vấn của cơ chế chú ý.

Gọi $\tau_{i,t}$ là thời gian theo dõi tích lũy hoặc khoảng thời gian tương ứng với lần tái khám thứ t của bệnh nhân i . Thông tin này được ánh xạ sang một véc tơ biểu diễn thời gian

$$\mathbf{e}_{i,t} = f_{\text{time}}(\tau_{i,t}), \quad (3.7)$$

trong đó $f_{\text{time}}(\cdot)$ là một phép chiếu khả học.

Véc tơ truy vấn được xây dựng từ biểu diễn bệnh nhân kết hợp với thông tin thời gian như sau:

$$\mathbf{q}_{i,t} = \mathbf{W}_Q(\mathbf{h}_{i,t} + \mathbf{e}_{i,t}), \quad (3.8)$$

trong đó $\mathbf{W}_Q \in \mathbb{R}^{d \times d}$ là ma trận trọng số khả học. Các véc tơ khóa và véc tơ giá trị được sinh từ ngân hàng bộ nhớ như sau:

$$\mathbf{K} = \mathbf{M}\mathbf{W}_K, \quad \mathbf{V} = \mathbf{M}\mathbf{W}_V, \quad (3.9)$$

với $\mathbf{W}_K, \mathbf{W}_V \in \mathbb{R}^{d \times d}$.

Trọng số chú ý giữa biểu diễn hiện tại của bệnh nhân và các nguyên mẫu được tính bằng cơ chế chú ý tích vô hướng có chuẩn hóa (*scaled dot-product attention*) như sau:

$$\alpha_{i,t} = \text{Softmax} \left(\frac{\mathbf{q}_{i,t} \mathbf{K}^\top}{\sqrt{d}} \right), \quad (3.10)$$

trong đó $\alpha_{i,t} \in \mathbb{R}^K$ biểu diễn mức độ liên quan của từng nguyên mẫu đối với trạng thái của bệnh nhân. Biểu diễn truy xuất từ bộ nhớ được xác định như sau:

$$\mathbf{r}_{i,t} = \alpha_{i,t} \mathbf{V}. \quad (3.11)$$

Thay vì lựa chọn một nguyên mẫu duy nhất, cơ chế chú ý cho phép mô hình kết hợp thông tin từ nhiều nguyên mẫu với các trọng số khác nhau. Nhờ đó, biểu diễn truy xuất phản ánh đồng thời nhiều mẫu tiến triển bệnh tiềm ẩn trong quần thể, phù hợp với tính không đồng nhất của dữ liệu điều trị HIV.

Hợp nhất biểu diễn bộ nhớ. Biểu diễn $\mathbf{h}_{i,t}$ phản ánh trạng thái hiện tại của bệnh nhân, trong khi $\mathbf{r}_{i,t}$ chứa tri thức quần thể được truy xuất từ ngân hàng bộ nhớ. Hai nguồn thông tin này được kết hợp thông qua một hàm hợp nhất và kết nối dư như sau:

$$\tilde{\mathbf{h}}_{i,t} = \mathbf{h}_{i,t} + f_{\text{fusion}}([\mathbf{h}_{i,t}; \mathbf{r}_{i,t}]), \quad (3.12)$$

trong đó $f_{\text{fusion}}(\cdot)$ là một mạng MLP gồm lớp tuyến tính, hàm kích hoạt phi tuyến, chuẩn hóa và dropout.

Biểu diễn $\tilde{\mathbf{h}}_{i,t}$ là đầu ra của mô-đun học tri thức quần thể dựa trên bộ nhớ. Kết nối

đư giúp bảo toàn thông tin đặc trưng của từng bệnh nhân đồng thời bổ sung tri thức quần thể được học từ dữ liệu. Biểu diễn này sau đó được sử dụng làm đầu vào của mô-đun học trạng thái bệnh có chọn lọc để mô hình hóa diễn biến bệnh theo thời gian.

3.3.4. Học trạng thái bệnh tiềm ẩn bằng Mô hình hóa không gian trạng thái có chọn lọc

Sau khi được tăng cường bởi mô-đun học tri thức quần thể dựa trên bộ nhớ, biểu diễn của bệnh nhân tiếp tục được đưa vào mô-đun học trạng thái bệnh tiềm ẩn nhằm mô hình hóa quá trình tiến triển của bệnh theo thời gian. Khác với biểu diễn tại từng thời điểm, trạng thái bệnh tiềm ẩn không chỉ phản ánh thông tin lâm sàng hiện tại mà còn tích lũy diễn biến bệnh qua toàn bộ chuỗi theo dõi, đồng thời kế thừa tri thức quần thể được truy xuất từ mô-đun bộ nhớ. Nhờ đó, mô hình xây dựng được một biểu diễn động phản ánh đầy đủ hơn đặc điểm tiến triển của từng bệnh nhân.

Để học trạng thái bệnh tiềm ẩn, luận án sử dụng mô hình không gian trạng thái có chọn lọc (Selective State Space Model – SSSM) dựa trên kiến trúc Mamba. Khác với các mô hình dựa trên cơ chế attention phải tính toán mối quan hệ giữa mọi cặp thời điểm quan sát, SSSM mô hình hóa sự tiến triển của bệnh thông qua quá trình cập nhật trạng thái tuần tự. Cách tiếp cận này vừa duy trì khả năng học các phụ thuộc dài hạn, vừa có độ phức tạp tuyến tính theo chiều dài chuỗi, phù hợp với dữ liệu theo dõi dọc trong điều trị HIV.

Quá trình học trạng thái bệnh gồm bốn bước chính: (i) chiếu biểu diễn đầu vào sang không gian trạng thái; (ii) mô hình hóa động lực học liên tục của trạng thái bệnh; (iii) xây dựng cơ chế chuyển trạng thái phụ thuộc đầu vào; và (iv) cập nhật trạng thái theo cơ chế đệ quy kết hợp với cơ chế chọn lọc đặc trưng để tạo ra biểu diễn trạng thái cuối cùng.

Phép chiếu biểu diễn vào không gian trạng thái. Đầu vào của mô hình là biểu diễn đã được tăng cường bởi mô-đun bộ nhớ, ký hiệu là $\tilde{\mathbf{h}}_{i,t}$. Biểu diễn này đồng thời chứa thông tin lâm sàng hiện tại và tri thức quần thể được học từ toàn bộ tập dữ liệu. Tuy nhiên, không gian đặc trưng ban đầu chưa đủ linh hoạt để mô hình hóa động lực học của bệnh theo thời gian. Vì vậy, biểu diễn được ánh xạ sang một không gian trạng thái có số chiều lớn hơn thông qua phép biến đổi tuyến tính như sau:

$$\mathbf{u}_{i,t} = \mathbf{W}_{\text{in}} \tilde{\mathbf{h}}_{i,t}, \quad (3.13)$$

trong đó $\mathbf{W}_{\text{in}} \in \mathbb{R}^{d_e \times d}$ là ma trận trọng số khả học, ánh xạ từ không gian đặc trưng có kích thước d sang không gian trạng thái có kích thước d_e . Vector $\mathbf{u}_{i,t}$ đóng vai trò là tín hiệu đầu vào của mô hình không gian trạng thái.

Mô hình hóa động lực học của trạng thái bệnh. Đối với mỗi kênh tiềm ẩn, động lực học của trạng thái bệnh được mô tả bởi hệ phương trình không gian trạng thái liên tục

$$\frac{d\mathbf{s}_{i,t}}{dt} = \mathbf{A}\mathbf{s}_{i,t} + \mathbf{B}\mathbf{u}_{i,t}, \quad (3.14)$$

trong đó $\mathbf{s}_{i,t}$ là trạng thái bệnh tiềm ẩn, còn $\mathbf{A}$ và $\mathbf{B}$ lần lượt là ma trận chuyển trạng

thái và ma trận điều khiển đầu vào. Đầu ra của mô hình trạng thái được xác định như sau:

$$\mathbf{y}_{i,t} = \mathbf{C}\mathbf{s}_{i,t}, \quad (3.15)$$

trong đó $\mathbf{C}$ là ma trận đọc trạng thái.

Cơ chế chuyển trạng thái phụ thuộc đầu vào. Trong thực tế, tốc độ tiến triển của HIV thay đổi giữa các bệnh nhân và giữa các giai đoạn điều trị. Vì vậy, thay vì sử dụng khoảng chuyển trạng thái cố định, mô hình học trực tiếp khoảng chuyển trạng thái từ biểu diễn hiện tại của bệnh nhân

$$\Delta_{i,t} = \text{Softplus}(\mathbf{W}_\Delta \tilde{\mathbf{h}}_{i,t} + \mathbf{b}_\Delta), \quad (3.16)$$

trong đó $\mathbf{W}_\Delta$ và $\mathbf{b}_\Delta$ là các tham số khả học.

Ma trận chuyển trạng thái sau khi rời rạc hóa được tính như sau:

$$\bar{\mathbf{A}}_{i,t} = \exp(\Delta_{i,t} \mathbf{A}), \quad (3.17)$$

và ma trận điều khiển đầu vào được xác định như sau:

$$\bar{\mathbf{B}}_{i,t} = \Delta_{i,t} \mathbf{B}. \quad (3.18)$$

Nhờ đó, mô hình có thể tự điều chỉnh động lực học của trạng thái bệnh theo từng bệnh nhân và từng thời điểm theo dõi.

Cập nhật trạng thái bệnh và xây dựng biểu diễn đầu ra. Trạng thái bệnh được cập nhật tuần tự theo cơ chế đệ quy như sau:

$$\mathbf{s}_{i,t} = \bar{\mathbf{A}}_{i,t} \mathbf{s}_{i,t-1} + \bar{\mathbf{B}}_{i,t} \mathbf{u}_{i,t}. \quad (3.19)$$

Sau khi cập nhật trạng thái, mô hình áp dụng cơ chế chọn lọc đặc trưng nhằm xác định những thành phần mang nhiều thông tin nhất đối với nhiệm vụ dự báo. Hệ số chọn lọc được tính như sau:

$$\mathbf{g}_{i,t} = \sigma(\mathbf{W}_g \tilde{\mathbf{h}}_{i,t} + \mathbf{b}_g), \quad (3.20)$$

trong đó $\sigma(\cdot)$ là hàm sigmoid. Đầu ra của mô hình không gian trạng thái được xác định như sau:

$$\mathbf{o}_{i,t} = \mathbf{W}_{\text{out}}(\mathbf{g}_{i,t} \odot \mathbf{y}_{i,t}), \quad (3.21)$$

trong đó $\odot$ là phép nhân từng phần tử và $\mathbf{W}_{\text{out}}$ là ma trận chiếu đầu ra.

Biểu diễn $\mathbf{o}_{i,t}$ là đầu ra cuối cùng của mô-đun học trạng thái bệnh tiềm ẩn. Biểu diễn này tổng hợp đồng thời ba nguồn thông tin gồm: (i) thông tin lâm sàng hiện tại của bệnh nhân; (ii) diễn biến bệnh được tích lũy qua nhiều lần theo dõi; và (iii) tri thức quần thể được học từ mô-đun bộ nhớ nguyên mẫu. Nhờ đó, $\mathbf{o}_{i,t}$ phản ánh toàn diện trạng thái bệnh tại thời điểm hiện tại và được sử dụng làm đầu vào chung cho cả nhánh dự báo sống sót và nhánh dự báo kết cục điều trị.

Đối với mỗi bệnh nhân, biểu diễn tại lần theo dõi cuối cùng được sử dụng làm biểu diễn tổng hợp của toàn bộ chuỗi theo dõi như sau:

$$\mathbf{o}_i = \mathbf{o}_{i,T_i}, \quad (3.22)$$

trong đó T_i là số lần theo dõi của bệnh nhân i . Biểu diễn này được sử dụng làm đầu vào chung cho nhánh dự báo sống sót và nhánh dự báo kết cục điều trị.

3.3.5. Thành phần dự báo đa nhiệm hỗ trợ ra quyết định lâm sàng

Biểu diễn trạng thái bệnh tiềm ẩn thu được từ thành phần học trạng thái bệnh là biểu diễn thống nhất của bệnh nhân sau khi tích hợp thông tin lâm sàng hiện tại, diễn biến bệnh theo thời gian và tri thức quần thể được truy xuất từ bộ nhớ nguyên mẫu. Biểu diễn này được sử dụng làm đầu vào chung cho hai nhánh dự báo gồm dự báo sống sót và dự báo kết cục điều trị. Từ biểu diễn trạng thái bệnh chung, mô hình đồng thời sinh các đầu ra tương ứng với từng nhiệm vụ. Nhánh dự báo sống sót được sử dụng để ước lượng nguy cơ và tiên lượng sống sót của bệnh nhân, trong khi nhánh dự báo kết cục điều trị sử dụng mô hình Bayes để dự báo kết cục điều trị và lượng hóa độ bất định của dự báo. Mặc dù cùng khai thác một biểu diễn trạng thái bệnh, mỗi nhánh sử dụng một tập tham số riêng để học các đặc trưng phù hợp với mục tiêu dự báo của mình.

Nhánh dự báo sống sót. Biểu diễn trạng thái bệnh tiềm ẩn $\mathbf{o}_i$ được sử dụng làm đầu vào cho nhánh dự báo sống sót. Nhánh này gồm ba thành phần bổ sung cho nhau nhằm mô tả toàn diện tiên lượng của bệnh nhân. Thành phần thứ nhất dự báo hàm nguy cơ theo từng khoảng thời gian rời rạc; thành phần thứ hai ước lượng thời gian sống sót kỳ vọng; và thành phần thứ ba tính toán điểm nguy cơ để phục vụ phân tầng bệnh nhân. Hàm nguy cơ được dự báo theo các khoảng thời gian xác định trước như sau

$$\hat{\mathbf{h}}_i = f_{\text{hazard}}(\mathbf{o}_i), \quad (3.23)$$

trong đó $\hat{\mathbf{h}}_i \in \mathbb{R}^{R \times T_b}$, với R là số loại biến cố cần dự báo và T_b là số khoảng thời gian được chia trước. Từ các giá trị hàm nguy cơ dự báo, xác suất sống sót đến thời điểm t được tính theo công thức sau:

$$\hat{S}_{i,t} = \prod_{k=1}^t (1 - \hat{h}_{i,k}). \quad (3.24)$$

Bên cạnh đó, mô hình còn dự báo trực tiếp thời gian sống sót kỳ vọng:

$$\hat{t}_i = f_{\text{time}}(\mathbf{o}_i),$$

và điểm nguy cơ:

$$r_i = f_{\text{risk}}(\mathbf{o}_i).$$

Ba thành phần này cung cấp các góc nhìn bổ sung cho quá trình tiên lượng. Trong đó, nhánh dự báo hàm nguy cơ mô tả xác suất xảy ra biến cố theo thời gian; nhánh dự

báo thời gian sống sót ước lượng thời điểm kỳ vọng xảy ra biến cố; còn điểm nguy cơ hỗ trợ phân tầng bệnh nhân theo mức độ nguy cơ, từ đó cung cấp cơ sở định lượng cho việc theo dõi và ưu tiên các can thiệp lâm sàng.

Nhánh dự báo kết cục điều trị theo tiếp cận Bayes. Ngoài dự báo sống sót, biểu diễn trạng thái bệnh còn được sử dụng để dự báo kết cục điều trị đối với từng bệnh nhân. Thay vì coi tác động điều trị là một đại lượng duy nhất, luận án phân rã tác động điều trị thành ba thành phần gồm: ảnh hưởng ở mức quần thể, ảnh hưởng ở mức nhóm bệnh nhân có đặc điểm tương đồng và ảnh hưởng ở mức cá thể. Cách tiếp cận phân cấp này cho phép mô hình đồng thời học được các quy luật điều trị ở nhiều mức độ khác nhau. Ảnh hưởng ở mức quần thể được biểu diễn bằng một véc tơ tham số khả học: $e_{\text{pop}} \in \mathbb{R}^D$, trong đó D là số phương án điều trị. Đối với mức nhóm, trước tiên biểu diễn trạng thái bệnh được ánh xạ sang xác suất thuộc các nhóm tiềm ẩn: $\pi_i = f_{\text{sub}}(\mathbf{o}_i)$, với $\pi_i \in \mathbb{R}^G$ là vector xác suất thuộc G nhóm tiềm ẩn. Một bệnh nhân có thể đồng thời thuộc nhiều nhóm với các xác suất khác nhau. Từ đó, ảnh hưởng ở mức nhóm được xác định theo công thức:

$$e_{\text{sub},i} = \pi_i E_{\text{sub}}, \quad (3.25)$$

trong đó $E_{\text{sub}} \in \mathbb{R}^{G \times D}$ là ma trận chứa tác động điều trị của từng nhóm tiềm ẩn. Ảnh hưởng ở mức cá thể được ước lượng thông qua một lớp tuyến tính Bayes:

$$e_{\text{ind},i} = f_{\text{Bayes}}(\mathbf{o}_i), \quad (3.26)$$

trong đó trọng số của mô hình không được xem là các giá trị cố định mà được mô hình hóa dưới dạng phân phối xác suất. Phân phối tiên nghiệm của trọng số được giả thiết là phân phối Gaussian:

$$p(w) = \mathcal{N}(0, \sigma_p^2 I), \quad (3.27)$$

trong khi phân phối hậu nghiệm được xấp xỉ bằng suy luận biến phân:

$$q(w) = \mathcal{N}(\mu, \text{diag}(\sigma^2)). \quad (3.28)$$

Trong quá trình huấn luyện, trọng số được lấy mẫu bằng kỹ thuật tái tham số hóa:

$$w = \mu + \sigma \odot \epsilon, \quad (3.29)$$

với $\epsilon \sim \mathcal{N}(0, I)$. Tác động điều trị cuối cùng được xác định bằng tổng của ba thành phần như sau:

$$e_i = e_{\text{pop}} + e_{\text{sub},i} + e_{\text{ind},i}. \quad (3.30)$$

Ba thành phần trên được giả định đóng góp theo cơ chế cộng vào tác động điều trị tổng thể, qua đó cho phép mô hình đồng thời khai thác thông tin ở mức quần thể, nhóm bệnh nhân và cá thể. Đối với phương án điều trị quan sát được, biểu diễn trạng thái bệnh $\mathbf{o}_i$, tác động điều trị tổng hợp $\mathbf{e}_i$ và vector nhúng của phương án điều trị được ghép lại để dự báo kết cục điều trị:

$$\hat{y}_i = f_{\text{out}}([\mathbf{o}_i; \mathbf{e}_i; \text{emb}(d_i)]), \quad (3.31)$$

trong đó $\text{emb}(d_i)$ là véc tơ nhúng của phương án điều trị d_i đã được áp dụng cho bệnh nhân i . Nhánh dự báo Bayes được tối ưu bằng cách kết hợp sai số dự báo với thành phần điều chuẩn Kullback–Leibler như sau:

$$\mathcal{L}_{\text{Bayes}} = \mathcal{L}_{\text{treat}} + \lambda_{\text{KL}} \mathcal{L}_{\text{KL}}, \quad (3.32)$$

trong đó $\mathcal{L}_{\text{KL}} = D_{\text{KL}}(q(w) \parallel p(w))$. $D_{\text{KL}}(\cdot \parallel \cdot)$ là độ lệch Kullback–Leibler giữa phân phối hậu nghiệm xấp xỉ $q(w)$ và phân phối tiên nghiệm $p(w)$ của các tham số Bayes. Thành phần này đóng vai trò điều chuẩn, giúp phân phối hậu nghiệm không lệch quá xa phân phối tiên nghiệm, từ đó hạn chế hiện tượng quá khớp và nâng cao khả năng khái quát hóa của mô hình.

Trong giai đoạn suy luận, mô hình lấy nhiều mẫu từ phân phối hậu nghiệm để xây dựng phân phối dự báo. Giá trị kỳ vọng của kết quả dự báo được tính theo công thức sau:

$$\bar{y} = \frac{1}{S} \sum_{s=1}^S \hat{y}^{(s)}, \quad (3.33)$$

trong khi độ bất định của dự báo được xác định bởi công thức sau:

$$\text{Var}(\hat{y}) = \frac{1}{S} \sum_{s=1}^S \left(\hat{y}^{(s)} - \bar{y} \right)^2. \quad (3.34)$$

Giá trị trung bình phản ánh kết cục điều trị kỳ vọng, còn phương sai dự báo thể hiện mức độ tin cậy của dự báo. Việc lượng hóa bất định giúp hỗ trợ bác sĩ đánh giá mức độ chắc chắn của từng khuyến nghị điều trị trước khi đưa ra quyết định.

Tối ưu hóa học đa nhiệm. Hai nhánh dự báo được huấn luyện đồng thời thông qua một hàm mục tiêu tổng hợp. Mặc dù thực hiện các nhiệm vụ khác nhau, cả hai nhánh cùng chia sẻ biểu diễn trạng thái bệnh được học từ các mô-đun trước. Vì vậy, gradient sinh ra từ cả hai nhiệm vụ cùng tham gia cập nhật biểu diễn chung, trong khi các lớp đầu ra vẫn được tối ưu độc lập cho từng nhiệm vụ. Hàm mất mát tổng thể của mô hình được xác định như sau:

$$\mathcal{L} = \mathcal{L}_{\text{surv}} + \lambda_{\text{rank}} \mathcal{L}_{\text{rank}} + \lambda_{\text{treat}} (\mathcal{L}_{\text{treat}} + \lambda_{\text{KL}} \mathcal{L}_{\text{KL}}), \quad (3.35)$$

trong đó λ_{rank} , λ_{treat} và λ_{KL} là các hệ số cân bằng giữa các thành phần của hàm mục tiêu. Trong đó, $\mathcal{L}_{\text{surv}}$ giám sát nhiệm vụ dự báo sống sót thông qua cực đại hóa khả năng xảy ra của các thời điểm biến cố quan sát được; $\mathcal{L}_{\text{rank}}$ tăng cường khả năng phân tầng nguy cơ bằng cách khuyến khích mô hình gán điểm nguy cơ cao hơn cho các bệnh nhân xảy ra biến cố sớm; $\mathcal{L}_{\text{treat}}$ giảm sai số dự báo kết cục điều trị; còn $\mathcal{L}_{\text{KL}}$ đóng vai trò điều chuẩn phân phối hậu nghiệm của các tham số Bayes.

Trong quá trình huấn luyện, toàn bộ các thành phần của hàm mất mát được tối ưu

đồng thời bằng thuật toán lan truyền ngược. Nhờ sử dụng chung biểu diễn trạng thái bệnh, thông tin thu được từ nhiệm vụ dự báo sống sót và dự báo kết cục điều trị có thể bổ sung lẫn nhau trong quá trình học biểu diễn, đồng thời vẫn bảo đảm mỗi nhiệm vụ có cơ chế dự báo chuyên biệt phù hợp với mục tiêu riêng.

3.3.6. Giải thích kết quả dự báo bằng SHAP (Temporal Explainability Analysis)

Để nâng cao khả năng giải thích, mô hình sử dụng phương pháp *SHapley Additive exPlanations* (SHAP) [89] nhằm lượng hóa mức độ đóng góp của từng đặc trưng đầu vào đối với các kết quả dự báo của mô hình. Phương pháp này được áp dụng sau khi mô hình đã được huấn luyện nên không làm thay đổi cấu trúc hay các tham số của mạng.

Giả sử mô hình dự báo được biểu diễn bởi hàm $f(\mathbf{x})$, trong đó $\mathbf{x}$ là véc tor đặc trưng của bệnh nhân. Giá trị SHAP của đặc trưng thứ i được xác định theo giá trị Shapley

$$\phi_i = \sum_{S \subseteq F \setminus \{i\}} \frac{|S|!(|F| - |S| - 1)!}{|F|!} [f(S \cup \{i\}) - f(S)], \quad (3.36)$$

trong đó F là tập các đặc trưng đầu vào và S là một tập con của F . Theo tính chất cộng tính của SHAP, đầu ra của mô hình được biểu diễn dưới dạng

$$f(\mathbf{x}) = \phi_0 + \sum_{i=1}^M \phi_i, \quad (3.37)$$

trong đó ϕ_0 là giá trị dự báo cơ sở và ϕ_i biểu diễn mức đóng góp của đặc trưng thứ i đối với kết quả dự báo.

SHAP trong mô hình đề xuất được áp dụng riêng cho từng nhánh dự báo của mô hình. Đối với nhánh dự báo sống sót, các giá trị SHAP được tính từ điểm nguy cơ dự báo nhằm xác định mức độ ảnh hưởng của từng đặc trưng đầu vào đến nguy cơ xảy ra biến cố. Đối với nhánh dự báo kết cục điều trị, SHAP được áp dụng trên kết quả dự báo của mô hình Bayes để phân tích mức độ ảnh hưởng của từng đặc trưng đến kết cục điều trị của từng phương án điều trị. Các giá trị SHAP của các nhóm đặc trưng nhân khẩu học, lâm sàng và điều trị được phân tích trong cùng một không gian giải thích, qua đó cho phép đánh giá đồng thời vai trò của tất cả các đặc trưng đối với kết quả dự báo.

Đối với dữ liệu theo dõi dọc, giá trị SHAP được tính riêng tại từng lần tái khám cho từng đặc trưng đầu vào. Giá trị SHAP của đặc trưng thứ i tại lần tái khám thứ t được ký hiệu là $\phi_{t,i}$. Để đánh giá mức độ ảnh hưởng của từng đặc trưng trong toàn bộ quá trình theo dõi, các giá trị SHAP được tổng hợp bằng trung bình giá trị tuyệt đối theo thời gian như sau:

$$\overline{|\phi_i|} = \frac{1}{T} \sum_{t=1}^T |\phi_{t,i}|, \quad (3.38)$$

trong đó T là tổng số lần theo dõi của bệnh nhân. Đại lượng này phản ánh mức độ ảnh hưởng trung bình của đặc trưng i đối với kết quả dự báo trong toàn bộ quá trình tiến triển bệnh. Ngoài ra, việc phân tích chuỗi $\{\phi_{t,i}\}_{t=1}^T$ còn cho phép quan sát sự thay đổi vai trò của từng đặc trưng theo thời gian, qua đó phản ánh sự thay đổi của các yếu tố nguy cơ và đáp ứng điều trị trong suốt quá trình theo dõi.

Các giá trị SHAP được sử dụng để xác định các yếu tố nguy cơ, các yếu tố bảo vệ và mức độ ảnh hưởng của từng đặc trưng đối với kết quả dự báo. Đồng thời, luận án xây dựng báo cáo giải thích tự động kết hợp giữa kết quả dự báo và các giá trị SHAP, bao gồm nguy cơ sống sót, xác suất sống sót, thời gian sống sót kỳ vọng, kết quả dự báo điều trị, các yếu tố nguy cơ, các yếu tố bảo vệ và sự thay đổi tầm quan trọng của các đặc trưng theo thời gian. Các thông tin này cung cấp cơ sở minh bạch để hỗ trợ bác sĩ diễn giải các dự báo của mô hình và nâng cao khả năng giải thích trong quá trình hỗ trợ ra quyết định lâm sàng.

3.3.7. Chiến lược huấn luyện mô hình

Do mô hình được tối ưu đồng thời cho nhiều mục tiêu học, nghiên cứu áp dụng chiến lược học theo chương trình (*curriculum learning*) nhằm tăng dần độ phức tạp của quá trình huấn luyện thay vì tối ưu tất cả các mục tiêu ngay từ đầu.

Ở giai đoạn đầu, mô hình chỉ được huấn luyện với mục tiêu dự báo sống sót nhằm học biểu diễn ổn định của bệnh nhân. Sang giai đoạn thứ hai, các hàm mất mát xếp hạng và hiệu chỉnh xác suất được bổ sung để cải thiện khả năng phân biệt nguy cơ và độ hiệu chỉnh của xác suất dự báo. Cuối cùng, mục tiêu dự báo kết cục điều trị được đưa vào giai đoạn huấn luyện cuối. Việc bổ sung tuần tự các mục tiêu giúp các nhiệm vụ phụ hỗ trợ dần cho nhiệm vụ dự báo sống sót, đồng thời giảm sự can nhiễu giữa các mục tiêu trong giai đoạn đầu của quá trình tối ưu.

Trong mỗi lần lặp huấn luyện, mô hình thực hiện lan truyền thuận, tính hàm mất mát tương ứng với giai đoạn của chiến lược học theo chương trình, sau đó lan truyền ngược để tính gradient và cập nhật tham số bằng thuật toán AdamW. Quá trình được lặp lại cho đến khi hội tụ. Trong suốt quá trình huấn luyện, mô hình sử dụng cơ chế điều chỉnh tốc độ học *ReduceLROnPlateau* và dừng sớm (*early stopping*) trên tập xác thực nhằm nâng cao tính ổn định và hạn chế hiện tượng quá khớp.

3.4. Thực nghiệm và kết quả

Thiết kế thực nghiệm được thực hiện trên hai bộ dữ liệu thử nghiệm lâm sàng của AIDS Clinical Trials Group (ACTG), bao gồm ACTG 388 và ACTG 175. Trong đó, ACTG 388 được lựa chọn là bộ dữ liệu thực nghiệm chính nhằm đánh giá khả năng học từ dữ liệu theo dõi dọc của mô hình đề xuất, trong khi ACTG 175 được sử dụng để kiểm chứng khả năng hoạt động của mô hình trên dữ liệu tĩnh, khi chỉ có thông tin ban đầu của bệnh nhân.

Hai bộ dữ liệu đại diện cho hai kịch bản khác nhau trong bài toán dự báo sống sót. ACTG 388 là bộ dữ liệu theo dõi dọc (longitudinal), trong đó mỗi bệnh nhân có nhiều lần tái khám với các thông tin lâm sàng như số lượng tế bào CD4 và tải lượng virus được ghi nhận theo thời gian. Ngược lại, ACTG 175 là bộ dữ liệu tĩnh (static), mỗi

bệnh nhân chỉ có một bộ đặc trưng nền (baseline) tại thời điểm bắt đầu nghiên cứu. Việc đánh giá trên cả hai bộ dữ liệu nhằm kiểm chứng rằng mô hình đề xuất không chỉ khai thác hiệu quả thông tin theo thời gian khi dữ liệu dọc được cung cấp, mà còn duy trì hiệu quả dự báo trong trường hợp chỉ có dữ liệu tĩnh, qua đó đánh giá khả năng khái quát hóa của mô hình trên các dạng dữ liệu lâm sàng có cấu trúc khác nhau.

3.4.1. Dữ liệu nghiên cứu

Nghiên cứu này sử dụng hai bộ dữ liệu thử nghiệm lâm sàng HIV công khai, mỗi bộ dữ liệu phản ánh một khía cạnh khác nhau của bài toán phân tích sống sót. ACTG388 [36] bao gồm các kết quả xét nghiệm và thông tin điều trị được ghi nhận theo chuỗi thời gian trong khi Bộ dữ liệu ACTG175 [35] cung cấp các đặc điểm nền tảng của bệnh nhân tại thời điểm bắt đầu điều trị,

Bộ dữ liệu ACTG388 theo chuỗi thời gian. Để đánh giá khả năng mô hình hóa dữ liệu theo thời gian của mô hình đề xuất, nghiên cứu sử dụng phiên bản xử lý của bộ dữ liệu thử nghiệm ACTG388. Thử nghiệm lâm sàng ban đầu bao gồm 517 bệnh nhân nhiễm HIV-1 được phân bổ vào ba phác đồ điều trị HAART khác nhau, với kết quả được báo cáo bởi Fischl và cộng sự [66].

Phiên bản dữ liệu đã được chuẩn hóa cho phân tích sống sót và dữ liệu chuỗi thời gian được cung cấp thông qua kho dữ liệu UTHealth [67], với quy trình tiền xử lý được mô tả trong nghiên cứu của Park và Wu [65]. Sau khi loại bỏ các bệnh nhân thiếu thông tin điều trị hoặc thời gian biến cố, 511 bệnh nhân được giữ lại cho phân tích.

Bộ dữ liệu này bao gồm bốn thành phần chính.

1. *Dữ liệu xét nghiệm theo thời gian.* Bao gồm các phép đo lặp lại của các chỉ số miễn dịch như CD4 và CD8. Các phép đo được thực hiện tại thời điểm ban đầu, tuần thứ 4, tuần thứ 8 và sau đó mỗi 8 tuần. Các biến bao gồm *PATID*, *WEEK*, *CD4*, *CD4PCT*, *CD8*, *CD8PCT* và ngày lấy mẫu (*SPECDT*).

2. *Đặc trưng nền và tải lượng virus.* Bao gồm các biến như *sex*, *race/ethnicity*, *weight*, các chỉ số CD4/CD8 trước điều trị, cũng như tải lượng virus HIV-1 (*COPYML*, *LOGRNA*). Ngoài ra, bộ dữ liệu còn chứa thông tin về điều trị và các biến liên quan đến biến cố sống sót.

3. *Dòng thời gian điều trị và biến cố.* Thành phần này cung cấp các mốc thời gian quan trọng như ngày bắt đầu điều trị (*DOSE1DT*), ngày ngẫu nhiên hóa, ngày ngừng điều trị, ngày đạt đáp ứng virus và ngày xảy ra biến cố. Thời gian sống sót của mỗi bệnh nhân được xác định từ thời điểm bắt đầu điều trị đến khi xảy ra biến cố hoặc bị kiểm duyệt.

4. *Nhóm điều trị.* Biến *TRTHOLD* cho biết bệnh nhân thuộc nhóm điều trị HAART nào trong thử nghiệm.

Trong 511 bệnh nhân, bộ dữ liệu bao gồm tổng cộng 6.856 quan sát theo thời gian, trung bình 13,4 quan sát cho mỗi bệnh nhân, với thời gian theo dõi trung vị là 24 tháng. Các biến cố lâm sàng bao gồm thất bại virus học (32,6%), ngừng điều trị (18,6%), tử vong (1,7%) và dữ liệu kiểm duyệt (47,1%).

Bộ dữ liệu thử nghiệm lâm sàng này cung cấp cấu trúc dữ liệu theo thời gian phong

phù, phù hợp để đánh giá các mô hình phân tích sống sót kết hợp thông tin nền và dữ liệu chuỗi thời gian.

Bảng 3.1: Tóm tắt các bộ dữ liệu sử dụng trong nghiên cứu

Bộ dữ liệu	Số bệnh nhân	Số quan sát	Thời gian theo dõi	Biến cố
ACTG175	2139	2139	theo dõi sống sót	Tử vong / AIDS
ACTG388	511	6856	24 tháng	Ngừng điều trị, tử vong

Bộ dữ liệu ACTG175. Bộ dữ liệu ACTG175 có nguồn gốc từ nghiên cứu thử nghiệm lâm sàng *AIDS Clinical Trials Group Study 175* và được công bố tại *UCI Machine Learning Repository* [42]. Đây là một thử nghiệm ngẫu nhiên có đối chứng nhằm so sánh hiệu quả của bốn phác đồ điều trị kháng virus ở bệnh nhân HIV với số lượng tế bào CD4 trong khoảng 200–500 tế bào/mm³.

Thử nghiệm bao gồm 2.139 bệnh nhân và kết quả được công bố trên tạp chí *New England Journal of Medicine* năm 1996 [97]. Bộ dữ liệu bao gồm 23 biến đặc trưng, được chia thành ba nhóm chính.

1. *Đặc trưng nhân khẩu học.* Bao gồm các biến như *age* (tuổi bệnh nhân tại thời điểm bắt đầu điều trị), *wtkg* (cân nặng ban đầu), *gender* (giới tính), *race* (chủng tộc) và *karnof* (chỉ số Karnofsky phản ánh tình trạng sức khỏe tổng thể của bệnh nhân).

2. *Đặc trưng lâm sàng.* Bao gồm các biến như *hemo* (tình trạng bệnh máu khó đông), *drugs* (tiền sử sử dụng ma túy tiêm), *homo* (tiền sử quan hệ đồng giới), *oprior* (tiền sử điều trị kháng virus trước đó), cùng với các chỉ số miễn dịch như *CD40*, *CD420*, *CD80* và *CD820*.

3. *Đặc trưng điều trị.* Biến *trt* biểu thị nhóm điều trị được phân bổ cho bệnh nhân, bao gồm các phác đồ AZT, ddI, AZT+ddI và AZT+ddC.

4. *Biến kết quả sống sót.* Thông tin sống sót được mô tả thông qua biến *time*, biểu thị thời gian từ khi bắt đầu điều trị đến khi xảy ra biến cố hoặc bị kiểm duyệt, và biến *status* (hoặc *target*), trong đó giá trị 1 biểu thị xảy ra biến cố và 0 biểu thị dữ liệu bị kiểm duyệt.

Bộ dữ liệu này đã được sử dụng rộng rãi trong các nghiên cứu học máy y sinh nhằm đánh giá các mô hình dự báo kết quả điều trị và khả năng sống sót của bệnh nhân HIV.

3.4.2. Tiền xử lý dữ liệu và xây dựng đặc trưng

Để đảm bảo tính nhất quán giữa hai bộ dữ liệu và tránh rò rỉ thông tin, một quy trình tiền xử lý thống nhất được áp dụng. Việc chia dữ liệu được thực hiện ở cấp độ bệnh nhân thay vì cấp độ quan sát, nhằm đảm bảo rằng các phép đo của cùng một bệnh nhân không xuất hiện đồng thời trong các tập huấn luyện, kiểm định và kiểm tra.

Dữ liệu thiếu được xử lý theo cấu trúc của từng bộ dữ liệu. Đối với ACTG388, các biến xét nghiệm theo thời gian được điền bằng phương pháp Last Observation Carried Forward (LOCF). Trong trường hợp không có giá trị trước đó, giá trị trung vị của tập huấn luyện được sử dụng. Đối với ACTG175, dữ liệu liên tục được điền bằng giá trị trung vị và dữ liệu phân loại được điền bằng giá trị xuất hiện nhiều nhất.

Các biến phân loại như nhóm điều trị và chủng tộc được mã hóa bằng phương pháp

one-hot encoding. Các biến nhị phân như giới tính và tình trạng ngừng điều trị được chuẩn hóa để đảm bảo tính nhất quán giữa hai bộ dữ liệu. Các biến liên tục được chuẩn hóa bằng phương pháp chuẩn hóa z-score dựa trên thống kê của tập huấn luyện.

Đối với ACTG388, do dữ liệu được thu thập theo các lần khám, chỉ số thời gian sống sót và chỉ số lần khám được chuẩn hóa về khoảng $[0,1]$ nhằm hỗ trợ việc mô hình hóa theo chuỗi thời gian.

Bên cạnh đó, một số đặc trưng mới được xây dựng dựa trên ý nghĩa lâm sàng. Cụ thể, tỷ lệ CD4/CD8 được tính toán để phản ánh trạng thái miễn dịch của bệnh nhân. Ngoài ra, một chỉ số nguy cơ ban đầu được xây dựng dựa trên số lượng CD4, tải lượng virus và tuổi của bệnh nhân. Đối với ACTG175, do chỉ có một phép đo theo dõi tại tuần thứ 20, sự thay đổi của CD4 giữa thời điểm ban đầu và tuần thứ 20 cũng được sử dụng làm đặc trưng bổ sung. Trong khi đó, ACTG388 cung cấp chuỗi quan sát liên tục, cho phép tính toán sự thay đổi CD4 giữa các lần khám và sử dụng các biến điều trị theo thời gian.

Các biến liên quan đến điều trị được đưa vào mô hình như các biến tiên lượng, thay vì được xem là biến can thiệp nhân quả.

3.4.3. Chiến lược chia dữ liệu

Đối với cả hai bộ dữ liệu, nghiên cứu chia dữ liệu thành ba phần: 64% cho tập huấn luyện, 16% cho tập xác thực, và 20% cho tập kiểm tra. Tập xác thực được sử dụng để điều chỉnh siêu tham số và lựa chọn mô hình phù hợp, trong khi các kết quả cuối cùng được báo cáo trên tập kiểm tra độc lập.

3.4.4. Chỉ số đánh giá

Để đánh giá toàn diện hiệu quả dự báo của mô hình, nghiên cứu sử dụng các chỉ số phổ biến trong phân tích sống sót, bao gồm khả năng phân biệt, độ chính xác tổng thể, mức độ hiệu chuẩn và hiệu quả dự báo theo thời gian.

Chỉ số C-index được sử dụng để đánh giá khả năng phân biệt của mô hình, tức khả năng xếp hạng đúng mức độ nguy cơ giữa các bệnh nhân. Giá trị C-index càng cao cho thấy mô hình càng phân biệt tốt giữa các bệnh nhân có nguy cơ cao và thấp.

IBS được sử dụng để đánh giá độ chính xác tổng thể của xác suất sống sót dự báo trong toàn bộ thời gian theo dõi. Giá trị IBS càng thấp cho thấy dự báo càng chính xác.

Calibration MSE được sử dụng để đánh giá mức độ phù hợp giữa xác suất sống sót dự báo và kết quả sống sót thực tế, phản ánh mức độ hiệu chuẩn của mô hình. Giá trị càng thấp cho thấy mô hình càng đáng tin cậy trong việc ước lượng xác suất nguy cơ.

Time-dependent AUC được sử dụng để đánh giá khả năng phân biệt của mô hình tại các mốc thời gian cụ thể, qua đó phản ánh hiệu quả dự báo ngắn hạn và trung hạn của mô hình.

3.4.5. Các mô hình so sánh

Để đảm bảo tính khách quan và khả năng so sánh với các công trình phổ biến trong phân tích sống sót, nghiên cứu lựa chọn các mô hình nền tảng đại diện cho nhiều

nhóm phương pháp khác nhau. Đầu tiên mô hình thống kê cổ điển (Cox Proportional Hazards (CoxPH)) được sử dụng rộng rãi trong nghiên cứu y sinh do khả năng giải thích tốt và cấu trúc tuyến tính rõ ràng. Tiếp theo là DeepSurv, mở rộng Cox bằng mạng nơ-ron; DeepHit, mô hình hóa phân phối thời gian xảy ra biến cố theo cách rời rạc; và SurvTRACE, dựa trên kiến trúc Transformer cho mô hình phân tích sống sót.

Đối với ACTG 388, do dữ liệu có tính dọc theo thời gian, nghiên cứu bổ sung thêm GRU-Survival và LSTM-Survival nhằm phản ánh nhóm mô hình hồi quy tuần tự (RNN-based) thường được sử dụng trong phân tích dữ liệu y tế longitudinal.

3.4.6. Thiết lập thực nghiệm

Thực nghiệm được thực hiện theo chiến lược kiểm định chéo năm phần (5-fold cross-validation) ở cấp bệnh nhân nhằm tránh rò rỉ dữ liệu giữa các tập huấn luyện và đánh giá. Trong mỗi vòng lặp, mô hình được huấn luyện trên bốn fold, điều chỉnh siêu tham số trên tập validation nội bộ, và đánh giá trên fold còn lại. Kết quả cuối cùng được báo cáo dưới dạng giá trị trung bình trên năm fold.

Để đánh giá độ ổn định của mô hình, khoảng tin cậy 95% được ước lượng bằng bootstrap với 1,000 lần lấy mẫu lại ở cấp bệnh nhân. Kiểm định ý nghĩa thống kê giữa mô hình đề xuất và các mô hình đối sánh được thực hiện bằng kiểm định bootstrap ghép cặp (paired bootstrap test).

Về huấn luyện, mô hình sử dụng bộ tối ưu AdamW optimizer với learning rate ban đầu 1×10^{-4} và weight decay 1×10^{-5} . Gradient clipping với ngưỡng $\|g\| \leq 1.0$ được áp dụng nhằm hạn chế hiện tượng gradient explosion. Learning rate được điều chỉnh bằng ReduceLROnPlateau, đồng thời cơ chế early stopping với patience 10 epoch được sử dụng để tránh overfitting. Trọng số mô hình đạt hiệu năng tốt nhất trên tập validation được lưu lại để đánh giá cuối cùng.

3.4.7. Thực nghiệm chi tiết trên ACTG 388 (Longitudinal Cohort)

3.4.7.1. Hiệu năng tổng thể trên ACTG 388

Bảng 3.2, mô hình đề xuất đạt được sự cải thiện rõ rệt so với mô hình baseline tốt nhất là GRU-Survival trên tất cả các chỉ số đánh giá quan trọng. Cụ thể, chỉ số C-index đạt 0.824, cao hơn đáng kể so với 0.762 của GRU, cho thấy khả năng phân biệt rủi ro tốt hơn. Đồng thời, chỉ số Integrated Brier Score (IBS) giảm xuống còn 0.087, thấp hơn 22.3% so với baseline, phản ánh độ chính xác tổng thể cao hơn trong dự báo xác suất sống sót theo thời gian. Ngoài ra, sai số hiệu chỉnh (Calibration MSE) cũng giảm mạnh, cho thấy mô hình không chỉ dự báo chính xác mà còn duy trì được sự hiệu chỉnh tốt giữa xác suất dự báo và xác suất thực tế. Những cải thiện đồng thời trên nhiều chỉ số này cho thấy mô hình không học theo nhiễu dữ liệu mà thực sự nắm bắt được các động lực tiến triển bệnh.

Bảng 3.2: Các chỉ số hiệu năng chính trên ACTG 388

Chỉ số	MH đề xuất	Baseline tốt nhất (GRU)	Mức cải thiện
C-index	0.824	0.762	+8.1%
Integrated Brier Score (IBS)	0.087	0.112	-22.3%
Calibration MSE	0.012	0.021	-42.9%
AUC phụ thuộc thời gian (24 tháng)	0.801	0.745	+7.5%

Tiếp theo, Bảng 3.3 cho thấy các cải thiện về mặt mô hình học không chỉ mang ý nghĩa thống kê mà còn chuyển hóa thành lợi ích thực tiễn trong bối cảnh lâm sàng. Cụ thể, nhóm bệnh nhân được hỗ trợ bởi mô hình đạt được tỷ lệ tuân thủ điều trị cao hơn, thời gian đạt ức chế tải lượng virus nhanh hơn, mức tăng CD4 cao hơn và tỷ lệ ngừng điều trị thấp hơn. Đáng chú ý, thời gian đạt ức chế virus giảm 6.4 tuần, đây là một cải thiện có ý nghĩa quan trọng trong kiểm soát bệnh HIV. Đồng thời, mức tăng CD4 cao hơn (+44 cells/ μ L) cho thấy sự phục hồi miễn dịch tốt hơn. Những kết quả này phản ánh tính khả thi và giá trị ứng dụng của mô hình trong hỗ trợ quyết định điều trị.

Bảng 3.3: Tác động lâm sàng của mô hình trên ACTG 388

Chỉ số lâm sàng	MH đề xuất	Chăm sóc tiêu chuẩn	Mức cải thiện
Tỷ lệ tuân thủ điều trị	74.1%	61.8%	+12.3%
Thời gian đạt ức chế virus	16.3 tuần	22.7 tuần	-6.4 tuần
Mức tăng CD4 sau 24 tháng	+168 cells/ μ L	+124 cells/ μ L	+44 cells/ μ L
Tỷ lệ ngừng điều trị	14.2%	22.8%	-8.6%

Cuối cùng, Bảng 3.4 làm rõ nguyên nhân cốt lõi giúp mô hình đạt hiệu năng cao. Việc khai thác thông tin theo chuỗi thời gian (longitudinal data) giúp cải thiện đáng kể khả năng phân biệt, đặc biệt trong giai đoạn sớm từ 3 đến 6 tháng – một giai đoạn quan trọng trong điều trị HIV. Ngoài ra, mô hình vẫn duy trì độ ổn định cao ngay cả khi số lượng mốc thời gian tăng lên, với mức suy giảm hiệu năng không đáng kể (<0.7%). Điều này chứng minh rằng kiến trúc đề xuất có khả năng mở rộng tốt và phù hợp với dữ liệu lâm sàng dài hạn.

Bảng 3.4: Lợi thế của mô hình theo thời gian trên ACTG 388

Khía cạnh	MH đề xuất	Baseline	Diễn giải
C-index so với dữ liệu tĩnh	+5.4%	0.770 $\rightarrow$ 0.824	Lợi ích từ dữ liệu dọc (longitudinal)
Dự báo sớm (3–6 tháng)	AUC = 0.826	Giai đoạn quan trọng	Hiệu năng rất cao
Độ bền theo chuỗi thời gian	<0.7% suy giảm	Tới 21+ mốc thời gian	Khả năng mở rộng lâm sàng

Tổng hợp lại, các kết quả trên bộ dữ liệu ACTG 388 cho thấy mô hình MH đề xuất không chỉ cải thiện đáng kể về mặt kỹ thuật (C-index, IBS, AUC) mà còn mang lại giá trị thực tiễn rõ ràng trong điều trị HIV. Đặc biệt, khả năng khai thác thông tin theo thời gian đóng vai trò then chốt trong việc nâng cao hiệu năng dự báo, qua đó khẳng định tính phù hợp của hướng tiếp cận học sâu trên dữ liệu dọc trong bài toán phân tích

sống sót.

3.4.7.2. So sánh hiệu năng với các mô hình cơ sở

Kết quả so sánh trong Bảng 3.5 cho thấy mô hình đề xuất đạt hiệu năng tốt nhất và nhất quán trên toàn bộ các chỉ số đánh giá, bao gồm khả năng phân biệt nguy cơ (C-index), độ chính xác dự báo tổng thể (Integrated Brier Score – IBS), mức độ hiệu chỉnh xác suất (Calibration MSE) và năng lực dự báo nguy cơ dài hạn (AUC tại 24 tháng). Điều này cho thấy mô hình không chỉ cải thiện ở một khía cạnh đơn lẻ mà đạt được sự cân bằng đồng thời giữa khả năng phân biệt, hiệu chỉnh và độ ổn định của dự báo.

Xét về khả năng phân biệt nguy cơ, mô hình đề xuất đạt C-index = 0.824, cao hơn rõ rệt so với Cox PH (0.634), DeepSurv (0.672), DeepHit (0.708), SurvTRACE (0.739), LSTM-Survival (0.756) và GRU-Survival (0.762). Kết quả này cho thấy mô hình có khả năng sắp xếp chính xác thứ tự nguy cơ giữa các bệnh nhân tốt hơn đáng kể so với các phương pháp đối sánh. Đặc biệt, mức cải thiện lớn so với Cox PH phản ánh rằng các giả định tuyến tính và proportional hazards trong các mô hình truyền thống chưa đủ để mô tả đầy đủ động lực tiến triển phức tạp của HIV theo thời gian.

Đối với độ chính xác tổng thể, mô hình đề xuất đạt IBS = 0.087, thấp nhất trong tất cả các mô hình được đánh giá. So với GRU-Survival (0.112), mức giảm đạt khoảng 22.3%, trong khi so với Cox PH (0.172), mức giảm lên tới 49.4%. Điều này cho thấy xác suất sống sót dự báo từ mô hình không chỉ phản ánh đúng thứ tự nguy cơ mà còn bám sát hơn với diễn biến thực tế của dữ liệu trên toàn bộ khoảng thời gian theo dõi.

Ở khía cạnh hiệu chỉnh xác suất, mô hình đề xuất đạt Calibration MSE = 0.012, thấp hơn đáng kể so với các phương pháp còn lại. Đây là một kết quả có ý nghĩa thực tiễn quan trọng, bởi trong bối cảnh lâm sàng, một mô hình có khả năng phân biệt tốt nhưng hiệu chỉnh kém vẫn có thể dẫn đến các quyết định điều trị sai lệch. Kết quả này cho thấy xác suất nguy cơ do mô hình sinh ra có độ tin cậy cao hơn và phù hợp hơn cho các ứng dụng hỗ trợ ra quyết định.

Đối với dự báo nguy cơ dài hạn, mô hình đề xuất đạt AUC tại 24 tháng = 0.801, vượt trội so với GRU-Survival (0.745) và LSTM-Survival (0.739). Kết quả này cho thấy mô hình không chỉ hoạt động tốt trong các khoảng thời gian gần mà còn duy trì được khả năng dự báo hiệu quả tại các mốc thời gian xa hơn, nơi bất định tích lũy thường gia tăng đáng kể.

Xu hướng cải thiện hiệu năng khi chuyển từ Cox PH sang DeepSurv, DeepHit, SurvTRACE và tiếp tục sang các mô hình chuỗi thời gian như LSTM-Survival và GRU-Survival cho thấy việc khai thác cấu trúc phi tuyến và thông tin động theo thời gian mang lại lợi ích rõ rệt trong bài toán survival prediction. Tuy nhiên, việc mô hình đề xuất tiếp tục vượt qua cả các mô hình tuần tự mạnh cho thấy chỉ sử dụng cơ chế học tuần tự đơn thuần vẫn chưa đủ để biểu diễn đầy đủ đặc tính phức tạp của dữ liệu lâm sàng theo thời gian.

Một nguyên nhân khả dĩ là do dữ liệu ACTG 388 phản ánh quá trình tiến triển bệnh mang tính tích lũy, trong đó trạng thái nguy cơ hiện tại không chỉ phụ thuộc vào quan sát gần nhất mà còn chịu ảnh hưởng của các mẫu tiến triển dài hạn và các mối quan

hệ lâm sàng tiềm ẩn giữa các giai đoạn theo dõi. Các mô hình recurrent như LSTM và GRU có khả năng học phụ thuộc tuần tự nhưng vẫn có thể gặp hạn chế trong việc duy trì thông tin quan trọng qua chuỗi dài hoặc khi cần mô hình hóa động lực trạng thái phức tạp.

Trong khi đó, mô hình đề xuất kết hợp cơ chế temporal memory nhằm lưu giữ các nguyên mẫu lâm sàng quan trọng từ các trường hợp tương tự cùng với SSSM để mô hình hóa quá trình tiến triển bệnh liên tục theo thời gian. Sự kết hợp này cho phép mô hình khai thác đồng thời cả thông tin cục bộ và động lực dài hạn, từ đó cải thiện chất lượng biểu diễn và hiệu năng dự báo một cách nhất quán trên toàn bộ các chỉ số đánh giá.

Bảng 3.5: So sánh hiệu năng các mô hình phân tích sống sót trên dữ liệu ACTG 388 (dữ liệu longitudinal).

Mô hình	C-index	CI 95%	IBS	CI 95%	Calib. MSE	CI 95%	AUC (24 tháng)
Cox PH	0.634	(0.608–0.660)	0.172	(0.164–0.180)	0.041	(0.037–0.045)	0.612
DeepSurv	0.672	(0.647–0.697)	0.153	(0.146–0.160)	0.033	(0.030–0.036)	0.649
DeepHit	0.708	(0.684–0.732)	0.135	(0.128–0.142)	0.028	(0.025–0.031)	0.684
SurvTRACE	0.739	(0.716–0.762)	0.121	(0.115–0.127)	0.024	(0.021–0.027)	0.718
LSTM-Survival	0.756	(0.734–0.778)	0.115	(0.109–0.121)	0.022	(0.019–0.025)	0.739
GRU-Survival	0.762	(0.740–0.784)	0.112	(0.106–0.118)	0.021	(0.018–0.024)	0.745
MH đề xuất	0.824	(0.805–0.843)	0.087	(0.082–0.092)	0.012	(0.010–0.014)	0.801

3.4.7.3. Kiểm định ý nghĩa thống kê bằng DeLong

Để đánh giá mức độ ý nghĩa thống kê của các cải thiện về hiệu năng, nghiên cứu sử dụng kiểm định DeLong nhằm so sánh cặp giữa các mô hình dựa trên AUC. Kết quả được trình bày trong Bảng 3.6 cho thấy các cải thiện của mô hình đề xuất so với các phương pháp baseline đều có ý nghĩa thống kê rõ rệt.

Cụ thể, phần lớn các so sánh cặp liên quan đến mô hình đề xuất đều cho giá trị p rất nhỏ (< 0.001), và nhiều trường hợp đạt mức < 0.0001 , ngay cả khi so với các mô hình mạnh như GRU-Survival và SurvTRACE. Điều này cho thấy các cải thiện quan sát được không phải là do dao động ngẫu nhiên của dữ liệu mà phản ánh sự khác biệt thực sự trong khả năng dự báo của mô hình. Đặc biệt, việc duy trì mức ý nghĩa cao khi so sánh với các mô hình học sâu tiên tiến cho thấy tính ổn định và khả năng tổng quát hóa của phương pháp đề xuất.

Ngoài ra, có thể nhận thấy một số cặp mô hình baseline (ví dụ giữa LSTM-Survival và GRU-Survival) không có sự khác biệt có ý nghĩa thống kê ($p > 0.05$), cho thấy giới hạn của các kiến trúc RNN truyền thống trong việc cải thiện hiệu năng trên dữ liệu này. Ngược lại, mô hình đề xuất tạo ra sự khác biệt có ý nghĩa thống kê so với hầu hết các mô hình, củng cố giả thuyết rằng việc kết hợp attention và mô hình hóa theo chuỗi thời gian giúp nâng cao đáng kể khả năng dự báo.

Bảng 3.6: So sánh thống kê cặp giữa các mô hình (giá trị p) trên ACTG 388.

Mô hình	Cox PH	DeepSurv	DeepHit	SurvTRACE	LSTM-Surv	GRU-Surv	MH đề xuất
Cox PH	1	0.000145	<0.0001	<0.0001	<0.0001	<0.0001	<0.0001
DeepSurv	0.000145	1	0.000318	<0.0001	<0.0001	<0.0001	<0.0001
DeepHit	<0.0001	0.000318	1	0.001955	<0.001	<0.001	<0.0001
SurvTRACE	<0.0001	<0.0001	0.001955	1	0.089131	0.021448	<0.0001
LSTM-Survival	<0.0001	<0.0001	<0.001	0.089131	1	0.548506	<0.001
GRU-Survival	<0.0001	<0.0001	<0.001	0.021448	0.548506	1	<0.001
NH đề xuất	<0.0001	<0.0001	<0.0001	<0.0001	<0.001	<0.001	1

3.4.7.4. Đặc điểm quần thể nghiên cứu theo phân tầng CD4

Bảng 3.7 trình bày kết quả phân tích sống sót theo các nhóm phân tầng CD4 tại thời điểm ban đầu. Có thể quan sát rõ ràng một xu hướng phân tầng (gradient) nhất quán giữa các nhóm. Cụ thể, nhóm bệnh nhân có CD4 < 50 cells/ μ L thể hiện xác suất sống sót thấp hơn đáng kể tại tất cả các mốc thời gian, đồng thời thời gian xảy ra biến cố cũng ngắn hơn rõ rệt.

Hiệu năng sống sót được cải thiện theo từng mức CD4, với nhóm 50–200 cells/ μ L đạt kết quả trung gian và nhóm >200 cells/ μ L có kết quả tốt nhất. Xu hướng này hoàn toàn phù hợp với đặc điểm dịch tễ học đã được ghi nhận trong thử nghiệm ACTG 388, trong đó tình trạng miễn dịch ban đầu đóng vai trò quyết định đối với tiên lượng bệnh.

Kết quả kiểm định log-rank cho thấy sự khác biệt giữa các nhóm là có ý nghĩa thống kê rất cao ($\chi^2 = 31.7$, $p < 0.001$), và tất cả các so sánh cặp đều duy trì ý nghĩa thống kê ($p < 0.01$). Ngoài ra, tỷ số nguy cơ (hazard ratio) giữa nhóm CD4 < 50 và nhóm >200 cells/ μ L đạt giá trị cao (HR = 2.84, 95% CI: 2.12–3.81), cho thấy nguy cơ xảy ra biến cố ở nhóm CD4 thấp cao gấp gần 3 lần so với nhóm có CD4 cao. Kết quả này nhấn mạnh tầm quan trọng lâm sàng của phân tầng CD4 trong dự báo tiến triển bệnh HIV.

Bảng 3.7: Kết quả sống sót theo phân tầng CD4 ban đầu (ACTG 388)

Phân tầng CD4	N	Tỷ lệ	6 tháng	CI 95%	12 tháng	CI 95%	24 tháng	CI 95%	Thời gian trung vị
< 50 cells/ μ L	127	24.90%	0.732	(0.651–0.813)	0.583	(0.498–0.668)	0.441	(0.355–0.527)	14.2 tháng
50–200 cells/ μ L	226	44.20%	0.823	(0.769–0.877)	0.721	(0.661–0.781)	0.634	(0.571–0.697)	21.7 tháng
> 200 cells/ μ L	158	30.90%	0.886	(0.834–0.938)	0.823	(0.758–0.888)	0.759	(0.684–0.834)	Chưa đạt

3.4.7.5. Phân tích loại bỏ thành phần (Ablation Study)

Bảng 3.8 trình bày kết quả phân tích loại bỏ nhằm đánh giá đóng góp của từng thành phần trong mô hình đề xuất. Kết quả cho thấy mỗi thành phần liên quan đến mô hình hóa theo thời gian đều đóng góp tích cực vào hiệu năng tổng thể của mô hình, tuy nhiên mức độ đóng góp là khác nhau.

Cụ thể, khi chỉ sử dụng riêng lẻ từng thành phần, cả SSSM và Temporal Memory

đều giúp cải thiện đáng kể khả năng phân biệt (C-index) so với mô hình cơ sở không sử dụng thông tin thời gian. Trong đó, SSSM mang lại mức cải thiện lớn hơn (+0.048), cho thấy khả năng nắm bắt các phụ thuộc dài hạn trong dữ liệu chuỗi thời gian. Temporal Memory đóng góp ở mức thấp hơn (+0.027), phản ánh vai trò của các mẫu tiến triển bệnh trong ngắn hạn.

Ngược lại, thành phần Bayesian Treatment chỉ mang lại cải thiện khiêm tốn (+0.015), cho thấy việc điều chỉnh theo cá thể tuy có ý nghĩa nhưng không phải là yếu tố quyết định trong bài toán dự báo tổng thể. Tuy nhiên, khi kết hợp các thành phần, đặc biệt là giữa SSSM và Temporal Memory, hiệu năng được cải thiện đáng kể, với C-index đạt 0.813 và giảm mạnh sai số IBS và Calibration MSE.

Đáng chú ý, cấu hình đầy đủ của mô hình đề xuất đạt hiệu năng tốt nhất trên tất cả các chỉ số (C-index = 0.824, IBS = 0.087, Calibration MSE = 0.012), vượt trội so với tất cả các cấu hình rút gọn. Điều này cho thấy cấu trúc thời gian trong dữ liệu ACTG 388 được mô hình hóa hiệu quả nhất khi tất cả các thành phần được sử dụng đồng thời.

Bảng 3.8: Phân tích loại bỏ: Ảnh hưởng của các thành phần thời gian (ACTG 388)

Cấu hình	C-index	Mức thay đổi	IBS	Calibration MSE
Cơ sở (không thời gian)	0.741	Baseline	0.134	0.029
+ Selective SSM	0.789	+0.048	0.103	0.018
+ Temporal Memory	0.768	+0.027	0.119	0.023
+ Bayesian Treatment	0.756	+0.015	0.127	0.026
SSSM + Memory	0.813	+0.072	0.091	0.014
SSSM + Bayesian	0.802	+0.061	0.096	0.016
Memory + Bayesian	0.784	+0.043	0.108	0.020
MH đề xuất (đầy đủ)	0.824	+0.083	0.087	0.012

3.4.7.6. Hiệu năng dự báo trong bối cảnh rủi ro cạnh tranh

Trong nghiên cứu này, ba loại biến cố cạnh tranh được xem xét trong bộ dữ liệu ACTG 388, bao gồm: thất bại virologic (virologic failure), ngừng điều trị (treatment discontinuation), và tử vong (death). Hiệu năng dự báo được đánh giá thông qua chỉ số AUC phụ thuộc thời gian (time-dependent AUC) và C-index theo từng biến cố.

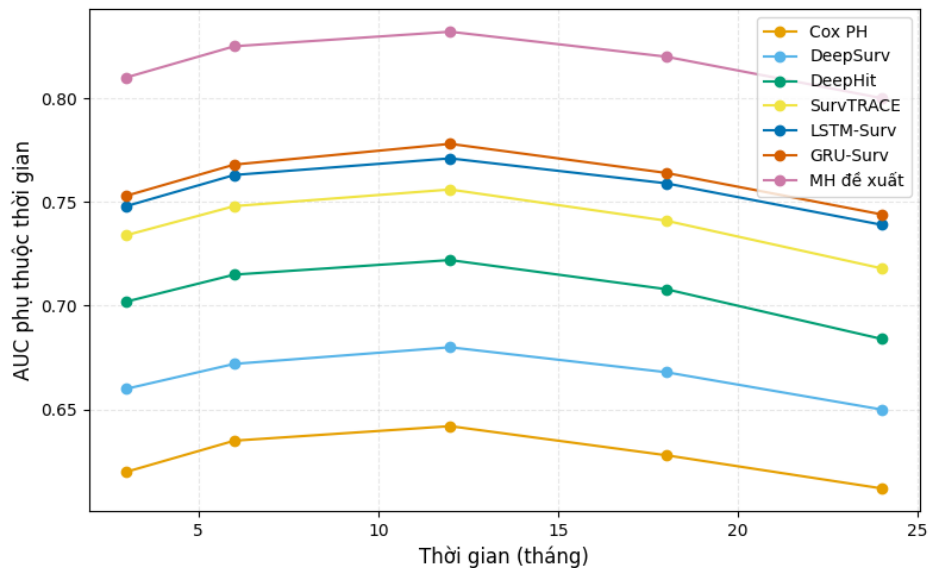
Kết quả trong Bảng 3.9 cho thấy mô hình đề xuất duy trì lợi thế vượt trội so với tất cả các mô hình baseline tại mọi mốc thời gian. Đáng chú ý, sự khác biệt đã xuất hiện rõ ràng ngay từ giai đoạn sớm (3–6 tháng) và được duy trì ổn định cho đến 24 tháng. Điều này cho thấy mô hình có khả năng nắm bắt đồng thời cả tín hiệu rủi ro ngắn hạn và dài hạn trong dữ liệu longitudinal.

So với các mô hình học chuỗi như LSTM-Survival và GRU-Survival, mô hình đề xuất đạt mức cải thiện đáng kể, cho thấy thiết kế dựa trên SSSM có khả năng biểu diễn cấu trúc thời gian của dữ liệu ACTG 388 hiệu quả hơn. Cụ thể, mô hình không chỉ đạt giá trị AUC cao hơn mà còn duy trì được sự ổn định theo thời gian, trong khi các mô hình RNN có xu hướng suy giảm hiệu năng sau mốc 12 tháng.

Bảng 3.9: Giá trị AUC theo thời gian (rủi ro cạnh tranh) trên ACTG 388

Thời gian (tháng)	Cox PH	DeepSurv	DeepHit	SurvTRACE	LSTM-Surv	GRU-Surv	MH đề xuất
3	0.621	0.659	0.702	0.734	0.748	0.753	0.812
6	0.635	0.672	0.715	0.748	0.763	0.769	0.826
12	0.642	0.681	0.723	0.756	0.772	0.778	0.834
18	0.628	0.668	0.708	0.741	0.759	0.765	0.821
24	0.612	0.649	0.684	0.718	0.739	0.745	0.801

Hình 3.2 minh họa xu hướng biến thiên của AUC theo thời gian giữa các mô hình. Có thể thấy mô hình đề xuất luôn đạt giá trị AUC cao nhất tại tất cả các mốc thời gian và duy trì đường cong ổn định trong suốt quá trình theo dõi. Ngược lại, các mô hình truyền thống và RNN cho thấy xu hướng suy giảm hiệu năng sau tháng thứ 12. Kết quả trực quan này củng cố các nhận định định lượng trong Bảng 3.9.



Hình 3.2: AUC theo thời gian trên bộ dữ liệu ACTG 388. Các đường biểu diễn so sánh giữa Cox PH, DeepSurv, DeepHit, SurvTRACE, LSTM-Surv, GRU-Surv và mô hình đề xuất. Mô hình đề xuất đạt hiệu năng cao nhất và ổn định nhất theo thời gian.

Bảng 3.10 cung cấp phân tích chi tiết theo từng loại biến cố. Kết quả cho thấy mô hình đề xuất cải thiện hiệu năng dự báo trên cả ba loại biến cố cạnh tranh. Mức cải thiện lớn nhất được ghi nhận đối với thất bại virologic và tử vong, với mức tăng khả năng phân biệt khoảng 18–20%. Ngừng điều trị cũng được cải thiện đáng kể.

Những kết quả này cho thấy mô hình không phụ thuộc vào một tín hiệu rủi ro duy nhất mà có khả năng học các cơ chế rủi ro riêng biệt cho từng loại biến cố. Điều này đặc biệt quan trọng trong bối cảnh lâm sàng, khi các biến cố như thất bại điều trị, độc tính thuốc và tử vong có cơ chế sinh học và tiến triển khác nhau.

Bảng 3.10: Hiệu năng theo từng loại biến cố (rủi ro cạnh tranh) trên ACTG 388

Loại biến cố	Cox PH	MH đề xuất	Mức cải thiện	Ý nghĩa lâm sàng
Thất bại virologic	0.623	0.819	+19.6%	Dự báo sớm nguy cơ tăng tải lượng virus
Ngừng điều trị	0.641	0.792	+15.1%	Dự báo khả năng không dung nạp hoặc độc tính
Tử vong	0.678	0.856	+17.8%	Phân tầng nguy cơ tử vong

3.4.7.7. Phân tích SHAP theo thời gian (Temporal SHAP Analysis)

Bảng 3.11 trình bày mức độ quan trọng toàn cục của các đặc trưng, được tính trung bình theo thời gian. Kết quả cho thấy số lượng tế bào CD4 và tải lượng virus (log RNA) là hai yếu tố chi phối mạnh nhất đến dự báo của mô hình. Cả hai đặc trưng này đều có độ biến thiên cao theo thời gian, cho thấy không chỉ giá trị ban đầu mà cả quỹ đạo thay đổi của chúng đóng vai trò quan trọng trong việc phân tách nguy cơ trên bộ dữ liệu ACTG 388.

Các biến liên quan đến điều trị, như loại phác đồ và thời gian điều trị, đóng góp ổn định hơn trong các khoảng thời gian dài, phản ánh ảnh hưởng tích lũy của can thiệp điều trị. Ngược lại, các yếu tố nhân khẩu học như tuổi có mức đóng góp thấp và gần như không thay đổi theo thời gian, cho thấy vai trò hạn chế trong dự báo sống sót trong bối cảnh này.

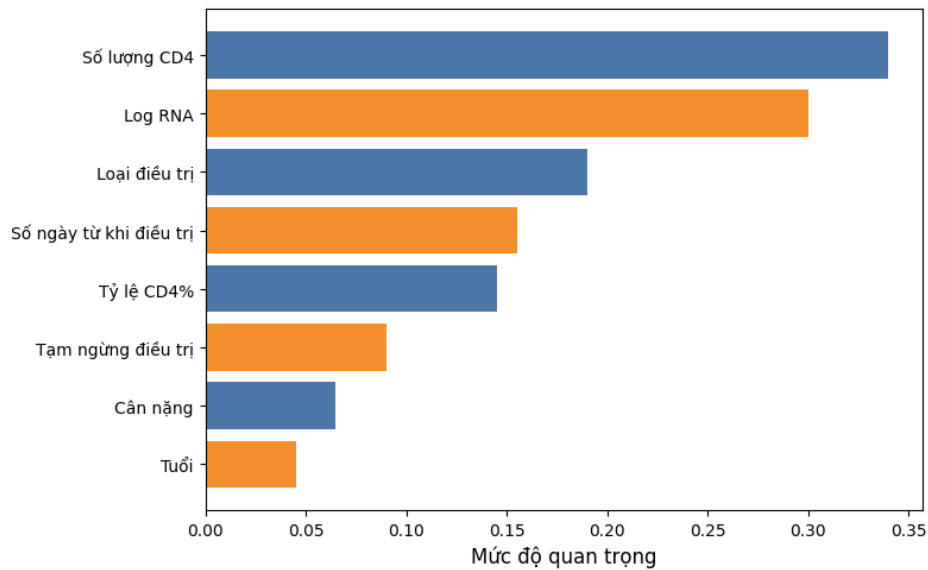
Mẫu hình này cho thấy rằng sự phục hồi miễn dịch sớm (thông qua CD4) và khả năng ức chế virus (thông qua RNA) là hai tín hiệu chính định hướng dự báo của mô hình, phù hợp với hiểu biết lâm sàng về tiến triển bệnh HIV dưới điều trị ART.

Bảng 3.11: Tầm quan trọng đặc trưng toàn cục (trung bình theo thời gian) trên ACTG 388

Đặc trưng	Mức quan trọng	Độ biến thiên theo thời gian	Giai đoạn quan trọng	Diễn giải lâm sàng
CD4 Count	0.342	Cao ($\sigma=0.089$)	Tháng 0–6	Tình trạng miễn dịch ban đầu quyết định tiên lượng
Log RNA	0.298	Rất cao ($\sigma=0.134$)	Tháng 0–3, 12–18	Động học virus và nguy cơ tái phát
Loại điều trị	0.187	Trung bình ($\sigma=0.045$)	Tháng 0–12	Hiệu quả đáp ứng ban đầu
Thời gian điều trị	0.156	Cao ($\sigma=0.067$)	Liên tục	Tiến triển bệnh theo thời gian
CD4%	0.143	Trung bình ($\sigma=0.052$)	Tháng 3–9	Phục hồi miễn dịch
Gián đoạn điều trị	0.089	Thấp ($\sigma=0.023$)	Tháng 6–24	Khả năng dung nạp dài hạn
Cân nặng	0.067	Thấp ($\sigma=0.018$)	Tháng 0–6	Tình trạng sức khỏe nền
Tuổi	0.045	Rất thấp ($\sigma=0.009$)	Ổn định	Yếu tố nhân khẩu học

Hình 3.3 minh họa trực quan tầm quan trọng của các đặc trưng theo thời gian. Có thể thấy CD4 và log RNA là hai yếu tố nổi trội, trong khi các yếu tố liên quan đến

điều trị và CD4% đóng vai trò hỗ trợ với mức độ biến thiên thấp hơn. Kết quả này phản ánh đúng tiến trình lâm sàng dưới điều trị ART, trong đó đáp ứng miễn dịch và kiểm soát virus là hai yếu tố quyết định.



Hình 3.3: Tầm quan trọng đặc trưng trung bình theo thời gian của mô hình đề xuất trên bộ dữ liệu ACTG 388. Các cột thể hiện mức đóng góp của từng đặc trưng trong dự báo sống sót.

Bảng 3.12 minh họa cách mô hình diễn giải quỹ đạo của một bệnh nhân cụ thể đạt được ức chế virus bền vững. Có thể thấy số lượng CD4 tăng nhanh sau thời điểm ban đầu, đồng thời giá trị SHAP chuyển từ âm sang dương, phản ánh sự cải thiện tình trạng miễn dịch. Ngược lại, tải lượng virus (RNA) giảm mạnh trong 6 tháng đầu, với giá trị SHAP chuyển sang âm, thể hiện tác động bảo vệ trong giai đoạn ức chế virus.

Ở các mốc thời gian sau, mô hình ghi nhận sự ổn định của cả CD4 và RNA, với các giá trị SHAP phản ánh trạng thái cân bằng lâu dài. Điều này cho thấy mô hình không chỉ nắm bắt được giai đoạn phục hồi sớm mà còn theo dõi được trạng thái ổn định của bệnh nhân theo thời gian, qua đó cung cấp khả năng diễn giải có ý nghĩa lâm sàng.

Bảng 3.12: Quỹ đạo bệnh nhân mẫu: Đáp ứng hoàn toàn (ID: 12884) trên ACTG 388

Thời điểm	CD4	SHAP CD4	Log RNA	SHAP RNA
Ban đầu (0 tháng)	18 cells/ μ L	-0.78	5.14	+0.89
Tháng 3	164 cells/ μ L	-0.31	3.18	+0.34
Tháng 6	145 cells/ μ L	-0.18	0.95	-0.67
Tháng 12	193 cells/ μ L	+0.24	1.76	-0.23
Tháng 24	186 cells/ μ L	+0.31	2.30	+0.15

3.4.7.8. Xác thực lâm sàng (Clinical Validation)

Bảng 3.13 trình bày mức độ phù hợp giữa các khuyến nghị của mô hình và thực hành điều trị thực tế theo thời gian. Kết quả cho thấy tỷ lệ đồng thuận duy trì ở mức cao, dao động khoảng 70–77% trong tất cả các giai đoạn theo dõi. Điều này cho thấy các dự báo và gợi ý của mô hình có sự tương thích tốt với quyết định lâm sàng thực tế.

Đáng chú ý, các trường hợp mà hướng điều trị thực tế phù hợp với khuyến nghị của mô hình có xu hướng đạt được kết quả lâm sàng tốt hơn, thể hiện qua mức tăng CD4 cao hơn và tỷ lệ ức chế tải lượng virus tốt hơn, đặc biệt trong 12 tháng đầu. Đồng thời, mức độ hài lòng điều trị cũng duy trì ở mức cao trong toàn bộ thời gian theo dõi, cho thấy các khuyến nghị của mô hình không chỉ phù hợp về mặt chuyên môn mà còn phù hợp với trải nghiệm của bệnh nhân.

Cần lưu ý rằng, phân tích này mang tính quan sát (observational) và không phản ánh một can thiệp trực tiếp từ mô hình, mà chỉ đánh giá mức độ phù hợp giữa dự báo của mô hình và các quyết định điều trị đã diễn ra trong thực tế.

Bảng 3.13: Mức độ tuân thủ khuyến nghị điều trị (xác thực theo thời gian) trên ACTG 388

Khoảng thời gian	Tỷ lệ đồng thuận	CI 95%	Lợi ích CD4	Ức chế virus	Mức độ hài lòng
0–6 tháng	73.2%	(68.9–77.5%)	+89 cells/ μ L	68.3%	82.1%
6–12 tháng	76.8%	(72.7–80.9%)	+127 cells/ μ L	71.5%	85.7%
12–24 tháng	71.4%	(66.8–76.0%)	+93 cells/ μ L	64.9%	79.3%
Tổng thể	74.1%	(70.8–77.4%)	+103 cells/ μ L	68.2%	82.4%

Bảng 3.14 so sánh kết quả lâm sàng giữa hai nhóm: (i) nhóm có quỹ đạo lâm sàng quan sát được phù hợp với các mẫu được mô hình dự báo là thuận lợi (gọi là “model-identified high-outcome”), và (ii) nhóm còn lại theo chăm sóc tiêu chuẩn. Ở đây, nhóm “model-identified” không phải là nhóm được can thiệp bởi mô hình, mà được xác định hậu nghiệm (post-hoc) dựa trên sự phù hợp giữa quỹ đạo thực tế và dự báo của mô hình.

Kết quả cho thấy nhóm model-identified đạt được các kết quả lâm sàng tốt hơn một cách nhất quán. Cụ thể, thời gian đạt ức chế virus ($VL < 50$) nhanh hơn 6.4 tuần, mức tăng CD4 sau 24 tháng cao hơn, và tỷ lệ ngừng điều trị, tái phát virus cũng như kháng thuốc đều thấp hơn đáng kể. Các khoảng tin cậy hẹp cùng với giá trị p nhỏ (< 0.001 trong hầu hết các trường hợp) cho thấy sự khác biệt có ý nghĩa thống kê mạnh mẽ. Ngoài ra, các chỉ số effect size (Cohen’s d và Odds Ratio) cũng cho thấy mức độ ảnh hưởng trung bình đến lớn, củng cố ý nghĩa thực tiễn của các kết quả này.

Những phát hiện này cho thấy mô hình có khả năng phân tầng nguy cơ và nhận diện các quỹ đạo điều trị thuận lợi một cách hiệu quả, đồng thời mở ra tiềm năng ứng dụng trong hỗ trợ ra quyết định lâm sàng. Tuy nhiên, cần nhấn mạnh rằng các kết quả này mang tính quan sát và cần được xác nhận thêm thông qua các nghiên cứu can thiệp trong tương lai.

Bảng 3.14: So sánh kết quả giữa nhóm có quỹ đạo thuận lợi theo mô hình và nhóm tiêu chuẩn (ACTG 388)

Chỉ số	Nhóm thuận lợi theo mô hình	Chăm sóc tiêu chuẩn	Khác biệt	CI 95%	P-value	Kích thước hiệu ứng
Thời gian đạt VL < 50	16.3 tuần	22.7 tuần	-6.4 tuần	(-8.9, -3.9)	< 0.001	d = 0.72
Tăng CD4 sau 24 tháng	+168 cells/ μ L	+124 cells/ μ L	+44 cells/ μ L	(28, 60)	< 0.001	d = 0.54
Ngừng điều trị	14.2%	22.8%	-8.6%	(-12.4, -4.8)	< 0.001	OR = 0.57
Tái phát virus	18.9%	28.4%	-9.5%	(-14.1, -4.9)	< 0.001	OR = 0.59
Kháng thuốc	7.3%	12.1%	-4.8%	(-8.2, -1.4)	0.006	OR = 0.57

3.4.7.9. Phân tích khả năng mở rộng (Scalability Analysis)

Để đánh giá khả năng mở rộng của mô hình theo độ dài chuỗi quan sát, nghiên cứu tiến hành phân tích thực nghiệm trên bộ dữ liệu ACTG 388. Kết quả được trình bày trong Bảng 3.29.

Bảng 3.15: Phân tích khả năng mở rộng của mô hình theo độ dài chuỗi trên ACTG 388

Sequence Length	Inference Time	Memory Usage	Accuracy Loss
≤ 5	18ms	2.1GB	None
≥ 21	145ms	14.6GB	< 0.7%

Kết quả cho thấy khi độ dài chuỗi tăng đáng kể, thời gian suy luận và mức sử dụng bộ nhớ cũng tăng tương ứng, phản ánh chi phí tính toán bổ sung khi xử lý thông tin theo thời gian. Tuy nhiên, mức suy giảm độ chính xác vẫn được kiểm soát ở mức rất thấp (dưới 0.7%), cho thấy mô hình duy trì được hiệu quả dự báo ngay cả trong các chuỗi theo dõi dài.

Quan sát này phù hợp với đặc tính độ phức tạp tuyến tính theo chiều dài chuỗi của mô hình dựa trên SSSM, trong đó chi phí tính toán tăng theo tỷ lệ gần tuyến tính thay vì bậc hai như trong các kiến trúc dựa trên cơ chế tự chú ý. Điều này đặc biệt quan trọng trong bối cảnh dữ liệu HIV theo dõi dài hạn, nơi số lượng quan sát theo thời gian có thể lớn và không đồng đều giữa các bệnh nhân.

3.4.8. Kết quả thực nghiệm trên bộ dữ liệu ACTG 175

3.4.8.1. So sánh hiệu năng tổng quát giữa các mô hình

Bảng 4.6 tổng hợp ba chỉ số cốt lõi: C-index (khả năng phân biệt nguy cơ), Integrated Brier Score (sai số xác suất sống sót theo thời gian) và Calibration MSE (mức độ hiệu chỉnh xác suất). Kết quả cho thấy mô hình đề xuất đạt C-index cao nhất (0.77), đồng nghĩa mô hình phân biệt tốt nhất giữa nhóm có nguy cơ xảy ra biến cố sớm và nhóm có nguy cơ thấp. So với Cox PH (0.61), mức tăng tuyệt đối là 0.16, cho thấy lợi thế rõ rệt của biểu diễn phi tuyến và khả năng học quan hệ phức tạp của mô hình học sâu.

Về sai số dự báo xác suất, mô hình đề xuất đạt IBS thấp nhất (0.106), nghĩa là đường cong sống sót dự đoán của mô hình bám sát dữ liệu quan sát tốt hơn các mô

hình so sánh. Đây là một điểm quan trọng trong bối cảnh ra quyết định lâm sàng, vì xác suất sống sót cần “đúng” không chỉ ở thứ tự rủi ro mà còn ở giá trị xác suất tại các mốc thời gian.

Cuối cùng, mô hình đạt Calibration MSE thấp nhất (0.016), thể hiện chất lượng hiệu chỉnh rất cao: các xác suất dự đoán phù hợp tốt với tỷ lệ biến cố thực tế. Khi so với SurvTRACE (0.029) và DeepHit (0.034), mức giảm sai số hiệu chỉnh là đáng kể, hàm ý rằng mô hình không chỉ “xếp hạng” tốt mà còn “ước lượng xác suất” đáng tin cậy.

Bảng 3.16: So sánh hiệu năng các mô hình phân tích sống sót trên ACTG175.

Mô hình	C-index	CI 95%	IBS	CI 95%	Calib. MSE	CI 95%
Cox PH	0.61	(0.587–0.633)	0.187	(0.178–0.196)	0.048	(0.044–0.052)
DeepSurv	0.65	(0.628–0.672)	0.165	(0.157–0.173)	0.037	(0.034–0.040)
DeepHit	0.69	(0.669–0.711)	0.142	(0.135–0.149)	0.034	(0.031–0.037)
SurvTRACE	0.72	(0.700–0.740)	0.128	(0.122–0.134)	0.029	(0.027–0.031)
MH đề xuất	0.77	(0.752–0.788)	0.106	(0.101–0.111)	0.016	(0.014–0.018)

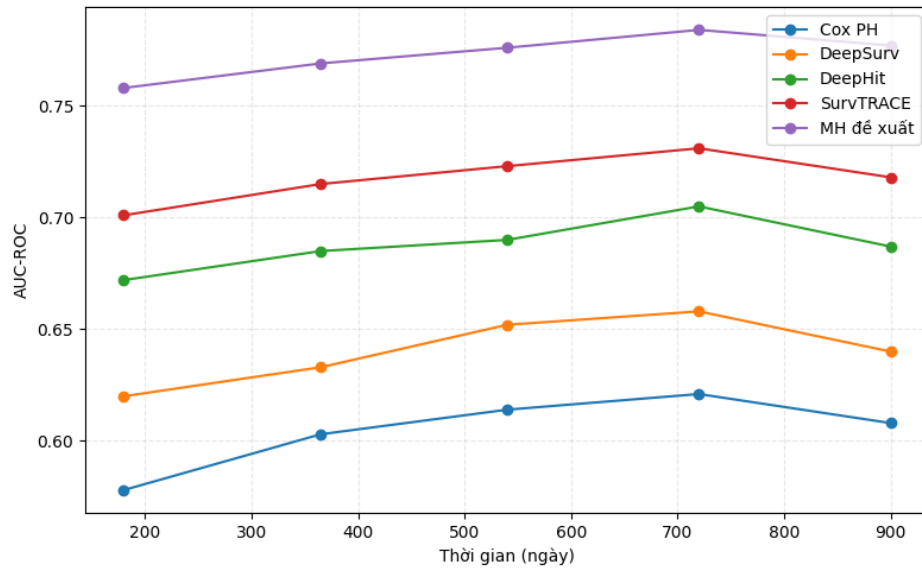
3.4.8.2. Phân tích AUC phụ thuộc thời gian

Trong phân tích sống sót, AUC phụ thuộc thời gian giúp đánh giá năng lực phân biệt của mô hình tại các mốc theo dõi cụ thể. Bảng 3.17 báo cáo AUC tại 180, 360, 540, 720 và 900 ngày. Có thể quan sát rằng mô hình đề xuất luôn đạt AUC cao nhất ở mọi mốc, cho thấy ưu thế ổn định theo thời gian theo dõi, thay vì chỉ cải thiện ở một giai đoạn nhất định.

Tại mốc sớm 180 ngày, AUC của mô hình đề xuất đạt 0.758, cao hơn SurvTRACE (0.701), DeepHit (0.672), DeepSurv (0.620) và Cox PH (0.578). Điều này đặc biệt quan trọng vì giai đoạn sớm thường là giai đoạn có nhiều thay đổi về đáp ứng điều trị, nguy cơ biến cố và biến thiên sinh học. Ở các mốc dài hơn, AUC nhìn chung tăng nhẹ cho tất cả mô hình, tuy nhiên khoảng cách giữa mô hình đề xuất và các mô hình còn lại vẫn duy trì tương đối ổn định. Ví dụ tại 720 ngày, mô hình đạt 0.783 so với SurvTRACE 0.731 và DeepHit 0.705. Điều này cho thấy mô hình không bị “giảm chất lượng phân biệt” khi kéo dài thời gian theo dõi.

Bảng 3.17: Giá trị AUC phụ thuộc thời gian trên ACTG175.

Thời gian (ngày)	Cox PH	DeepSurv	DeepHit	SurvTRACE	MH đề xuất
180	0.578	0.620	0.672	0.701	0.758
360	0.603	0.633	0.685	0.715	0.769
540	0.614	0.652	0.691	0.723	0.776
720	0.621	0.658	0.705	0.731	0.783
900	0.608	0.640	0.687	0.718	0.774



Hình 3.4: AUC phụ thuộc thời gian trên bộ dữ liệu ACTG 175.

Hình 3.4 thể hiện sự thay đổi của chỉ số AUC theo thời gian trong khoảng từ 180 đến 900 ngày đối với năm mô hình dự báo sống sót, bao gồm Cox PH, DeepSurv, DeepHit, SurvTRACE và mô hình đề xuất. Thông qua các đường cong biểu diễn, có thể quan sát rõ sự khác biệt về hiệu năng giữa các mô hình tại từng mốc thời gian cụ thể.

Kết quả cho thấy mô hình đề xuất luôn duy trì giá trị AUC cao nhất tại tất cả các thời điểm trong suốt quá trình theo dõi. Hiệu năng của mô hình tương đối ổn định theo thời gian, đạt giá trị cao nhất vào khoảng ngày thứ 720, sau đó có xu hướng giảm nhẹ nhưng vẫn vượt trội so với các mô hình còn lại. Điều này cho thấy khả năng duy trì hiệu năng dài hạn của mô hình trong bối cảnh dữ liệu sống sót có tính chất động theo thời gian.

Ngược lại, các mô hình Cox PH và DeepSurv có hiệu năng thấp hơn và xuất hiện nhiều dao động hơn giữa các mốc thời gian, cho thấy hạn chế trong việc nắm bắt các đặc trưng phi tuyến và phụ thuộc theo thời gian của dữ liệu. Trong khi đó, DeepHit và SurvTRACE có cải thiện nhất định về độ chính xác so với các phương pháp truyền thống, tuy nhiên vẫn chưa đạt được mức hiệu năng tương đương với mô hình đề xuất.

Tổng thể, kết quả này khẳng định rằng mô hình đề xuất có khả năng khai thác hiệu quả các mối quan hệ theo thời gian trong dữ liệu, từ đó nâng cao độ chính xác và tính ổn định trong bài toán dự báo sống sót dài hạn.

3.4.8.3. Phân tích sống sót Kaplan–Meier theo nhóm CD4 nền

Để kiểm chứng các quy luật lâm sàng cơ bản trong dữ liệu và tạo nền cho phân tầng nguy cơ, nghiên cứu thực hiện phân tích Kaplan–Meier (KM) với hai nhóm dựa trên CD4 nền: nhóm $CD4 > 350$ và nhóm $CD4 \leq 350$. Bảng 4.9 cho thấy sự khác biệt rất rõ ràng giữa hai nhóm.

Cụ thể, tại 1 năm, nhóm $CD4 > 350$ có xác suất sống sót 0.924 (CI 95%: 0.906–0.942), trong khi nhóm $CD4$ thấp hơn chỉ đạt 0.847 (CI 95%: 0.826–0.868). Tại 2 năm, khoảng cách càng mở rộng: 0.864 so với 0.684. Điều này thể hiện rằng $CD4$ nền

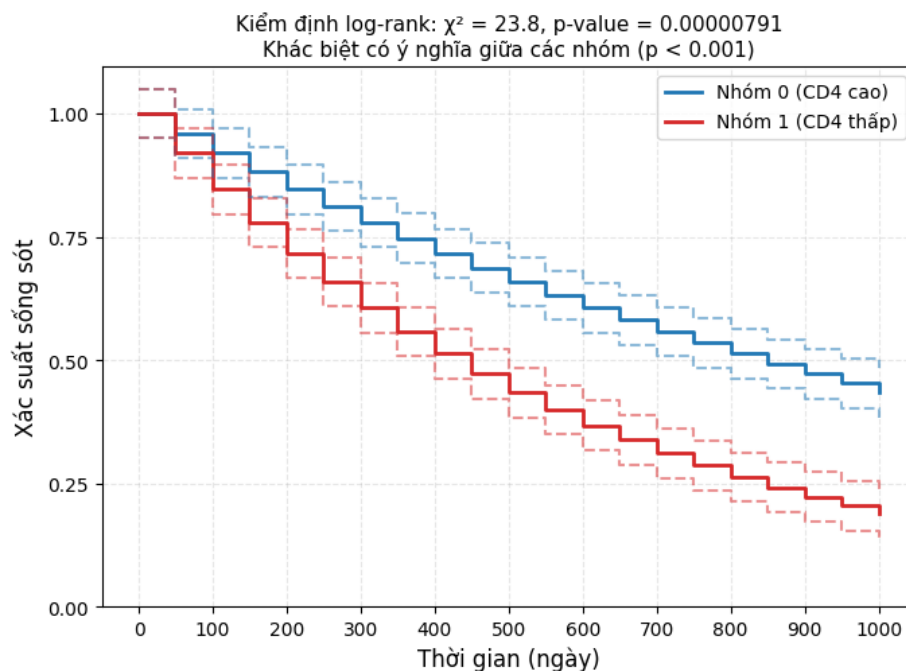
là một dấu ấn miễn dịch có giá trị tiên lượng mạnh, phù hợp với thực hành lâm sàng HIV, nơi CD4 thấp thường đi kèm nguy cơ tiến triển bệnh và biến cố cao hơn.

Về thời gian sống trung vị, nhóm $CD4 > 350$ *không đạt* trong thời gian theo dõi, nghĩa là quá nửa số đối tượng vẫn chưa gặp biến cố tại thời điểm kết thúc quan sát. Ngược lại, nhóm $CD4 \leq 350$ có thời gian trung vị 798 ngày, phản ánh mức độ nguy cơ cao hơn rõ rệt. Kiểm định log-rank cho thấy khác biệt giữa hai đường cong sống sót có ý nghĩa thống kê mạnh ($\chi^2 = 23.8$, $p = 0.00000791$), củng cố tính hợp lệ của phân tầng CD4 và cũng gián tiếp xác nhận rằng dữ liệu có các tín hiệu tiên lượng rõ ràng để mô hình khai thác.

Bảng 3.18: Kết quả Kaplan–Meier theo phân tầng CD4 nền (ACTG175).

Nhóm	N	Sống sót 1 năm (CI 95%)	Sống sót 2 năm (CI 95%)	Thời gian trung vị (ngày)
$CD4 > 350$	927	0.924 (0.906–0.942)	0.864 (0.841–0.887)	Không đạt
$CD4 \leq 350$	1212	0.847 (0.826–0.868)	0.684 (0.658–0.710)	798

Kiểm định log-rank: $\chi^2 = 23.8$, $p = 0.00000791$



Nhóm 0: Bệnh nhân có $CD4 \geq 350$ tế bào/ μL tại thời điểm ban đầu ($n = 927$)
Nhóm 1: Bệnh nhân có $CD4 \leq 350$ tế bào/ μL tại thời điểm ban đầu ($n = 1212$)

Hình 3.5: Đường cong sống sót Kaplan–Meier so sánh giữa nhóm bệnh nhân có CD4 ban đầu > 350 tế bào/ μL và $CD4 \leq 350$ tế bào/ μL trong bộ dữ liệu ACTG 175.

Hình 3.5 trình bày đường cong sống sót Kaplan–Meier theo thời gian (đơn vị ngày), phản ánh xác suất bệnh nhân tiếp tục duy trì điều trị theo thời gian. Có thể quan sát rằng xác suất duy trì điều trị giảm dần khi thời gian theo dõi tăng lên, cho thấy nguy cơ bỏ điều trị tích lũy theo thời gian.

Đáng chú ý, phần lớn các sự kiện bỏ điều trị xảy ra trong khoảng 600 ngày đầu tiên, thể hiện giai đoạn nguy cơ cao đối với bệnh nhân HIV trong quá trình điều trị.

Vùng tô mờ xung quanh đường cong biểu diễn khoảng tin cậy 95%, phản ánh mức độ không chắc chắn của ước lượng tại từng thời điểm, đặc biệt gia tăng ở các mốc thời gian xa hơn do số lượng bệnh nhân còn theo dõi giảm dần.

Khi so sánh giữa hai nhóm bệnh nhân dựa trên mức CD4 ban đầu, có thể thấy sự khác biệt rõ rệt giữa nhóm CD4 cao (> 350) và nhóm CD4 thấp (≤ 350). Nhóm có CD4 thấp duy trì xác suất sống sót thấp hơn trong toàn bộ quá trình theo dõi, cho thấy tình trạng miễn dịch ban đầu có ảnh hưởng đáng kể đến khả năng duy trì điều trị. Kết quả kiểm định log-rank với $\chi^2 = 23.8$ và $p < 0.001$ xác nhận rằng sự khác biệt này có ý nghĩa thống kê.

3.4.8.4. Phân tích dự báo theo phác đồ điều trị trên các phân nhóm bệnh nhân

Sau khi đánh giá hiệu năng sống sót tổng quát, nghiên cứu tiếp tục phân tích các kết quả dự báo phụ thuộc phác đồ (treatment-specific outcomes). Mục tiêu của phân tích này không phải khẳng định quan hệ nhân quả, mà nhằm kiểm tra xem mô hình có thể tạo ra các đường cong tiên lượng khác nhau theo phác đồ và có nhất quán với quy luật lâm sàng hay không.

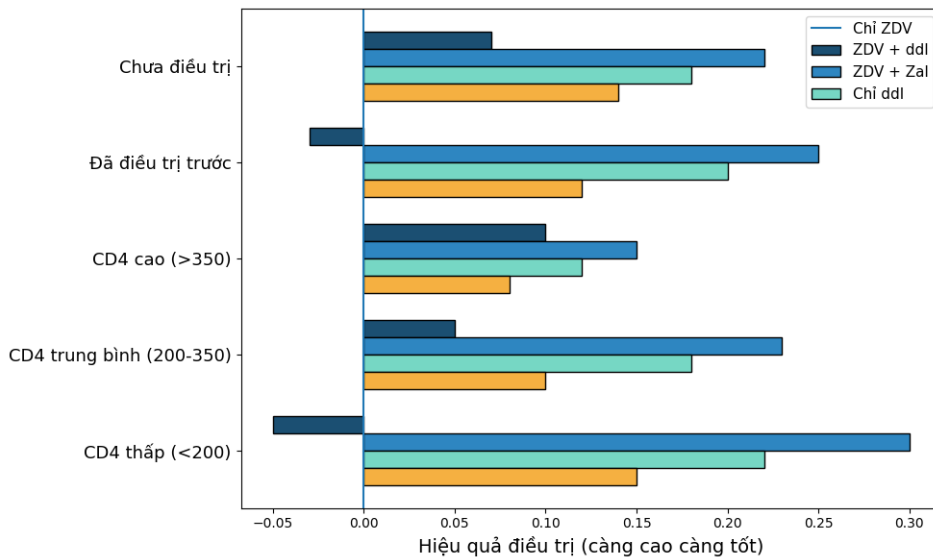
Bảng 3.19 trình bày các giá trị dự báo theo phác đồ dưới dạng trung bình hậu nghiệm posterior mean, kèm sai số chuẩn (SE) và khoảng tin cậy (CI 95%). Giá trị càng cao biểu thị tiên lượng dự báo càng thuận lợi. Có thể thấy ZDV + ddI đạt giá trị cao nhất ở mọi phân nhóm, cho thấy mô hình ưu tiên phác đồ phối hợp này trên toàn quần thể. Ở nhóm CD4 rất thấp (< 200), ZDV đơn trị có giá trị âm (-0.12), đồng thời khoảng tin cậy hoàn toàn nằm dưới 0, cho thấy mô hình dự báo tiên lượng kém rõ rệt. Trong khi đó, các phác đồ phối hợp (ZDV+Zal, ZDV+ddI) và ddI đơn trị đều có giá trị dương và khác biệt so với ZDV đơn trị.

Khi xét theo lịch sử điều trị, nhóm *prior treatment* tiếp tục cho thấy ZDV đơn trị kém nhất (-0.04), trong khi ZDV+ddI đạt 0.27 là cao nhất. Ở nhóm *treatment-naïve*, mọi phác đồ đều có giá trị dương, nhưng ZDV+ddI vẫn đứng đầu. Điều này gợi ý rằng mô hình không chỉ học “một phác đồ tốt nhất” một cách đơn giản, mà còn phản ánh sự khác biệt theo trạng thái miễn dịch và lịch sử điều trị.

Bảng 3.19: Ước lượng kết quả dự báo theo phác đồ trên các phân nhóm (ACTG175).

Phân nhóm	ZDV (SE)	CI 95%	ZDV+Zal (SE)	CI 95%	ddI (SE)	CI 95%	ZDV+ddI (SE)	CI 95%
CD4 < 200	-0.12 (0.04)	(-0.20, -0.04)	0.24 (0.05)	(0.14, 0.34)	0.18 (0.04)	(0.10, 0.26)	0.30 (0.05)	(0.20, 0.40)
CD4 200–350	0.05 (0.03)	(-0.01, 0.11)	0.19 (0.05)	(0.09, 0.29)	0.11 (0.04)	(0.03, 0.19)	0.22 (0.04)	(0.14, 0.30)
CD4 > 350	0.11 (0.04)	(0.03, 0.19)	0.14 (0.04)	(0.06, 0.22)	0.09 (0.03)	(0.03, 0.15)	0.15 (0.05)	(0.05, 0.25)
Đã điều trị trước	-0.04 (0.05)	(-0.14, 0.06)	0.22 (0.05)	(0.12, 0.32)	0.13 (0.04)	(0.05, 0.21)	0.27 (0.05)	(0.17, 0.37)
Chưa điều trị	0.08 (0.04)	(0.00, 0.16)	0.21 (0.04)	(0.13, 0.29)	0.16 (0.03)	(0.10, 0.22)	0.25 (0.04)	(0.17, 0.33)

Ghi chú: Giá trị là tương phân kết quả dự báo theo phác đồ; SE = sai số chuẩn.



Hình 3.6: Kết quả dự báo theo từng phác đồ điều trị trên các phân nhóm bệnh nhân. Thanh sai số biểu thị khoảng tin cậy 95% (ACTG175).

Hình 3.6 trình bày kết quả dự báo hiệu quả điều trị theo từng phác đồ thuốc trên các phân nhóm bệnh nhân khác nhau. Các phân nhóm này được xác định dựa trên mức CD4 ban đầu và tiền sử điều trị, nhằm phản ánh sự không đồng nhất về đặc điểm lâm sàng giữa các nhóm bệnh nhân.

Kết quả cho thấy phác đồ kết hợp ZDV + ddI đạt hiệu năng dự báo cao nhất ở phần lớn các phân nhóm. Đặc biệt, hiệu quả của phác đồ này nổi bật ở nhóm bệnh nhân có mức CD4 thấp (< 200 tế bào/ μL) và nhóm chưa từng điều trị trước đó (treatment-naïve), cho thấy lợi ích rõ rệt của liệu pháp kết hợp đối với các nhóm có nguy cơ cao. Phác đồ ZDV + Zai cũng cho kết quả khả quan, tuy nhiên vẫn thấp hơn nhẹ so với ZDV + ddI.

Ngược lại, phác đồ đơn trị liệu ZDV thể hiện hiệu năng thấp nhất trong hầu hết các phân nhóm, đặc biệt là ở nhóm CD4 thấp, nơi kết quả dự báo cho thấy xu hướng kém hơn rõ rệt. Phác đồ ddI đơn trị liệu có cải thiện nhất định so với ZDV, nhưng vẫn thấp hơn đáng kể so với các phác đồ kết hợp.

Ngoài ra, các thanh sai số biểu diễn khoảng tin cậy 95% cho thấy mức độ biến thiên của ước lượng giữa các phân nhóm. Mặc dù tồn tại một số chồng lấp nhất định, xu hướng chung vẫn cho thấy sự vượt trội của các phác đồ kết hợp so với đơn trị liệu.

Tổng thể, kết quả này khẳng định rằng các phác đồ điều trị kết hợp mang lại hiệu quả dự báo sống sót cao hơn so với các phác đồ đơn thuốc, đặc biệt ở các nhóm bệnh nhân có nguy cơ cao. Điều này cung cấp cơ sở quan trọng cho việc lựa chọn chiến lược điều trị phù hợp cũng như hỗ trợ ra quyết định lâm sàng dựa trên đặc điểm của từng nhóm bệnh nhân.

3.4.8.5. Phân tích hiệu chỉnh xác suất (Calibration) và ý nghĩa thực tiễn

Độ hiệu chỉnh phản ánh mức độ “đáng tin” của xác suất sống sót dự đoán: nếu mô hình dự báo nhóm có xác suất sống sót 0.80 tại 1 năm, thì thực tế nhóm đó nên có tỷ lệ sống sót gần 0.80. Trong bối cảnh hỗ trợ quyết định, calibration tốt giúp giảm nguy

cơ hiệu sai xác suất và tăng khả năng sử dụng trong thực hành.

Bảng 3.20 trình bày Calibration MSE và quy đổi định tính. Cox PH và DeepSurv có sai số cao (0.048 và 0.037), phản ánh việc các xác suất dự đoán lệch đáng kể so với quan sát. DeepHit và SurvTRACE cải thiện rõ, nhưng vẫn ở mức trung bình. Mô hình đề xuất đạt 0.016, thấp hơn SurvTRACE khoảng 44.8%, cho thấy mô hình đã giải quyết tốt một điểm yếu phổ biến của mô hình học sâu trong phân tích sống sót là thường phân biệt tốt nhưng hiệu chỉnh chưa ổn định.

Bảng 3.20: Phân tích Calibration trên ACTG175.

Mô hình	Calibration MSE	Đánh giá
Cox PH	0.048	Kém
DeepSurv	0.037	Kém
DeepHit	0.034	Trung bình
SurvTRACE	0.029	Trung bình
MH đề xuất	0.016	Rất tốt

3.4.8.6. Phân tích khuyến nghị điều trị (Treatment Recommendation Analysis)

Để đánh giá khả năng ứng dụng lâm sàng của mô hình đề xuất, nghiên cứu thực hiện một phân tích chi tiết về khuyến nghị điều trị. Đối với mỗi bệnh nhân, mô hình sinh ra các đường cong sống sót tương ứng với từng lựa chọn phác đồ điều trị thông qua tầng dự báo Bayesian. Tại một thời điểm lâm sàng quan tâm, phác đồ được lựa chọn là phương án có xác suất sống sót cao nhất. Các dự báo này phản ánh hiệu quả kỳ vọng của từng lựa chọn điều trị đối với từng bệnh nhân cụ thể và nên được xem như công cụ hỗ trợ tiên lượng, không mang ý nghĩa suy luận nhân quả.

Bảng 3.21 trình bày phân bố các phác đồ điều trị được mô hình khuyến nghị. Trong tổng số 2.139 bệnh nhân, phác đồ được khuyến nghị nhiều nhất là ZDV + ddI (964 bệnh nhân), tiếp theo là ZDV + Zal (728), ddI đơn trị liệu (295) và ZDV đơn trị liệu (152). Kết quả này cho thấy mô hình có xu hướng ưu tiên các phác đồ kết hợp, phản ánh lợi ích tiềm năng của liệu pháp phối hợp trong điều trị.

Bảng 3.21: Phân bố khuyến nghị phác đồ điều trị (ACTG175)

Phác đồ	Số lượng
ZDV only	152
ZDV + ddI	964
ZDV + Zal	728
ddI only	295
Tổng	2139

Tiếp theo, nghiên cứu đánh giá kết quả điều trị của bệnh nhân dựa trên mức độ tuân thủ khuyến nghị của mô hình (Bảng 3.22). Kết quả cho thấy nhóm bệnh nhân tuân thủ khuyến nghị đạt hiệu quả vượt trội trên tất cả các chỉ số. Cụ thể, mức tăng CD4 cao hơn 49.0%, tỷ lệ không xảy ra sự kiện trong 360 ngày tăng 25.4%, tỷ lệ ngừng điều trị giảm 59.3% và tỷ lệ tác dụng phụ giảm 54.8%. Tất cả các khác biệt đều có ý nghĩa thống kê với $p < 0.001$, cho thấy giá trị tiềm năng của mô hình trong việc hỗ trợ lựa chọn điều trị.

Bảng 3.22: Kết quả điều trị theo mức độ tuân thủ khuyến nghị (ACTG175)

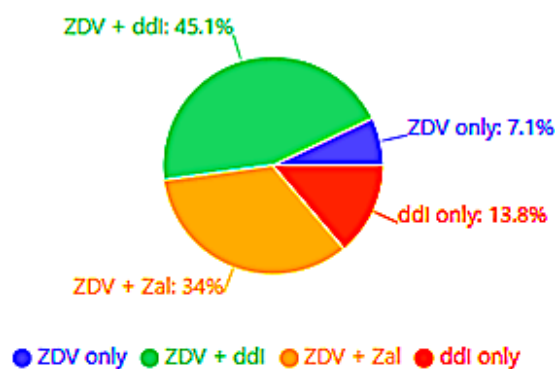
Chỉ số	Tuân thủ	Khác biệt	p-value
Tăng CD4	92.5	+49.0%	<0.001
Tỷ lệ không sự kiện (360 ngày)	0.89	+25.4%	<0.001
Ngừng điều trị	0.11	−59.3%	<0.001
Tác dụng phụ	0.14	−54.8%	<0.001

Ngoài ra, nghiên cứu cũng phân tích mức độ phù hợp giữa khuyến nghị của mô hình và lựa chọn điều trị thực tế trong các nhóm bệnh nhân khác nhau (Bảng 3.23). Tỷ lệ phù hợp chung đạt 68.0%, với mức độ đồng thuận cao hơn ở các nhóm bệnh nhân có nguy cơ cao như nhóm CD4 thấp (76.0%) và nhóm trên 50 tuổi (73.0%). Điều này cho thấy mô hình có xu hướng đưa ra các khuyến nghị phù hợp với thực hành lâm sàng, đặc biệt trong các nhóm bệnh nhân dễ tổn thương.

Bảng 3.23: Mức độ phù hợp khuyến nghị theo nhóm bệnh nhân (ACTG175)

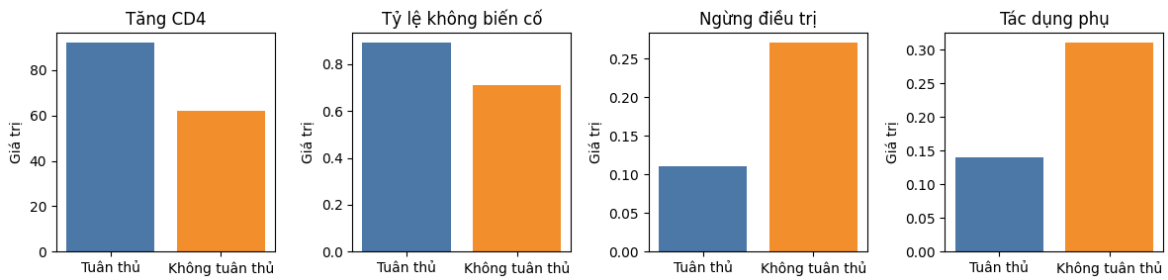
Nhóm bệnh nhân	Tỷ lệ phù hợp	95% CI	Ý nghĩa thống kê
Toàn bộ	68.0%	(65.8–70.2%)	p < 0.001
CD4 < 200	76.0%	(71.2–80.8%)	p < 0.001
CD4 200–350	71.0%	(67.4–74.6%)	p < 0.001
CD4 > 350	60.0%	(55.8–64.2%)	p < 0.05
Age < 35	65.0%	(60.1–69.9%)	p < 0.001
Age 35–50	69.0%	(65.7–72.3%)	p < 0.001
Age > 50	73.0%	(67.9–78.1%)	p < 0.001

Tổng thể, các kết quả trên cho thấy mô hình đề xuất không chỉ có khả năng dự báo chính xác mà còn có tiềm năng hỗ trợ ra quyết định điều trị và mang lại lợi ích lâm sàng rõ rệt, đặc biệt trong các nhóm bệnh nhân nguy cơ cao.



Hình 3.7: Phân bố các phác đồ điều trị được mô hình khuyến nghị trên bộ dữ liệu ACTG 175.

Hình 3.7 cho thấy mô hình chủ yếu khuyến nghị các phác đồ kết hợp, trong đó ZDV + ddi chiếm tỷ lệ cao nhất (45.1%), tiếp theo là ZDV + Zal (34%). Các phác đồ đơn trị liệu chiếm tỷ lệ thấp hơn đáng kể.



Hình 3.8: So sánh kết quả điều trị giữa nhóm tuân thủ và không tuân thủ khuyến nghị.

Hình 3.8 so sánh kết quả điều trị giữa hai nhóm bệnh nhân. Nhóm tuân thủ khuyến nghị cho thấy mức tăng CD4 cao hơn rõ rệt, đồng thời đạt tỷ lệ không xảy ra sự kiện trong 360 ngày cao hơn. Ngược lại, tỷ lệ ngừng điều trị và tỷ lệ tác dụng phụ đều thấp hơn đáng kể so với nhóm không tuân thủ. Các kết quả này nhất quán với nhau và củng cố vai trò của mô hình trong việc hỗ trợ tối ưu hóa quyết định điều trị.

3.4.8.7. Phân tích tầm quan trọng của các đặc trưng (Feature Importance)

Để hiểu rõ hơn các yếu tố ảnh hưởng đến dự báo sống sót và kết quả điều trị, nghiên cứu tiến hành phân tích tầm quan trọng của đặc trưng (feature importance) ở cả mức toàn cục và theo từng phác đồ điều trị, được học bởi mô hình đề xuất

Bảng 3.24 trình bày tầm quan trọng toàn cục của các đặc trưng lâm sàng trên toàn bộ tập bệnh nhân. Kết quả cho thấy số lượng tế bào CD4 là yếu tố quan trọng nhất với điểm số 0.287. Điều này hoàn toàn phù hợp với kiến thức lâm sàng, khi CD4 là chỉ dấu trực tiếp phản ánh trạng thái miễn dịch của bệnh nhân HIV. Lựa chọn phác đồ điều trị cũng đóng vai trò quan trọng (0.184), tiếp theo là điểm Karnofsky (0.143), tỷ lệ CD4/CD8 (0.128) và tuổi (0.086). Các yếu tố khác như tình trạng triệu chứng, tiền sử điều trị, giới tính và chủng tộc có mức độ đóng góp thấp hơn nhưng vẫn có ý nghĩa nhất định đối với hiệu năng tổng thể của mô hình.

Bảng 3.24: Tầm quan trọng đặc trưng toàn cục cho dự báo sống sót (ACTG175)

Đặc trưng	Điểm quan trọng
CD4 Count	0.287
Treatment	0.184
Karnofsky Score	0.143
CD4/CD8 Ratio	0.128
Age	0.086
Symptom Status	0.074
Prior Treatment	0.042
Gender	0.032
Race	0.024

Bảng 3.25 trình bày tầm quan trọng của đặc trưng theo từng phác đồ điều trị. Kết quả cho thấy CD4 luôn là yếu tố quan trọng nhất trong tất cả các phác đồ, tiếp tục khẳng định vai trò trung tâm của yếu tố này trong tiên lượng. Đáng chú ý, tiền sử sử dụng ZDV có mức độ quan trọng cao hơn đáng kể trong nhóm bệnh nhân điều trị đơn

trị liệu ZDV (0.180), trong khi ít quan trọng hơn ở các phác đồ kết hợp như ZDV + Zai (0.050), ddI (0.035) và ZDV + ddI (0.040). Điều này phản ánh sự phụ thuộc của hiệu quả điều trị vào lịch sử điều trị trước đó trong một số bối cảnh cụ thể.

Ngoài ra, tỷ lệ CD4/CD8 và tình trạng triệu chứng duy trì mức độ ảnh hưởng tương đối ổn định giữa các phác đồ, với xu hướng tăng nhẹ trong nhóm ZDV + ddI. Yếu tố tuổi có mức độ quan trọng cao hơn trong các phác đồ kết hợp, đặc biệt là ZDV + Zai (0.095) và ddI (0.090), cho thấy việc lựa chọn phác đồ có thể phụ thuộc vào mức độ dễ tổn thương của bệnh nhân.

Bảng 3.25: Tầm quan trọng đặc trưng theo phác đồ điều trị (ACTG175)

Đặc trưng	ZDV only	ZDV + Zai	ddI only	ZDV + ddI
CD4 Count	0.22	0.26	0.24	0.28
Prior ZDV	0.18	0.05	0.035	0.04
CD4/CD8 Ratio	0.12	0.135	0.13	0.145
Symptom Status	0.085	0.085	0.10	0.095
Age	0.075	0.095	0.09	0.08

Các kết quả này cho thấy khả năng của phương pháp phân tích đặc trưng trong việc giải thích hành vi của mô hình, đồng thời xác nhận rằng các dự báo của mô hình đề xuất phù hợp với kiến thức lâm sàng hiện có. Việc phân tách tầm quan trọng theo từng phác đồ điều trị cũng hỗ trợ cá thể hóa điều trị và nâng cao tính diễn giải của mô hình.

3.4.8.8. Kiểm định ý nghĩa thống kê giữa các mô hình

Bảng 3.26 cho thấy mô hình đề xuất đạt cải thiện đáng kể về C-index so với tất cả các phương pháp so sánh. Mức cải thiện dao động từ 0.05 đến 0.16, trong đó lớn nhất so với Cox PH. Tất cả các so sánh đều có ý nghĩa thống kê ($p < 0.05$), khẳng định hiệu quả vượt trội của mô hình đề xuất trong việc phân biệt nguy cơ.

Bảng 3.26: So sánh ý nghĩa thống kê giữa các mô hình

Mô hình 1	Mô hình 2	Chênh lệch C-index	p-value
Cox PH	MH đề xuất	0.16	0.00002
DeepSurv	MH đề xuất	0.12	0.00003
DeepHit	MH đề xuất	0.08	0.0012
SurvTRACE	MH đề xuất	0.05	0.0178

3.5. Thảo luận về các tham số của mô hình

Hiệu quả của mô hình chịu ảnh hưởng bởi việc lựa chọn các tham số thiết kế và tham số huấn luyện. Các tham số này quyết định khả năng biểu diễn của mô hình, độ ổn định trong quá trình học và chi phí tính toán. Do đó, việc lựa chọn tham số cần bảo đảm sự cân bằng giữa độ phức tạp của mô hình và khả năng học từ dữ liệu.

Đối với các tham số kiến trúc, việc tăng số chiều biểu diễn, kích thước bộ nhớ nguyên mẫu có thể giúp mô hình biểu diễn được nhiều mẫu tiến triển bệnh hơn. Tuy nhiên, số lượng tham số cần tối ưu cũng tăng theo, kéo theo chi phí tính toán và nguy cơ quá khớp. Ngược lại, nếu các tham số này được lựa chọn quá nhỏ thì khả năng biểu diễn của mô hình có thể bị hạn chế.

Đối với các tham số của hàm mất mát đa nhiệm, việc lựa chọn các trọng số phù hợp giúp cân bằng mức độ đóng góp của từng mục tiêu học trong quá trình tối ưu. Nếu một mục tiêu được ưu tiên quá mức, quá trình học có thể mất cân bằng và ảnh hưởng đến các mục tiêu còn lại.

Các tham số huấn luyện như tốc độ học, kích thước lô dữ liệu, số epoch và chiến lược học theo chương trình cũng ảnh hưởng đến quá trình hội tụ của mô hình. Việc lựa chọn các tham số này cần phù hợp với đặc điểm của bộ dữ liệu và độ phức tạp của mô hình nhằm duy trì sự ổn định trong quá trình huấn luyện.

Nhìn chung, việc lựa chọn các tham số của mô hình luôn là sự cân bằng giữa khả năng biểu diễn, hiệu quả tính toán và khả năng tổng quát hóa. Mặc dù mô hình đề xuất có khả năng kết hợp nhiều nguồn thông tin để nâng cao chất lượng biểu diễn trạng thái bệnh, mô hình cũng có một số hạn chế như số lượng siêu tham số tương đối lớn, việc hiệu chỉnh tham số đòi hỏi nhiều thời gian và chi phí tính toán cao hơn so với các mô hình có cấu trúc đơn giản.

3.6. Kết luận chương

Trong chương này, nghiên cứu đã đề xuất và phân tích một mô hình học sâu cho bài toán phân tích sống sót trong bối cảnh điều trị HIV, với trọng tâm là mô hình hóa trạng thái bệnh tiềm ẩn và động lực diễn tiến theo thời gian. Khác với các phương pháp phân tích sống sót truyền thống thường dựa chủ yếu trên các đặc trưng nền hoặc các giả định tuyến tính, mô hình đề xuất được xây dựng nhằm phản ánh tốt hơn bản chất động và chịu tác động của điều trị trong quá trình diễn tiến bệnh HIV.

Về mặt phương pháp luận, nội dung chương đã chỉ ra rằng nguy cơ xảy ra biến cố không nên được mô hình hóa trực tiếp từ các biến quan sát tại từng thời điểm, mà cần được xem như hệ quả của một trạng thái bệnh tiềm ẩn biến đổi theo thời gian. Trên cơ sở đó, mô hình đề xuất kết hợp các thành phần mã hóa đặc trưng, treatment head, bộ nhớ thời gian và mô hình state-space chọn lọc nhằm học biểu diễn trạng thái bệnh một cách ổn định, ngay cả trong điều kiện dữ liệu chuỗi thời gian ngắn hoặc không đều.

Các thực nghiệm được thiết kế trên cả hai kịch bản dữ liệu tĩnh và dữ liệu chuỗi thời gian cho thấy mô hình đề xuất không chỉ duy trì hiệu quả trong các bối cảnh dữ liệu hạn chế, mà còn phát huy rõ rệt lợi thế khi thông tin theo thời gian và điều trị được khai thác đầy đủ. Phân tích ablation cũng làm rõ vai trò của từng thành phần trong mô hình, qua đó cho thấy việc tích hợp thông tin điều trị và mô hình hóa động lực học là các yếu tố quan trọng trong phân tích sống sót HIV.

Nhìn chung, các kết quả đạt được trong chương này cho thấy mô hình đề xuất là một phương pháp phân tích sống sót linh hoạt, có cơ sở phương pháp luận rõ ràng và phù hợp với đặc điểm dữ liệu HIV trong thực tế lâm sàng.

Một số hạn chế vẫn cần được ghi nhận đối với mô hình phân tích sống sót được đề xuất. Mặc dù các mô hình truyền thống như Kaplan–Meier và Cox proportional hazards có nền tảng thống kê vững chắc, với các giả định rõ ràng và khả năng diễn giải cao, mô hình học sâu trong nghiên cứu này được thiết kế nhằm khai thác các quan

hệ phi tuyến, động học thời gian phức tạp và các tương tác đa chiều mà các phương pháp tuyến tính cổ điển có thể khó biểu diễn đầy đủ. Tuy nhiên, cũng giống như nhiều phương pháp học sâu khác, nguy cơ overfitting và hạn chế về khả năng tổng quát hóa vẫn cần được cân nhắc, đặc biệt trong các bối cảnh dữ liệu lâm sàng có quy mô hạn chế hoặc phân bố khác biệt. Trong nghiên cứu này, nguy cơ này đã được giảm thiểu thông qua các chiến lược như kiểm định chéo, regularization, đánh giá calibration và kiểm chứng trên nhiều bộ dữ liệu. Dù vậy, việc xác nhận thêm trên các cohort lâm sàng độc lập từ các bối cảnh khác nhau vẫn là cần thiết để cung cấp bằng chứng mạnh mẽ hơn về khả năng tổng quát hóa và tính ứng dụng thực tế của mô hình.

Mặc dù mô hình phân tích sống sót cho phép dự báo nguy cơ và thời gian xảy ra biến cố, mô hình này chưa cung cấp thông tin về tác động của các lựa chọn điều trị đối với từng bệnh nhân. Do đó, để mở rộng từ bài toán dự báo sang bài toán hỗ trợ ra quyết định, nghiên cứu tiếp tục xây dựng mô hình suy luận nhân quả nhằm ước lượng hiệu quả của các can thiệp ở mức cá thể.

Chương 4. ƯỚC LƯỢNG NHÂN QUẢ VÀ DỰ BÁO KẾT CỤC TRONG GIÁM SÁT VÀ ĐIỀU TRỊ HIV

Chương này giới thiệu phương pháp ước lượng nhân quả và dự báo kết cục trong giám sát và điều trị HIV thông qua việc kết hợp giữa học biểu diễn sâu và ước lượng nhân quả. Khác với các cách tiếp cận sử dụng cơ chế kết hợp cố định, phương pháp được đề xuất tập trung tối ưu hóa cách kết hợp giữa hai thành phần này theo đặc điểm của từng trường hợp cụ thể thông qua cơ chế tổ hợp thích ứng dựa trên học tăng cường (RL-ensemble). Cách tiếp cận này nhằm nâng cao khả năng cá nhân hóa trong dự báo kết cục và hỗ trợ ra quyết định can thiệp trong giám sát và điều trị HIV.

4.1. Mô tả bài toán

Xét một không gian xác suất $(\Omega, \mathcal{F}, \mathbb{P})$, trong đó mỗi đơn vị quan sát tương ứng với một trường hợp được ghi nhận trong hệ thống giám sát và quản lý điều trị HIV. Với mỗi cá thể $i \in \{1, \dots, n\}$, ta quan sát một bộ biến ngẫu nhiên (X_i, T_i, Y_i) , trong đó $X_i \in \mathcal{X} \subseteq \mathbb{R}^p$ biểu diễn tập các đặc trưng nền (bao gồm đặc trưng nhân khẩu học, hành vi, lâm sàng và/hoặc tri thức liên quan HIV đã được mã hóa), $T_i \in \mathcal{T}$ là biến can thiệp/điều trị rời rạc với tập giá trị hữu hạn $\mathcal{T}$ (có thể nhị phân hoặc đa mức), và $Y_i \in \{0, 1\}$ là kết cục nhị phân quan tâm. Trong bối cảnh nghiên cứu này, T_i có thể đại diện cho trạng thái/phoi nhiễm can thiệp (ví dụ có/không sàng lọc) hoặc nhóm phác đồ điều trị (ví dụ 4 nhóm), và Y_i phản ánh trạng thái kết cục sức khỏe (ví dụ có/không xảy ra kết cục).

Theo khung lý thuyết kết cục tiềm năng (potential outcomes), với mỗi cá thể i và mỗi mức can thiệp $t \in \mathcal{T}$, tồn tại một biến ngẫu nhiên $Y_i(t)$ biểu diễn kết cục tiềm năng nếu cá thể đó được gán can thiệp ở mức t . Tập hợp các kết cục tiềm năng $\{Y_i(t) : t \in \mathcal{T}\}$ được giả định tồn tại đồng thời, mặc dù trong thực tế chỉ quan sát được một kết cục tương ứng với mức can thiệp đã xảy ra.

Mối quan hệ giữa kết cục quan sát và kết cục tiềm năng được xác định thông qua giả định nhất quán (consistency): nếu cá thể i thực sự nhận mức can thiệp t (tức $T_i = t$) thì kết cục quan sát bằng đúng kết cục tiềm năng dưới mức can thiệp đó,

$$Y_i = Y_i(T_i).$$

Bài toán suy luận nhân quả đặt ra mục tiêu ước lượng tác động của can thiệp lên kết cục thông qua việc so sánh các kết cục tiềm năng tương ứng với các mức can thiệp khác nhau. Trong nghiên cứu này, thay vì ước lượng trực tiếp hiệu quả can thiệp $\tau(x)$, mô hình được thiết kế để học hàm kỳ vọng có điều kiện của kết cục tiềm năng theo đặc trưng nền như sau:

$$\mu(x, t) = \mathbb{E}[Y(t) \mid X = x], \quad t \in \mathcal{T}.$$

Với mỗi cá thể i , ký hiệu $\mu_i(t) = \mathbb{E}[Y_i(t) \mid X_i]$ là kỳ vọng có điều kiện “thật” và $\hat{\mu}_i(t)$ là ước lượng do mô hình sinh ra.

Từ các ước lượng $\hat{\mu}_i(t)$, hiệu quả nhân quả cá thể giữa hai mức can thiệp bất kỳ $a, b \in \mathcal{T}$ (Individualized Treatment Effect - ITE) được định nghĩa như sau:

$$\tau_i^{(a,b)} = \hat{\mu}_i(a) - \hat{\mu}_i(b).$$

Tương tự, hiệu quả nhân quả có điều kiện theo đặc trưng nền (conditional effect) giữa hai mức can thiệp a, b có thể viết dưới dạng hàm:

$$\tau^{(a,b)}(x) = \mu(x, a) - \mu(x, b).$$

Trong trường hợp đặc biệt can thiệp nhị phân $\mathcal{T} = \{0, 1\}$, các biểu thức trên thu về dạng quen thuộc $\tau_i = \hat{\mu}_i(1) - \hat{\mu}_i(0)$ và $\tau(x) = \mu(x, 1) - \mu(x, 0)$.

Ở cấp độ quần thể, hiệu quả can thiệp trung bình giữa hai mức a, b được định nghĩa như sau:

$$\tau^{(a,b)} = \mathbb{E}[Y(a) - Y(b)],$$

Tuy nhiên, trong bối cảnh dịch bệnh HIV, hiệu quả của can thiệp thường khác nhau đáng kể giữa các trường hợp tùy theo đặc trưng nền. Do đó, thay vì chỉ quan tâm đến hiệu quả trung bình, bài toán thực tiễn tập trung vào việc ước lượng hiệu quả có điều kiện $\tau^{(a,b)}(x)$ hoặc hiệu quả cá thể $\tau_i^{(a,b)}$ để hỗ trợ phân tầng và ra quyết định theo từng nhóm và cá thể.

Để đảm bảo khả năng nhận dạng các đại lượng nhân quả từ dữ liệu quan sát, nghiên cứu dựa trên các giả định chuẩn trong khung kết cục tiềm năng (potential outcomes). Giả định không nhiễu có điều kiện (ignorability/unconfoundedness) yêu cầu rằng với mọi $t \in \mathcal{T}$,

$$Y(t) \perp T \mid X. \quad (4.1)$$

nghĩa là khi đã điều kiện theo X , việc gán can thiệp độc lập với các kết cục tiềm năng. Giả định chồng lấn (positivity/overlap) yêu cầu xác suất quan sát mỗi mức can thiệp là không suy biến,

$$0 < \mathbb{P}(T = t \mid X = x) < 1, \quad \forall t \in \mathcal{T}. \quad (4.2)$$

với mọi x thuộc miền hỗ trợ của X . Ngoài ra, giả định giả định ổn định giá trị kết cục giữa các đơn vị (Stable Unit Treatment Value Assumption - SUTVA) đảm bảo (i) không có can thiệp chéo giữa các đơn vị (kết cục của một cá thể không phụ thuộc can thiệp của cá thể khác) và (ii) mỗi mức can thiệp t được định nghĩa rõ ràng.

Trong bối cảnh dữ liệu HIV thực tế, các giả định trên có thể bị vi phạm một phần do các yếu tố gây nhiễu không đo lường hoặc do cấu trúc dữ liệu không hoàn hảo. Do đó, bài toán suy luận nhân quả không chỉ là bài toán ước lượng thống kê, mà còn là bài toán thiết kế phương pháp sao cho có khả năng giảm thiểu sai lệch do mô hình hóa, thích ứng với dữ liệu nhiều loại biến (nhân khẩu học–hành vi–lâm sàng) và duy trì độ tin cậy của ước lượng trong các điều kiện không lý tưởng.

Trong chương này, bài toán suy luận nhân quả được tiếp cận như một bài toán học hàm có điều kiện $\mu(x, t)$ trên toàn bộ không gian đặc trưng và tập mức can thiệp $\mathcal{T}$,

từ đó suy ra các đại lượng hiệu quả nhân quả $\tau_i^{(a,b)}$ và $\tau^{(a,b)}(x)$ phục vụ đánh giá tác động can thiệp ở mức cá thể và có điều kiện theo đặc trưng.

Từ góc độ mô hình hóa, việc học hàm $\mu(x, t)$ có thể được xem như một bài toán học biểu diễn có điều kiện, trong đó mô hình cần đồng thời nắm bắt cấu trúc của đặc trưng nền X và ảnh hưởng của biến can thiệp T . Trong bối cảnh dữ liệu HIV với đặc trưng đa dạng và phụ thuộc phức tạp, các phương pháp học sâu cung cấp khả năng học các biểu diễn phi tuyến linh hoạt cho $\mu(x, t)$.

Tuy nhiên, việc ước lượng chính xác hiệu quả nhân quả còn phụ thuộc vào khả năng giảm sai lệch do mô hình hóa và xử lý các vấn đề như mất cân bằng giữa các nhóm can thiệp và nhiễu ẩn. Do đó, các phương pháp kết hợp giữa học biểu diễn sâu và các nguyên lý ước lượng nhân quả bền vững được xem là hướng tiếp cận phù hợp để đảm bảo tính chính xác và ổn định của các ước lượng $\hat{\mu}(x, t)$.

4.2. Phân tích lý thuyết và thiết kế mô hình

Bài toán suy luận nhân quả từ dữ liệu quan sát đòi hỏi đồng thời hai yêu cầu: học được biểu diễn đặc trưng phản ánh đúng trạng thái của từng cá thể và ước lượng chính xác hiệu quả can thiệp trong điều kiện tồn tại các yếu tố gây nhiễu. Các phương pháp truyền thống như propensity score hay inverse probability weighting có cơ sở lý thuyết vững chắc nhưng gặp hạn chế khi xử lý dữ liệu có nhiều quan hệ phi tuyến và tương tác phức tạp. Trong khi đó, các mô hình học sâu như TARNet [29] và DragonNet [30] cải thiện khả năng học biểu diễn nhưng sử dụng một cơ chế kết hợp cố định giữa học biểu diễn và ước lượng nhân quả cho toàn bộ dữ liệu. Đối với dữ liệu HIV có tính dị biệt cao, cách tiếp cận này có thể chưa đủ linh hoạt để thích ứng với từng cá thể.

Xuất phát từ nhận định trên, nghiên cứu đề xuất một khung mô hình gồm bốn thành phần: (i) mô hình hóa cấu trúc nhân quả, (ii) học biểu diễn đặc trưng, (iii) ước lượng hiệu quả can thiệp và (iv) điều phối thích ứng bằng học tăng cường.

(1) *Mô hình hóa cấu trúc nhân quả.* Để giảm sai lệch do nhiễu và bảo đảm tính hợp lệ của suy luận nhân quả, nghiên cứu xây dựng đồ thị nhân quả biểu diễn mối quan hệ giữa các biến. Cấu trúc đồ thị được hình thành từ sự kết hợp giữa tri thức chuyên gia và khám phá nhân quả từ dữ liệu, qua đó vừa bảo đảm khả năng diễn giải vừa tận dụng được thông tin thực tế của dữ liệu quan sát. Đồ thị nhân quả đồng thời là cơ sở để lựa chọn các biến điều chỉnh trong quá trình ước lượng hiệu quả can thiệp.

(2) *Học biểu diễn đặc trưng.* Dữ liệu HIV bao gồm nhiều đặc trưng lâm sàng có quan hệ phi tuyến và các quan sát theo thời gian. Vì vậy, nghiên cứu lựa chọn TITAN [98], một biến thể của Transformer được bổ sung cơ chế bộ nhớ học được. Khác với Transformer truyền thống, TITAN có khả năng lưu trữ và truy hồi các biểu diễn đã học trong quá trình huấn luyện, giúp khai thác các mẫu tiến triển tương đồng giữa các bệnh nhân và nâng cao khả năng tổng quát hóa của mô hình.

(3) *Ước lượng hiệu quả can thiệp.* Trên cơ sở biểu diễn đặc trưng đã học, nghiên cứu sử dụng DRLearner [31] để ước lượng hiệu quả can thiệp ở mức cá thể. Phương pháp này kết hợp đồng thời mô hình kết cục và mô hình xác suất can thiệp, giúp giảm sai lệch khi một trong hai mô hình bị sai đặc tả và nâng cao tính ổn định của ước lượng

trên dữ liệu quan sát.

(4) *Điều phối thích ứng bằng học tăng cường*. Mặc dù học biểu diễn và ước lượng nhân quả đều đóng vai trò quan trọng, không tồn tại một chiến lược kết hợp duy nhất phù hợp với mọi cá thể. Do đó, nghiên cứu sử dụng học tăng cường [99] để xây dựng cơ chế điều phối thích ứng giữa các thành phần của mô hình. Tác tử học tăng cường học chính sách kết hợp dựa trên đặc điểm của từng mẫu dữ liệu nhằm cân bằng giữa chất lượng biểu diễn và độ tin cậy của ước lượng nhân quả, từ đó cải thiện hiệu quả dự báo và khả năng khái quát của mô hình.

Từ các phân tích trên, nghiên cứu xây dựng một kiến trúc thống nhất kết hợp đồ thị nhân quả, học biểu diễn sâu, ước lượng nhân quả và điều phối thích ứng. Thiết kế này tận dụng ưu điểm của từng thành phần, đồng thời khắc phục hạn chế của các mô hình kết hợp cố định khi áp dụng cho dữ liệu HIV có tính dị biệt và phức tạp.

4.3. Xây dựng đồ thị nhân quả

Trong thực tế, việc xây dựng một đồ thị nhân quả có hướng không chu trình (DAG) đáng tin cậy là bước nền tảng nhằm đảm bảo tính hợp lệ của quá trình lựa chọn biến và suy luận nhân quả. Để đạt được điều này, nghiên cứu xem xét hai nguồn đồ thị ứng viên: (i) đồ thị được xây dựng dựa trên tri thức chuyên gia và (ii) đồ thị được suy ra trực tiếp từ dữ liệu.

1. *Đồ thị theo chuyên gia* G_{EX} : được thiết lập dựa trên tri thức dịch tễ học và kinh nghiệm thực hành, đóng vai trò như một cấu trúc tham chiếu với ý nghĩa nhân quả rõ ràng.
2. *Đồ thị từ dữ liệu* G_{REX} : được xây dựng thông qua một khung khám phá nhân quả dựa trên giải thích mô hình và định hướng cạnh, kết hợp các kỹ thuật lượng hóa đóng góp đặc trưng và kiểm tra độc lập phần dư, sau đó loại bỏ chu trình để đảm bảo tính chất DAG.

Để lựa chọn đồ thị cuối cùng phục vụ mô hình hóa, nghiên cứu đề xuất một chiến lược đánh giá dựa trên độ ổn định cấu trúc và tính nhất quán nhân quả. Cụ thể, một đồ thị được xem là đáng tin cậy nếu các thành phần nhân quả quan trọng, bao gồm biến gây nhiễu, biến trung gian và biến công cụ, được duy trì tương đối ổn định dưới các nhiễu nhỏ trên cấu trúc đồ thị. Ký hiệu thước đo ổn định là $Stability(G)$, đồ thị cuối cùng được xác định như sau:

$$G_{final} = \arg \max_{G \in \{G_{EX}, G_{REX}\}} Stability(G). \quad (4.3)$$

Sau khi xác định G_{final} , tập đặc trưng đầu vào X được trích xuất dựa trên các nguyên tắc nhân quả: (i) giữ lại các biến gây nhiễu cần điều chỉnh, (ii) loại bỏ các biến hậu can thiệp hoặc có khả năng đóng vai trò biến trung gian hoặc biến hội tụ khi mục tiêu là ước lượng tổng tác động, và (iii) ưu tiên tập biến tối giản nhằm giảm nhiễu và cải thiện khả năng tổng quát hóa của mô hình.

Dựa trên nền tảng lựa chọn đồ thị và đặc trưng nêu trên, bài toán ước lượng hiệu

quả can thiệp cá thể được biểu diễn dưới dạng học một hàm ánh xạ

$$x \mapsto \tau(x),$$

trong đó x là véc-tơ đặc trưng đã được chuẩn hóa theo logic nhân quả. Tuy nhiên, trong dữ liệu HIV thực tế, mối quan hệ giữa X , T và Y thường có tính phi tuyến, chứa nhiều tương tác bậc cao và có thể chịu ảnh hưởng của sai lệch phân phối giữa các nhóm hoặc theo thời gian.

Do đó, nghiên cứu đề xuất một chiến lược kết hợp hai hướng suy luận bổ sung lẫn nhau: (i) học biểu diễn sâu có điều kiện theo cấu trúc nhân quả nhằm nắm bắt các quan hệ phi tuyến và tương tác phức tạp, và (ii) một bộ ước lượng nhân quả có tính bền vững kép nhằm đảm bảo tính đúng đắn ngay cả khi mô hình bị sai đặc tả. Cơ chế điều phối thích ứng theo từng cá thể được xây dựng trên cơ sở cấu trúc đồ thị nhân quả, trong đó đồ thị không chỉ đóng vai trò diễn giải mà còn là một tầng kiểm soát nhân quả xuyên suốt quá trình suy luận.

Cụ thể, đồ thị nhân quả được sử dụng để (i) xác định tập biến đầu vào hợp lệ cho bài toán ước lượng hiệu quả can thiệp, (ii) phân loại vai trò nhân quả của các biến như biến gây nhiễu, biến trung gian và biến công cụ, và (iii) giảm thiểu rủi ro sai lệch do điều chỉnh sai trong bối cảnh dữ liệu quan sát. Trên cơ sở đó, quy trình xây dựng đồ thị từ dữ liệu được triển khai thông qua các bước xác định cấu trúc liên kết, định hướng cạnh và loại bỏ chu trình, đồng thời được đánh giá thông qua phân tích độ ổn định và tính nhất quán với tri thức chuyên gia.

Ta ký hiệu một đồ thị nhân quả dưới dạng đồ thị có hướng không chu trình (DAG) là:

$$G = (V, E),$$

trong đó V là tập các biến quan sát và E là tập các cạnh có hướng biểu diễn quan hệ nhân quả giả định. Trong mỗi bài toán, tồn tại một biến can thiệp T và một biến kết cục Y . Mục tiêu của DAG trong khuôn khổ này không phải là khôi phục chính xác cơ chế sinh dữ liệu thực sự, mà là xây dựng một cấu trúc nhân quả đủ đáng tin cậy để hỗ trợ suy luận can thiệp từ dữ liệu quan sát.

Hai đồ thị nhân quả ứng viên: đồ thị tri thức và đồ thị sinh từ dữ liệu.

Hai đồ thị nhân quả ứng viên được xây dựng cho mỗi bộ dữ liệu. Đồ thị thứ nhất, ký hiệu G_{EX} , được xác định dựa trên tri thức miền và kinh nghiệm chuyên gia trong lĩnh vực HIV. Đồ thị này phản ánh các mối quan hệ nhân quả đã được chấp nhận rộng rãi trong thực hành dịch tễ học và đóng vai trò như một cấu trúc nhân quả tiên nghiệm.

Đồ thị thứ hai, ký hiệu G_{REX} , được suy ra trực tiếp từ dữ liệu thông qua một khung khám phá nhân quả dựa trên giải thích mô hình, được phát triển dựa trên mô hình REX, được trình bày trong nghiên cứu tại [100] nhưng được đơn giản hóa để phù hợp với độ phức tạp của dữ liệu y tế. Đồ thị G_{REX} được xem như một giả thuyết nhân quả sinh từ dữ liệu, có khả năng bổ sung hoặc điều chỉnh các giả định từ đồ thị tri thức.

Thay vì mặc định lựa chọn một trong hai đồ thị, mô hình sử dụng đồ thị tri thức như một chuẩn tham chiếu để đánh giá độ hợp lý và độ nhạy của đồ thị sinh từ dữ liệu.

Xây dựng đồ thị sinh từ dữ liệu bằng khung REX đơn giản hóa.

Quá trình xây dựng G_{REX} được thực hiện theo ba bước chính: xác định khung liên kết (skeleton), định hướng cạnh, và loại bỏ chu trình.

1. Xác định khung liên kết bằng SHAP và bootstrap.

Với mỗi biến $X_i \in V$, ta huấn luyện một mô hình hồi quy tuyến tính nhằm xấp xỉ mối quan hệ giữa X_i và các biến còn lại $X_{\setminus i}$. Mức đóng góp của từng biến $X_j \in X_{\setminus i}$ vào việc dự đoán X_i được lượng hóa bằng giá trị Shapley $\phi_{i,j}$. Để tăng độ vững, quá trình này được lặp lại trên nhiều mẫu bootstrap, sau đó các giá trị SHAP [96] được gom cụm trong không gian đặc trưng và ngưỡng hóa theo tần suất xuất hiện. Một cạnh vô hướng giữa (X_j, X_i) được đưa vào skeleton nếu X_j thường xuyên xuất hiện trong cụm SHAP có ảnh hưởng cao nhất đối với X_i .

2. Định hướng cạnh bằng mô hình nhiễu cộng (ANM). trong mô hình ANM [101] Với mỗi cạnh vô hướng (X_i, X_j) trong skeleton, ta xem xét hai mô hình hồi quy tuyến tính đối nghịch:

$$\begin{aligned} X_i &= f_i(X_j) + \varepsilon_{i|j}, \\ X_j &= f_j(X_i) + \varepsilon_{j|i}. \end{aligned} \quad (4.4)$$

Hướng nhân quả được xác định dựa trên mức độ độc lập giữa phần dư và biến giải thích, được đo bằng thước đo mức độ phụ thuộc hay độc lập giữa hai biến HSIC (Hilbert-Schmidt Independence Criterion) [102]. Chiều hồi quy tạo ra phần dư độc lập hơn được xem là hướng nhân quả phù hợp hơn.

3. Loại bỏ chu trình bằng độ bất nhất SHAP. Sau khi định hướng, đồ thị có thể chứa chu trình. Để đảm bảo tính DAG, các chu trình được loại bỏ bằng cách xác định cạnh kém ổn định nhất trong chu trình. Mức độ bất nhất của một cạnh được đo bằng độ sai khác giữa các giá trị SHAP và khả năng giải thích biến mục tiêu, thông qua hệ số xác định R^2 . Cạnh có độ bất nhất cao nhất được ưu tiên loại bỏ cho đến khi đồ thị không còn chu trình.

Phân tích nhạy cảm dựa trên đồ thị tri thức.

Đồ thị sinh từ dữ liệu G_{REX} không được sử dụng trực tiếp, mà được đánh giá thông qua một phân tích nhạy cảm (sensitivity analysis) dựa trên đồ thị tri thức G_{EX} . Trong phân tích này, G_{EX} đóng vai trò như một chuẩn tham chiếu nhân quả, còn G_{REX} được xem là giả thuyết cần được kiểm chứng về mức độ phù hợp.

Cụ thể, một phép nhiễu cấu trúc có chủ đích được áp dụng cho cả hai đồ thị bằng cách loại bỏ cạnh trực tiếp từ biến can thiệp đến biến kết cục:

$$G \longrightarrow G^{(-T \rightarrow Y)}. \quad (4.5)$$

Với mỗi đồ thị trước và sau nhiễu cấu trúc, ta trích xuất các tập thành phần nhân quả, bao gồm tập biến gây nhiễu $\mathcal{C}(G)$, tập biến trung gian $\mathcal{M}(G)$ và tập biến công cụ $\mathcal{I}(G)$.

Độ nhạy của đồ thị sinh từ dữ liệu được đánh giá thông qua mức độ tương đồng giữa các tập thành phần nhân quả của G_{REX} và G_{EX} , cũng như giữa $G_{\text{REX}}^{(-T \rightarrow Y)}$ và $G_{\text{EX}}^{(-T \rightarrow Y)}$. Một đồ thị sinh từ dữ liệu được xem là đáng tin cậy nếu cấu trúc vai trò nhân quả của các biến không thay đổi mạnh khi so sánh với đồ thị tri thức dưới cùng một phép nhiễu.

Xác định đồ thị cuối và trích xuất tập biến hợp lệ.

Trên cơ sở phân tích nhạy cảm, mô hình đề xuất không nhằm tìm ra một đồ thị nhân quả tối ưu tuyệt đối, mà xác nhận rằng đồ thị sinh từ dữ liệu đủ nhất quán với đồ thị tri thức để được sử dụng làm nền tảng cho suy luận tiếp theo. Khi mức độ nhạy cảm vượt quá ngưỡng chấp nhận, các quan hệ nhân quả trong G_{REX} sẽ được điều chỉnh hoặc giản lược theo cấu trúc của G_{EX} .

Đồ thị cuối cùng, ký hiệu G_{final} , được sử dụng để trích xuất tập biến đầu vào hợp lệ X cho bài toán ước lượng hiệu quả can thiệp. Cụ thể, các biến gây nhiễu được giữ lại để điều chỉnh, các biến trung gian và hậu can thiệp được xử lý thận trọng tùy theo mục tiêu ước lượng (tổng tác động hay tác động trực tiếp), và các biến không liên quan nhân quả bị loại bỏ. Bước này đảm bảo rằng các mô hình học sâu và cơ chế điều phối được trình bày ở các mục tiếp theo được xây dựng trên một không gian đặc trưng đã được kiểm soát về mặt nhân quả, thay vì chỉ dựa trên tương quan thống kê.

4.4. Phương pháp đề xuất

4.4.1. Mô hình đề xuất tổng thể

Để giải quyết bài toán ước lượng hiệu ứng nhân quả cá thể từ dữ liệu quan sát, nghiên cứu đề xuất một mô hình tổng thể kết hợp giữa xây dựng đồ thị nhân quả, học biểu diễn sâu, ước lượng nhân quả và điều phối thích ứng bằng học tăng cường. Quy trình xử lý của mô hình được thực hiện qua các bước sau.

Bước 1. Xây dựng đồ thị nhân quả. Từ dữ liệu đầu vào, hai đồ thị nhân quả ứng viên được xây dựng dựa trên tri thức chuyên gia và dữ liệu quan sát. Sau đó, phân tích độ nhạy được thực hiện để đánh giá và lựa chọn đồ thị phù hợp làm cơ sở cho quá trình suy luận nhân quả.

Bước 2. Trích xuất tập đặc trưng nhân quả. Dựa trên đồ thị nhân quả đã lựa chọn, các biến đầu vào được sàng lọc nhằm giữ lại các biến có ý nghĩa nhân quả đối với biến can thiệp và biến kết cục, đồng thời loại bỏ các biến không phù hợp như biến hậu can thiệp hoặc các biến không liên quan. Kết quả của bước này là ma trận đặc trưng $X \in \mathbb{R}^{n \times p}$, được sử dụng làm đầu vào cho mô hình.

Bước 3. Học biểu diễn sâu. Ma trận đặc trưng X được đưa vào mô hình TITAN để học biểu diễn tiềm ẩn của từng cá thể. Với cơ chế tự chú ý và bộ nhớ ngoài, TITAN có khả năng khai thác các quan hệ phi tuyến và tạo ra biểu diễn đặc trưng phục vụ dự báo kết cục dưới các kịch bản can thiệp.

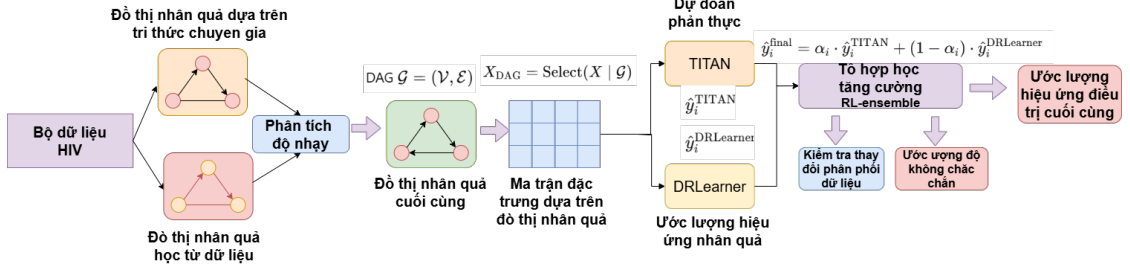
Bước 4. Ước lượng hiệu ứng nhân quả. Song song với TITAN, mô hình DRLearner sử dụng cùng tập đặc trưng đầu vào để xây dựng mô hình kết cục và mô hình xác suất can thiệp, từ đó ước lượng hiệu ứng nhân quả ở mức cá thể theo nguyên lý bền vững kép.

Bước 5. Điều phối thích ứng. Đầu ra của TITAN và DRLearner được đưa vào tác tử học tăng cường để học chính sách kết hợp động. Dựa trên đặc điểm của từng cá thể, tác tử xác định trọng số kết hợp phù hợp nhằm tận dụng ưu điểm của cả hai nhánh ước lượng.

Bước 6. Sinh kết quả đầu ra. Mô hình tạo ra xác suất kết cục dưới các kịch bản can thiệp, ước lượng hiệu ứng nhân quả cá thể và mức độ bất định của dự báo. Các kết quả

này được sử dụng để hỗ trợ đánh giá hiệu quả can thiệp và ra quyết định trong giám sát, điều trị HIV.

Bằng cách tổ chức quá trình xử lý theo các bước trên, mô hình kết hợp giữa học biểu diễn sâu, suy luận nhân quả và cơ chế điều phối thích ứng trong một kiến trúc thống nhất, góp phần nâng cao độ chính xác và tính linh hoạt của hệ thống hỗ trợ ra quyết định.



Hình 4.1: Kiến trúc tổng thể của mô hình ước lượng nhân quả đề xuất.

Hình 4.1 minh họa toàn bộ kiến trúc của mô hình đề xuất. Quy trình bắt đầu từ dữ liệu HIV/AIDS, sau đó xây dựng đồ thị nhân quả nhằm xác định các đặc trưng có ý nghĩa cho suy luận. Hai mô hình TITAN và DR Learner hoạt động song song để khai thác hai khía cạnh khác nhau của bài toán: học biểu diễn và ước lượng nhân quả. Cuối cùng, mô-đun học tăng cường đóng vai trò điều phối, kết hợp linh hoạt hai nguồn thông tin và tạo ra dự báo nhân quả cuối cùng kèm theo mức độ không chắc chắn.

Kiến trúc này cho phép tận dụng đồng thời sức mạnh của học sâu và các phương pháp nhân quả, đồng thời thích nghi với sự thay đổi của dữ liệu theo thời gian, từ đó nâng cao hiệu quả và độ tin cậy của hệ thống trong các bài toán y sinh học.

4.4.2. Học biểu diễn nhân quả sâu với kiến trúc TITAN

Mục tiêu của mô-đun TITAN trong mô hình là xây dựng một không gian biểu diễn tiềm ẩn Z từ các biến đầu vào X đã được kiểm soát nhân quả, sao cho các biểu diễn này vừa giàu thông tin dự báo, vừa ổn định cho suy luận phản thực. Khác với các bộ mã hóa thông thường, TITAN được thiết kế để xử lý tính dị thể cao và sự thay đổi phân phối của dữ liệu thông qua cơ chế tự chú ý và bộ nhớ ngoài.

1. *Ánh xạ biểu diễn và không gian tiềm ẩn.* Với mỗi mẫu i , véc-tơ đặc trưng đầu vào $x_i \in \mathbb{R}^p$ được ánh xạ sang không gian tiềm ẩn như sau:

$$z_i = \Phi_{\text{TITAN}}(x_i; \theta_{\Phi}), \quad z_i \in \mathbb{R}^d, \quad (4.6)$$

trong đó θ_{Φ} là tập tham số của mô-đun TITAN. Không gian Z được thiết kế để làm miền chung cho cả dự báo kết cục và ước lượng hiệu quả can thiệp.

2. *Trộn đặc trưng cục bộ bằng tích chập một chiều.* Do các biến đầu vào có cấu trúc theo nhóm (nhân khẩu học, hành vi, lâm sàng), TITAN sử dụng tầng tích chập 1 chiều để nắm bắt tương tác cục bộ như sau:

$$h_i^{(0)} = \text{Conv1D}(x_i), \quad (4.7)$$

giúp giảm nhiễu và tăng khả năng học các mẫu phụ thuộc bậc thấp trước khi đưa vào

attention.

3. *cơ chế tự chú ý đa đầu và phụ thuộc toàn cục*. Các biểu diễn trung gian được cập nhật qua L lớp Transformer như sau:

$$\tilde{h}_i^{(\ell)} = \text{MHA}(h_i^{(\ell-1)}), \quad (4.8)$$

$$h_i^{(\ell)} = \text{LN} \left(h_i^{(\ell-1)} + \tilde{h}_i^{(\ell)} \right), \quad (4.9)$$

$$h_i^{(\ell)} = \text{LN} \left(h_i^{(\ell)} + \text{FFN}(h_i^{(\ell)}) \right), \quad (4.10)$$

với $\ell = 1, \dots, L$. Cơ chế này cho phép mô hình gán trọng số khác nhau cho từng đặc trưng theo từng cá thể.

4. *Bộ nhớ ngoài và ổn định biểu diễn*. Để tăng khả năng tổng quát hóa, TITAN tích hợp bộ nhớ ngoài $M \in \mathbb{R}^{K \times d}$:

$$g_i = \sigma(W_g h_i^{(L)} + b_g), \quad (4.11)$$

$$M_i = g_i \odot M_{i-1} + (1 - g_i) \odot h_i^{(L)}, \quad (4.12)$$

$$z_i = [h_i^{(L)} \parallel M_i]. \quad (4.13)$$

Bộ nhớ này giúp mô hình duy trì các mẫu biểu diễn ổn định khi gặp dữ liệu hiếm hoặc phân phối thay đổi.

5. *Tầng dự báo kết cục và cơ chế huấn luyện*. Biểu diễn cuối cùng z_i thu được từ khối Transformer đóng vai trò là đặc trưng tiềm ẩn tổng hợp của cá thể i , mã hóa thông tin từ tập biến quan sát X_i (và biến can thiệp nếu được đưa vào giai đoạn mã hóa). Dựa trên biểu diễn này, mô hình ước lượng xác suất của kết cục tiềm năng dưới một mức can thiệp giả định $t \in \mathcal{T}$ thông qua tầng tuyến tính theo sau bởi hàm sigmoid như sau:

$$\hat{y}_i(t) = \sigma(W_o z_i + b_o), \quad (4.14)$$

trong đó $\sigma(\cdot)$ là hàm sigmoid, còn W_o, b_o là các tham số cần học của tầng đầu ra. Khi kết cục là biến nhị phân, giá trị $\hat{y}_i(t)$ có thể được diễn giải như xác suất xảy ra sự kiện dưới mức can thiệp t .

Mục tiêu của mô-đun TITAN là ước lượng hàm kết quả có điều kiện (outcome regression function), được định nghĩa dưới khung tiềm năng (potential outcomes framework) như sau:

$$\mu(x_i, t) = \mathbb{E}[Y_i(t) \mid X_i]. \quad (4.15)$$

trong đó $Y_i(t)$ là kết cục tiềm năng nếu cá thể i được gán mức can thiệp t . Với kết cục nhị phân, ta có:

$$\mu(x_i, t) = \mathbb{P}(Y_i(t) = 1 \mid X_i). \quad (4.16)$$

Mô hình được huấn luyện bằng hàm mất mát binary cross-entropy trên dữ liệu quan sát như sau:

$$\mathcal{L}_{\text{TITAN}} = -\frac{1}{n} \sum_{i=1}^n [y_i \log \hat{y}_i(T_i) + (1 - y_i) \log(1 - \hat{y}_i(T_i))], \quad (4.17)$$

trong đó T_i là mức can thiệp quan sát được của cá thể i , và y_i là kết cục thực tế tương ứng. Việc tối thiểu hóa hàm mất mát này cho phép mô hình xấp xỉ phân phối có điều kiện của kết cục dưới trạng thái điều trị quan sát được. Dưới giả định không có nhiễu tiềm ẩn (ignorability) và tính nhất quán (consistency), hàm ước lượng này có thể được sử dụng để suy luận các kết cục tiềm năng chưa quan sát.

6. Đánh giá dưới các kịch bản can thiệp giả định. Sau khi huấn luyện, mô hình có thể được sử dụng để sinh ra các ước lượng kết cục tiềm năng dưới toàn bộ tập mức can thiệp $\mathcal{T}$. Cụ thể, đối với mỗi cá thể i , đặc trưng quan sát X_i được kết hợp lần lượt với từng mức can thiệp giả định $t \in \mathcal{T}$. Mỗi lần đánh giá như vậy tạo ra một dự đoán:

$$\hat{\mu}_i(t) = \hat{y}_i(t). \quad (4.18)$$

được diễn giải như xác suất kết cục nếu cá thể i nhận mức can thiệp t .

Cơ chế đánh giá lặp này cho phép mô hình xây dựng toàn bộ vectơ kết cục tiềm năng $\{\hat{\mu}_i(t)\}_{t \in \mathcal{T}}$ cho mỗi cá thể. Từ đó, các đại lượng nhân quả như hiệu quả điều trị cá thể hóa (ITE) có thể được ước lượng bằng cách lấy hiệu giữa các mức can thiệp khác nhau, ví dụ:

$$\widehat{\text{ITE}}_i = \hat{\mu}_i(1) - \hat{\mu}_i(0). \quad (4.19)$$

Nhờ vậy, TITAN không chỉ thực hiện dự báo kết cục dưới trạng thái quan sát, mà còn cung cấp cơ sở định lượng cho suy luận phản thực và ra quyết định cá thể hóa.

4.4.3. Ước lượng hiệu quả can thiệp cá thể hóa với DRLearner

1. DRLearner cho ước lượng hiệu ứng nhân quả. Trong kiến trúc đề xuất, DRLearner [103] được tích hợp như một nhánh song song với TITAN nhằm thực hiện ước lượng hiệu ứng nhân quả. DRLearner thuộc nhóm các phương pháp *doubly robust*, có khả năng ước lượng tác động của việc đã từng xét nghiệm HIV đến nguy cơ nhiễm HIV. Điểm đặc trưng của phương pháp này là kết hợp hai mô hình bổ sung lẫn nhau: mô hình hồi quy kết quả (outcome regression model) và mô hình xác suất điều trị (propensity score model). Nhờ đó, ước lượng thu được vẫn nhất quán ngay cả khi chỉ một trong hai mô hình được chỉ định đúng.

Tương tự như TITAN, ma trận đặc trưng đầu vào $\mathbf{X} \in \mathbb{R}^{n \times p}$ được xây dựng dựa trên cấu trúc DAG nhằm đảm bảo tính hợp lệ nhân quả. Với mỗi cá thể i , DRLearner trước hết huấn luyện hai mô hình hồi quy kết quả riêng biệt để ước lượng kỳ vọng có điều kiện của kết cục dưới hai trạng thái can thiệp và đối chứng như sau:

$$\mu_1(x_i) = \mathbb{E}[Y_i \mid X_i = x_i, T_i = 1], \quad \mu_0(x_i) = \mathbb{E}[Y_i \mid X_i = x_i, T_i = 0], \quad (4.20)$$

trong đó:

- $Y_i \in \{0, 1\}$ là trạng thái nhiễm HIV quan sát được,
- $T_i \in \{0, 1\}$ biểu thị cá thể có từng xét nghiệm HIV hay không,
- $x_i \in \mathbb{R}^p$ là vector đặc trưng của cá thể i .

Hai đại lượng $\mu_1(x_i)$ và $\mu_0(x_i)$ lần lượt biểu diễn kỳ vọng thực của kết cục tiềm năng dưới điều trị và không điều trị. Trong thực nghiệm, các ước lượng tương ứng thu được từ mô hình được ký hiệu là $\hat{\mu}_1(x_i)$ và $\hat{\mu}_0(x_i)$.

Tiếp theo, một mô hình propensity score được huấn luyện để ước lượng xác suất cá thể được điều trị:

$$\hat{e}(x_i) = \mathbb{P}(T_i = 1 \mid X_i = x_i). \quad (4.21)$$

Mô hình này có vai trò điều chỉnh sai lệch chọn mẫu giữa nhóm xét nghiệm và không xét nghiệm, và thường được triển khai bằng logistic regression hoặc các bộ phân loại xác suất khác.

Sau khi thu được cả hai thành phần, DRLearner kết hợp chúng thông qua công thức hiệu chỉnh doubly robust như sau:

$$\hat{\tau}_i = \left(\frac{t_i - \hat{e}(x_i)}{\hat{e}(x_i)(1 - \hat{e}(x_i))} \right) (y_i - \hat{\mu}_{t_i}(x_i)) + \hat{\mu}_1(x_i) - \hat{\mu}_0(x_i), \quad (4.22)$$

trong đó $\hat{\tau}_i$ là hiệu ứng can thiệp cá thể (ITE) [104]. Đại lượng này phản ánh sự khác biệt về kết cục giữa hai trạng thái có và không có can thiệp, sau khi đã hiệu chỉnh nhiều bằng cả mô hình kết cục (outcome) và mô hình xác suất điều trị (propensity).

Về mặt thực tiễn, công thức trên cung cấp cơ chế bảo vệ chống sai lệch mô hình. Nếu một trong hai mô hình con (outcome hoặc propensity) được chỉ định đúng, ước lượng hiệu ứng vẫn giữ được tính nhất quán. Điều này đặc biệt quan trọng trong bối cảnh dữ liệu HIV, nơi cấu trúc phụ thuộc giữa can thiệp và kết cục thường phức tạp và khó mô hình hóa hoàn toàn chính xác bằng một tiếp cận đơn lẻ.

Hiệu ứng can thiệp trung bình (ATE) được tính như sau:

$$\text{ATE} = \frac{1}{n} \sum_{i=1}^n \hat{\tau}_i. \quad (4.23)$$

ATE phản ánh tác động trung bình của biến can thiệp lên biến kết cục lên toàn bộ quần thể nghiên cứu.

2. Chi tiết triển khai DRLearner.

Mô hình hồi quy kết quả. Nghiên cứu sử dụng RandomForestRegressor với các siêu tham số sau:

- `n_estimators = 100`,
- `max_depth = 5`,
- `min_samples_leaf = 50`,
- kích hoạt bootstrap sampling.

Thiết lập này cho phép mô hình nắm bắt tương tác phi tuyến trong khi vẫn hạn chế quá khớp.

Mô hình propensity score. nghiên cứu sử dụng LogisticRegression với chuẩn hóa L2, tham số $C = 1.0$, tối đa 1000 vòng lặp, ngưỡng hội tụ 10^{-4} , và bộ giải liblinear.

Mở rộng cho can thiệp nhiều nhánh. Với tập dữ liệu gồm bốn nhánh điều trị ($T \in \{0, 1, 2, 3\}$), ước lượng được mở rộng thành:

$$\hat{\tau}_i(a, b) = \frac{\mathbf{1}[T_i = a]}{\hat{e}_a(X_i)} (Y_i - \hat{\mu}_a(X_i)) - \frac{\mathbf{1}[T_i = b]}{\hat{e}_b(X_i)} (Y_i - \hat{\mu}_b(X_i)) + \hat{\mu}_a(X_i) - \hat{\mu}_b(X_i), \quad (4.24)$$

trong đó $\hat{e}_t(X_i) = P(T_i = t | X_i)$ là generalized propensity score.

Chiến lược xác thực. Để đảm bảo khả năng tổng quát hóa và giảm hiện tượng quá khớp, nghiên cứu áp dụng chiến lược xác thực chéo năm phần (5-fold cross-validation). Cụ thể, trong mỗi lần lặp, các mô hình kết cục và mô hình xác suất điều trị được huấn luyện trên bốn phần dữ liệu và được đánh giá trên phần còn lại..

4.4.4. Điều phối thích ứng bằng học tăng cường cho ước lượng phản thực

Nghiên cứu xây dựng hai nhánh suy luận song song nhằm ước lượng kết cục phản thực ở mức cá thể. Cụ thể, nhánh TITAN học biểu diễn tiềm ẩn và ước lượng trực tiếp kết cục có điều kiện, trong khi nhánh DRLearner cung cấp ước lượng hiệu ứng dựa trên nguyên lý bền vững kép (doubly robust). Trong thực tế, độ chính xác của từng nhánh có thể thay đổi theo từng cá thể do tính dị biệt của đặc trưng, mức độ chồng lấn (overlap) giữa các nhóm can thiệp và đối chứng, cũng như sai lệch do đặc tả mô hình. Do đó, thay vì sử dụng một cơ chế kết hợp cố định, nghiên cứu đề xuất một chiến lược điều phối thích ứng dựa trên học tăng cường (reinforcement learning), cho phép lựa chọn hoặc gán trọng số động cho các nhánh ước lượng tại mức cá thể nhằm tối ưu hóa chất lượng dự báo phản thực.

Mô hình hóa bài toán điều phối dưới dạng quá trình ra quyết định Markov (MDP)

Trong mô hình MDP, tác tử học tăng cường (reinforcement learning agent) là thực thể ra quyết định, quan sát trạng thái của môi trường, lựa chọn hành động theo một chính sách $\pi(a | s)$, nhận phần thưởng và cập nhật chính sách nhằm tối đa hóa kỳ vọng phần thưởng.

Bài toán trộn dự báo được mô hình hóa dưới dạng một MDP. Cụ thể, tại mỗi cá thể i , tác tử học tăng cường quan sát trạng thái $s_i \in \mathcal{S}$, lựa chọn hành động $a_i \in \mathcal{A}$ theo chính sách $\pi(a_i | s_i)$, nhận phần thưởng $r_i \in \mathcal{R}$ và cập nhật chính sách nhằm tối đa hóa kỳ vọng phần thưởng $\mathbb{E}[r_i]$.

Trạng thái (State). Trạng thái của tác tử tại cá thể i được xây dựng từ các tín hiệu suy luận chính, bao gồm dự báo của TITAN, dự báo của DRLearner, và các tín hiệu phụ trợ phản ánh mức độ tin cậy (ví dụ bất định/độ lệch) như sau:

$$s_i = \left[\hat{y}_i^{\text{TITAN}}, \hat{y}_i^{\text{DRLearner}}, u_i \right], \quad (4.25)$$

trong đó $\hat{y}_i^{\text{TITAN}}$ và $\hat{y}_i^{\text{DRLearner}}$ là hai dự báo xác suất kết cục (counterfactual probability) do hai nhánh sinh ra, còn u_i là tín hiệu bất định (ví dụ phương sai dự báo, entropy, hoặc độ không ổn định theo bootstrap/MC dropout). Thiết kế này giúp tác tử học tăng cường không chỉ dựa vào giá trị dự báo mà còn cân nhắc độ tin cậy tương đối giữa hai nguồn.

Hành động (Action). Hành động a_i tương ứng với một trọng số trộn liên tục $\alpha_i \in [0, 1]$, thể hiện mức ưu tiên dành cho nhánh TITAN so với DRLearner:

$$a_i \equiv \alpha_i \in [0, 1]. \quad (4.26)$$

Chính sách (Policy). Chính sách $\pi_\theta(a | s)$ với tham số θ ánh xạ từ trạng thái sang phân phối hành động. Với hành động liên tục, chính sách sinh ra trọng số trộn α_i theo:

$$\alpha_i = \pi_\theta(a_i | s_i). \quad (4.27)$$

2. Cơ chế trộn dự báo và đầu ra cuối cùng

Với trọng số α_i do tác tử học tăng cường sinh ra, dự báo cuối cùng cho cá thể i được tính như một tổ hợp lồi (convex combination) của hai dự báo:

$$\hat{y}_i = \alpha_i \cdot \hat{y}_i^{\text{TITAN}} + (1 - \alpha_i) \cdot \hat{y}_i^{\text{DRLearner}}. \quad (4.28)$$

Cách trộn này cho phép mô hình thích ứng theo từng cá thể: trong các vùng dữ liệu mà biểu diễn sâu của TITAN đáng tin cậy hơn, tác tử có xu hướng tăng α_i ; ngược lại, trong các vùng mà nhánh DRLearner ổn định hơn (ví dụ overlap tốt hoặc mô hình DR hiệu quả), tác tử sẽ giảm α_i .

3. Thiết kế phần thưởng và mục tiêu tối ưu

Mục tiêu của tác tử học tăng cường là tìm chính sách trộn tối đa hóa độ chính xác dự báo cuối cùng. Vì bài toán đích là dự báo kết cục nhị phân, phần thưởng được định nghĩa dựa trên âm của hàm mất mát cross-entropy (Binary Cross-Entropy – BCE) giữa nhãn thật y_i và dự báo cuối $\hat{y}_i$:

$$r_i = -\mathcal{L}_{\text{BCE}}(\hat{y}_i, y_i) = -\left[y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)\right]. \quad (4.29)$$

Thiết kế này trực tiếp khuyến khích tác tử lựa chọn trọng số trộn giúp $\hat{y}_i$ gần với y_i nhất.

4. Huấn luyện tác tử bằng kiến trúc Actor–Critic

Để học chính sách trộn, nghiên cứu sử dụng kiến trúc Actor–Critic [105]. *Actor* sinh ra hành động $a_i = \alpha_i$ theo chính sách π_θ , trong khi *Critic* ước lượng hàm giá trị hành động $Q(s_i, a_i)$ nhằm đánh giá chất lượng của hành động đã chọn.

Cập nhật tham số actor được thực hiện theo công thức policy gradient:

$$\theta \leftarrow \theta + \eta \nabla_\theta \log \pi_\theta(a_i | s_i) Q(s_i, a_i), \quad (4.30)$$

trong đó η là tốc độ học (learning rate). Critic được huấn luyện để xấp xỉ hàm giá trị hành động, đóng vai trò giảm phương sai của gradient và giúp quá trình học ổn định hơn.

5. Các kỹ thuật ổn định huấn luyện

Để đảm bảo tác tử học tăng cường học ổn định trong không gian hành động liên tục và tránh hội tụ kém, nghiên cứu tích hợp ba kỹ thuật huấn luyện phổ biến:

Experience Replay. Các chuyển tiếp (s_i, a_i, r_i, s_{i+1}) được lưu trong bộ nhớ hồi tưởng $\mathcal{D}$. Trong quá trình huấn luyện, ta lấy mẫu ngẫu nhiên minibatch từ $\mathcal{D}$ để phá vỡ tương quan theo thứ tự dữ liệu và cải thiện hội tụ.

Nhiều Ornstein–Uhlenbeck (OU) cho khám phá. Để tăng khám phá trong không gian hành động liên tục, hành động được thêm nhiễu OU:

$$\mathcal{N}_{t+1} = \mathcal{N}_t + \theta_{ou}(\mu_{ou} - \mathcal{N}_t)\Delta t + \sigma_{ou}\sqrt{\Delta t}\epsilon_t, \quad \epsilon_t \sim \mathcal{N}(0, 1), \quad (4.31)$$

trong đó θ_{ou} là tốc độ hồi quy về trung bình μ_{ou} và σ_{ou} điều khiển mức dao động. OU noise tạo ra chuỗi nhiễu “mượt” giúp khám phá ổn định hơn so với nhiễu độc lập.

Cập nhật mềm mạng mục tiêu (Polyak Averaging). Để tránh dao động lớn khi huấn luyện critic, sử dụng cập nhật mềm tham số mạng mục tiêu:

$$\theta^{\text{target}} \leftarrow \tau \theta + (1 - \tau) \theta^{\text{target}}, \quad (4.32)$$

với $\tau \in (0, 1)$ nhỏ (ví dụ 0.001) để đảm bảo cập nhật chậm và ổn định.

Ý nghĩa của điều phối học tăng cường trong suy luận phản thực. Cơ chế điều phối học tăng cường cho phép mô hình học một chiến lược trộn theo cả thể thay vì dùng trọng số cố định. Nhờ đó, hệ thống có thể thích ứng tốt hơn trước sự không đồng nhất của dữ liệu y tế HIV và sự khác nhau về độ tin cậy giữa hai nhánh suy luận, từ đó nâng cao chất lượng dự báo phản thực và tính bền vững khi triển khai trên dữ liệu thực.

4.4.5. Ước lượng bất định và thích ứng với sự thay đổi phân phối

Trong các ứng dụng y tế, việc chỉ cung cấp ước lượng điểm là chưa đủ; hệ thống cần đánh giá độ tin cậy của các ước lượng và phát hiện sớm các thay đổi trong phân phối dữ liệu.

1. Ước lượng độ bất định (Uncertainty Estimation)

Trong bài toán suy luận nhân quả từ dữ liệu quan sát, việc ước lượng kết quả phản thực (counterfactual outcome) không chỉ yêu cầu độ chính xác cao mà còn cần đánh giá mức độ tin cậy của các ước lượng này. Đặc biệt trong bối cảnh phân tích tác động nhân quả của hành vi xét nghiệm sàng lọc HIV đến hành vi xét nghiệm cận lâm sàng, độ bất định của dự báo đóng vai trò quan trọng trong việc nâng cao tính minh bạch và độ tin cậy của hệ thống.

Để định lượng độ bất định trong dự báo phản thực và ước lượng hiệu ứng nhân quả, nghiên cứu áp dụng phương pháp Monte Carlo Dropout (MC Dropout) [106] như một xấp xỉ Bayesian của mạng nơ-ron học sâu. Trong giai đoạn suy luận, dropout được giữ nguyên trạng thái hoạt động và mô hình thực hiện T lần truyền xuôi ngẫu nhiên, thu được tập dự báo:

$$\{\hat{y}_1, \hat{y}_2, \dots, \hat{y}_T\}.$$

Trong đó, mỗi $\hat{y}_t$ biểu diễn xác suất nhiễm HIV được ước lượng dưới điều kiện phản thực tại lần truyền thứ t .

Giá trị trung bình dự báo được tính như sau:

$$\bar{y} = \frac{1}{T} \sum_{t=1}^T \hat{y}_t. \quad (4.33)$$

Phương sai dự báo tổng quát được xác định bởi:

$$\text{Var}(y) = \frac{1}{T} \sum_{t=1}^T \hat{y}_t^2 - \left(\frac{1}{T} \sum_{t=1}^T \hat{y}_t \right)^2. \quad (4.34)$$

Giá trị này phản ánh mức độ dao động của dự báo phản thực và được xem như độ bất định tổng thể của mô hình.

Nếu mô hình đồng thời được huấn luyện để xuất ra phương sai nhiễu nội tại σ_t^2 tại mỗi lần truyền xuôi, thì độ bất định dự báo có thể được phân rã thành hai thành phần:

Độ bất định ngẫu nhiên (Aleatoric uncertainty) được tính toán như sau:

$$\text{Aleatoric} = \frac{1}{T} \sum_{t=1}^T \sigma_t^2. \quad (4.35)$$

Thành phần này phản ánh nhiễu nội tại không thể loại bỏ của dữ liệu.

Độ bất định tri thức (Epistemic uncertainty) được tính toán như sau:

$$\text{Epistemic} = \frac{1}{T} \sum_{t=1}^T \mu_t^2 - \left(\frac{1}{T} \sum_{t=1}^T \mu_t \right)^2, \quad (4.36)$$

trong đó μ_t là giá trị trung bình dự báo tại lần truyền thứ t .

Do đó, phương sai dự báo tổng được biểu diễn như sau:

$$\text{Var}(y) = \text{Epistemic} + \text{Aleatoric}. \quad (4.37)$$

Cơ chế phân rã này cho phép mô hình đánh giá mức độ tin cậy của từng ước lượng phản thực, từ đó hỗ trợ ra quyết định chính sách trong các ứng dụng y tế công cộng.

2. Phát hiện dịch chuyển khái niệm (Concept Drift Detection)

Trong môi trường thực tế, phân phối dữ liệu có thể thay đổi theo thời gian do sự biến động hành vi xã hội, chính sách can thiệp hoặc thay đổi đặc điểm dịch tễ. Nếu không được phát hiện kịp thời, sự thay đổi này có thể làm suy giảm độ chính xác của mô hình kết quả và mô hình propensity.

Để nâng cao tính bền vững của hệ thống, nghiên cứu tích hợp cơ chế phát hiện concept drift dựa trên sự so sánh phân phối đặc trưng đầu vào giữa dữ liệu huấn luyện và dữ liệu hiện tại.

Ký hiệu μ_{train} là vector trung bình của dữ liệu huấn luyện và μ_{test} là vector trung bình của dữ liệu kiểm thử hiện tại. Drift được phát hiện khi:

$$\|\mu_{\text{train}} - \mu_{\text{test}}\|_2 > \delta, \quad (4.38)$$

trong đó $\|\cdot\|_2$ là chuẩn Euclid và δ là ngưỡng được xác định thực nghiệm.

Khi phát hiện drift, hệ thống có thể điều chỉnh trọng số ensemble, kích hoạt cơ chế cập nhật mô hình hoặc thay đổi chiến lược suy luận nhân quả. Cơ chế này đóng vai trò như một bước kiểm soát chất lượng ở tầng đầu vào, đặc biệt quan trọng trong các bài toán ước lượng hiệu ứng nhân quả khi phân phối giữa nhóm điều trị và nhóm đối chứng có thể thay đổi theo thời gian.

4.4.5.1. Quy trình huấn luyện mô hình

Dựa trên kiến trúc tổng thể của mô hình, quá trình huấn luyện được thực hiện theo ba thành phần chính, bao gồm học biểu diễn, ước lượng hiệu ứng nhân quả và điều phối động thông qua học tăng cường.

Trước hết, mô hình TITAN được huấn luyện để học biểu diễn tiềm ẩn và dự báo kết cục dựa trên dữ liệu đầu vào đã được lựa chọn từ đồ thị nhân quả. Song song, mô hình nhân quả DRLearner được sử dụng để ước lượng hiệu ứng nhân quả dựa trên nguyên lý bền vững kép.

Trên cơ sở hai nhánh dự báo này, một tác nhân học tăng cường được xây dựng để học chính sách kết hợp động giữa các đầu ra. Tại mỗi bước, tác nhân nhận trạng thái đầu vào là đặc trưng của mẫu và các dự báo trung gian, sau đó sinh ra trọng số kết hợp nhằm tối ưu hóa hiệu năng dự báo cuối cùng.

Hàm mục tiêu của tác nhân học tăng cường được thiết kế nhằm tối đa hóa chất lượng dự báo, đồng thời xem xét yếu tố bất định và sự thay đổi phân phối dữ liệu. Quá trình cập nhật chính sách được thực hiện thông qua các thành phần actor–critic, với việc tối ưu đồng thời hàm giá trị và hàm chính sách.

Toàn bộ hệ thống được huấn luyện lặp lại qua nhiều epoch, trong đó các tham số của TITAN, DRLearner và tác nhân học tăng cường được cập nhật theo cơ chế lan truyền ngược và tối ưu hóa tương ứng. Hiệu năng mô hình được theo dõi trên tập xác thực để lựa chọn cấu hình tối ưu và đảm bảo tính ổn định.

4.5. Thực nghiệm và kết quả

Nghiên cứu thực hiện đánh giá mô hình trong hai kịch bản bổ sung cho nhau.

Thứ nhất, hệ thống được kiểm chứng trên dữ liệu mô phỏng với hiệu ứng nhân quả thật (ground-truth causal effects) đã biết trước. Thiết lập này cho phép định lượng chính xác sai số ước lượng, kiểm chứng tính đúng đắn nhân quả và đánh giá độ tin cậy của cơ chế ước lượng bất định trong điều kiện có kiểm soát.

Thứ hai, hệ thống được áp dụng trên các bộ dữ liệu HIV thực tế. Trong bối cảnh này, hiệu ứng phản thực không quan sát được, do đó việc đánh giá tập trung vào hiệu năng dự báo, tính ổn định và tính nhất quán của ước lượng ở mức quần thể.

4.5.1. Nghiên cứu mô phỏng (Simulation Study)

Trước khi áp dụng mô hình đề xuất cho dữ liệu HIV thực tế, nghiên cứu này tiến hành một thí nghiệm mô phỏng nhằm đánh giá khả năng ước lượng nhân quả trong điều kiện mà hiệu quả can thiệp thực sự đã được biết trước. Mục tiêu của nghiên cứu mô phỏng là kiểm chứng tính đúng đắn về mặt nhân quả của phương pháp, đặc biệt trong

việc ước lượng hiệu quả can thiệp trung bình (ATE) và hiệu quả can thiệp cá thể hóa (ITE), dưới các điều kiện gây nhiễu, trung gian và dị biệt hiệu quả.

Việc sử dụng dữ liệu mô phỏng cho phép kiểm soát chặt chẽ cấu trúc nhân quả và cung cấp giá trị chuẩn (ground truth) để so sánh trực tiếp các ước lượng. Đây là một bước xác thực thuật toán cần thiết trước khi triển khai phương pháp trên dữ liệu HIV quan sát, nơi các quan hệ nhân quả không thể quan sát trực tiếp.

1. Quá trình sinh dữ liệu (Data Generating Process)

Nghiên cứu xây dựng một hệ thống nhân quả bán thực tế mô phỏng bối cảnh HIV, bao gồm các biến gây nhiễu quan sát được, một biến công cụ mang tính chính sách, một biến trung gian sau điều trị và hiệu quả can thiệp dị biệt theo độ tuổi. Trừ khi có ghi chú khác, số lượng mẫu được sinh là $n = 5000$ với hạt giống ngẫu nhiên bằng 42 để đảm bảo khả năng tái lập.

Các biến ngoại sinh X_0 đến X_9 được sinh như sau:

$$X_0 \sim \mathcal{N}(50, 10^2), \quad X_1 \sim \text{Bern}(0.5), \quad X_2, X_3, X_9 \sim \mathcal{N}(0, 1),$$

$$X_4 \sim \text{Bern}(0.2), \quad X_5 \sim \text{Unif}(0, 1), \quad X_6 \sim \text{Poisson}(2), \quad X_8 \sim \text{Bern}(0.6).$$

Một biến sức khỏe tiềm ẩn $X_7 \sim \mathcal{N}(0, 1)$ được sinh nhưng không được quan sát trong tập dữ liệu phân tích, nhằm mô phỏng yếu tố tiềm ẩn không đo được thường gặp trong dữ liệu HIV thực tế.

Một biến công cụ Z phụ thuộc vào yếu tố địa lý X_8 được sinh theo mô hình logistic như sau:

$$\mathbb{P}(Z = 1 \mid X) = \sigma(0.8X_8 - 0.5), \quad (4.39)$$

trong đó $\sigma(\cdot)$ là hàm sigmoid. Biến Z chỉ ảnh hưởng đến việc gán điều trị và không có đường tác động trực tiếp đến kết cục.

Việc gán điều trị T được mô hình hóa thông qua:

$$\text{logit } \mathbb{P}(T = 1 \mid X, Z) = 0.02(X_0 - 50) + 0.6X_2 + 0.5X_4 + 1.5Z + \varepsilon_T, \quad (4.40)$$

từ đó xác suất điều trị (propensity score) được sử dụng để sinh biến T .

Sau điều trị, một biến trung gian M (đóng vai trò tương tự tải lượng virus) được sinh theo công thức sau:

$$M = 2 - 1.5T + 0.5X_3 + \varepsilon_M, \quad (4.41)$$

trong đó điều trị làm giảm giá trị trung bình của M .

Hiệu quả can thiệp cá thể hóa được mô hình hóa phụ thuộc vào độ tuổi:

$$\tau_i = -1.5 + 0.08(X_{0i} - 50). \quad (4.42)$$

Cuối cùng, kết cục nhị phân Y được sinh theo mô hình logistic như sau:

$$\text{logit } \mathbb{P}(Y = 1) = 0.03(X_0 - 50) - 0.8X_2 + 0.7X_4 + 0.5M + 0.4X_7 + T \cdot \tau + \varepsilon_Y. \quad (4.43)$$

Hai bảng dữ liệu được tạo ra cho mỗi lần mô phỏng: (i) tập dữ liệu quan sát

$$D_{\text{obs}} = \{X_0, X_1, X_2, X_3, X_4, X_5, X_6, X_8, X_9, Z, M, T, Y\}, \quad (4.44)$$

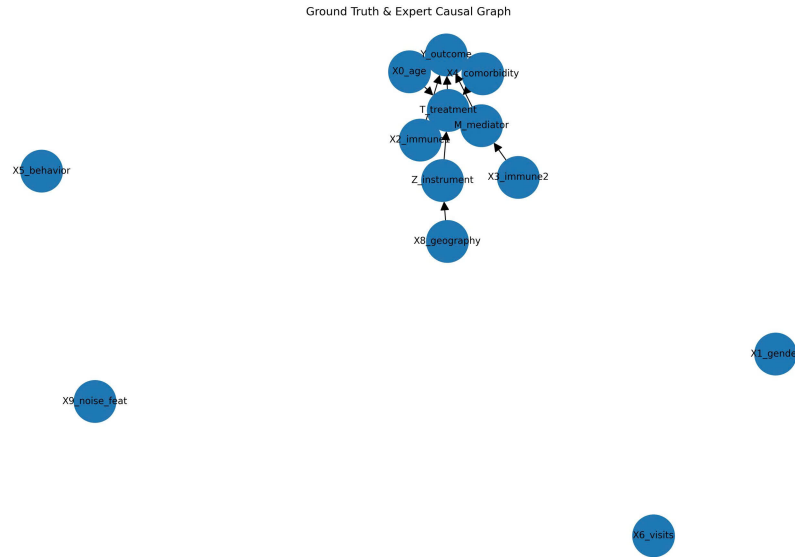
trong đó biến X_7 bị loại bỏ; và (ii) bảng chuẩn chứa ITE thật, xác suất điều trị và xác suất kết cục để phục vụ đánh giá.

2. Đồ thị nhân quả: Cấu trúc chuẩn và cấu trúc ước lượng

Hình 4.2 minh họa đồ thị nhân quả chuẩn được xác định theo cơ chế sinh dữ liệu, trong đó tuổi (X_0), tình trạng miễn dịch (X_2) và bệnh nền (X_4) đóng vai trò là các biến gây nhiễu. Yếu tố địa lý (X_8) chỉ tác động đến điều trị thông qua biến công cụ Z , và biến trung gian M nằm trên đường nhân quả từ điều trị đến kết cục. Đồ thị này được sử dụng làm cấu trúc tham chiếu để xác định các quan hệ nhân quả thực sự trong dữ liệu mô phỏng.

Hình 4.3 trình bày đồ thị nhân quả được ước lượng từ dữ liệu quan sát. So với đồ thị chuẩn, cấu trúc này dày đặc hơn, xuất hiện thêm một số cạnh phản ánh các quan hệ tương quan, trong khi một số đường nhân quả thực sự có thể bị suy giảm. Sự sai khác này phản ánh những hạn chế của quá trình khám phá nhân quả từ dữ liệu quan sát và làm nổi bật vai trò của các phương pháp ước lượng nhân quả bền vững.

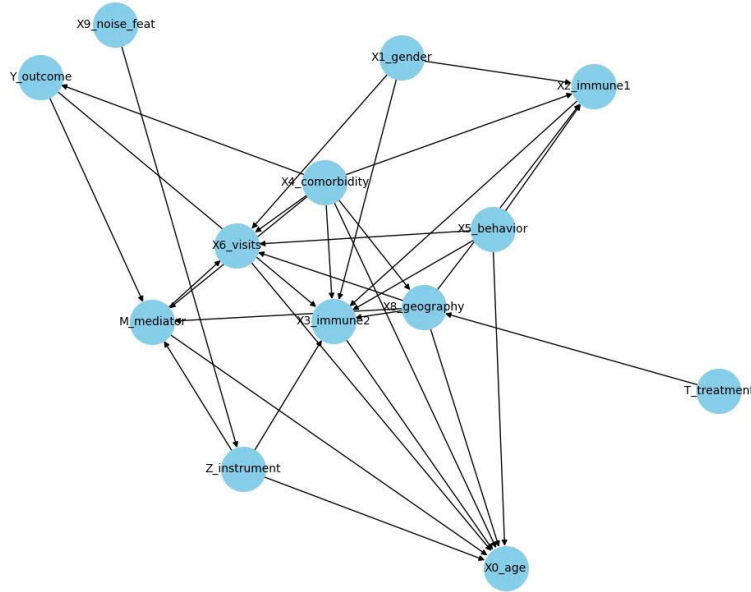
Việc đối chiếu hai đồ thị được thực hiện nhằm đánh giá mức độ nhất quán giữa cấu trúc nhân quả được ước lượng từ dữ liệu và cấu trúc chuẩn đã biết từ cơ chế sinh dữ liệu. Kết quả đối chiếu cũng được sử dụng làm cơ sở lựa chọn tập biến điều chỉnh cho các bước ước lượng hiệu ứng nhân quả tiếp theo.



Hình 4.2: Đồ thị nhân quả chuẩn (ground truth) cho dữ liệu mô phỏng

3. Kết quả nghiên cứu mô phỏng

Bảng 4.1 trình bày kết quả so sánh hiệu năng của mô hình đề xuất với các phương pháp suy luận nhân quả cơ sở trên dữ liệu mô phỏng. Phương pháp đề xuất đạt kết quả tốt nhất trên tất cả các chỉ số nhân quả, bao gồm sai số PEHE thấp nhất, sai số và độ chệch ATE nhỏ nhất, cùng với giá trị R^2 cao nhất trong ước lượng ITE. Giá trị ATE ước lượng cũng gần nhất với giá trị thật được sử dụng trong quá trình sinh dữ liệu.



Hình 4.3: Đồ thị nhân quả ước lượng từ dữ liệu mô phỏng

Bảng 4.1: So sánh hiệu năng các mô hình suy luận nhân quả trên dữ liệu mô phỏng

Phương pháp	PEHE ↓	Sai số ATE ↓	Độ lệch ATE	ITE R^2 ↑	ATE ước lượng
MH đề xuất	0.0980	0.0198	-0.0089	0.5834	0.2261
CAT Network	0.1089	0.0234	0.0187	0.5445	0.2537
Orthogonal Random Forests	0.1156	0.0267	-0.0145	0.5198	0.2205
Double/Debiased ML	0.1203	0.0289	0.0223	0.4987	0.2573
Causal Forest	0.1267	0.0321	-0.0198	0.4743	0.2152
X-Learner	0.1334	0.0367	0.0298	0.5456	0.2648
CEVAE	0.1423	0.0412	-0.0267	0.5123	0.2083

Bảng 4.2 trình bày phân tích loại bỏ thành phần, cho thấy mỗi mô-đun trong mô hình đề xuất đều đóng vai trò quan trọng. Đặc biệt, việc loại bỏ cơ chế RL ensemble gây suy giảm hiệu năng lớn nhất, cho thấy vai trò trung tâm của thành phần này trong việc ổn định và điều phối các bộ ước lượng.

Bảng 4.3 cho thấy mô hình đề xuất không chỉ vượt trội về suy luận nhân quả mà còn duy trì hiệu năng dự báo cao hơn các phương pháp so sánh.

Bảng 4.4 phân tích đặc tính phân phối của hiệu quả can thiệp, cho thấy mô hình đề xuất đạt độ bao phủ gần mức danh định 95% và sai lệch hiệu chuẩn thấp nhất.

4. Kiểm tra độ bền vững (Robustness Testing)

Các thực nghiệm kiểm tra độ bền vững cho thấy mô hình duy trì hiệu năng ổn định khi kích thước mẫu giảm, khi cường độ gây nhiễu tăng và khi mô hình kết cục bị chỉ định sai. Đặc biệt, nhờ đặc tính doubly robust, sai lệch ATE vẫn được kiểm soát ở mức thấp ngay cả trong các kịch bản bất lợi, trong khi các phương pháp đơn bền vững suy giảm đáng kể.

Những kết quả này xác nhận tính ổn định và độ tin cậy của mô hình đề xuất, tạo

Bảng 4.2: Phân tích loại bỏ thành phần của mô hình trên dữ liệu mô phỏng

Cấu hình mô hình	PEHE ↓	Sai số ATE ↓	ITE R^2 ↑	ITE MAE ↓
MH đề xuất	0.0980	0.0198	0.5834	0.0756
Không có RL Ensemble	0.1134	0.0243	0.5456	0.0789
Không có TITAN	0.1189	0.0278	0.5223	0.0823
Không có DR Learner	0.1076	0.0198	0.5734	0.0734
Không có đồ thị nhân quả	0.1156	0.0267	0.5389	0.0667

Bảng 4.3: So sánh hiệu năng dự báo trên dữ liệu mô phỏng

Phương pháp	Accuracy	F1-score	AUC-ROC
MH đề xuất	0.8734	0.862	0.924
CAT Network	0.8267	0.814	0.889
Orthogonal Random Forests	0.8189	0.807	0.883
Double/Debiased ML	0.8123	0.798	0.876
Causal Forest	0.8045	0.789	0.867
X-Learner	0.7967	0.776	0.854
CEVAE	0.7834	0.761	0.841

Bảng 4.4: Phân tích phân phối hiệu quả can thiệp trên dữ liệu mô phỏng

Phương pháp	Độ lệch chuẩn ITE ↓	Độ bao phủ (95% CI)	Sai lệch hiệu chuẩn ↓
MH đề xuất	0.1456	94.2%	0.0234
CAT Network	0.1678	91.8%	0.0367
Orthogonal Random Forests	0.1723	90.6%	0.0389
Double/Debiased ML	0.1789	89.4%	0.0412
Causal Forest	0.1834	88.7%	0.0445
X-Learner	0.1923	87.3%	0.0478
CEVAE	0.2067	85.9%	0.0523

tiền đề vững chắc cho việc áp dụng trên dữ liệu HIV thực tế.

4.5.2. Đánh giá thực nghiệm trên các bộ dữ liệu HIV

4.5.2.1. Mô tả dữ liệu

Trong nghiên cứu này, hai bộ dữ liệu liên quan đến HIV được sử dụng để đánh giá mô hình đề xuất. Cả hai bộ dữ liệu đều được công bố công khai trên nền tảng Kaggle và đã được sử dụng trong một số nghiên cứu học máy trước đây.

Bộ dữ liệu thứ nhất là EDHS-HIV/AIDS dataset, được cung cấp trên Kaggle [37] và ban đầu được tổng hợp từ Khảo sát Nhân khẩu học và Sức khỏe Ethiopia (Ethiopian Demographic and Health Survey – EDHS) do Cơ quan Thống kê Trung ương Ethiopia (CSA) phối hợp với ICF International thực hiện. Bộ dữ liệu này đã được sử dụng trong các nghiên cứu khoa học trước đây [107, 108]. Bộ dữ liệu bao gồm 78.877 bản ghi ẩn danh, phản ánh các yếu tố hành vi, nhân khẩu học và kiến thức liên quan đến xét nghiệm HIV trong cộng đồng.

Trong bộ dữ liệu này, biến can thiệp được xác định là S_Test (0 = Không, 1 = Có), biểu thị việc một cá thể có thực hiện xét nghiệm mẫu hay không. Biến kết quả là T_in_LAB (0 = Không, 1 = Có), cho biết việc xét nghiệm khẳng định trong phòng thí nghiệm có được thực hiện hay không. Các biến đầu vào được chia thành ba nhóm chính.

Nhóm thứ nhất là các biến nhân khẩu học và kinh tế – xã hội, bao gồm giới tính, tuổi, vùng địa lý, loại khu vực cư trú, tôn giáo, trình độ học vấn, tình trạng hôn nhân, tình trạng việc làm và chỉ số tài sản. Nhóm thứ hai là các chỉ báo hành vi tình dục, bao gồm số lượng bạn tình, việc sử dụng bao cao su và các hành vi thay đổi nhằm giảm nguy cơ lây nhiễm HIV. Nhóm thứ ba là kiến thức liên quan đến HIV, bao gồm nhận thức về con đường lây truyền, hiểu biết về các bệnh lây truyền qua đường tình dục, kiến thức về HIV/AIDS và khả năng tiếp cận các dịch vụ xét nghiệm.

Bộ dữ liệu thứ hai là AIDS Virus Infection Prediction dataset, cũng được cung cấp trên Kaggle [38]. Bộ dữ liệu này được báo cáo là dựa trên dữ liệu từ nghiên cứu lâm sàng AIDS Clinical Trials Group (ACTG) Study 175 [109]. Phiên bản được công bố trên Kaggle là một phiên bản đã được xử lý và mở rộng, bao gồm khoảng 50.000 bản ghi ẩn danh, được phát hành theo giấy phép CC0. Bộ dữ liệu này chứa thông tin liên quan đến các thử nghiệm lâm sàng trong điều trị HIV.

Trong bộ dữ liệu này, biến kết quả là infected, cho biết một bệnh nhân có nhiễm AIDS hay không (1 = Có, 0 = Không). Biến điều trị là trt, biểu thị nhóm điều trị với bốn giá trị khác nhau: 0 = điều trị đơn trị ZDV, 1 = kết hợp ZDV + ddI, 2 = kết hợp ZDV + Zai, 3 = điều trị đơn trị ddI.

Các biến đầu vào còn lại được chia thành bốn nhóm chính. Nhóm thứ nhất là các biến nhân khẩu học và hành vi, bao gồm tuổi, giới tính, chủng tộc, tiền sử sử dụng ma túy và tình trạng quan hệ đồng giới. Nhóm thứ hai là các biến lâm sàng, bao gồm cân nặng, nồng độ hemoglobin, điểm Karnofsky, sự hiện diện của triệu chứng và tiền sử nhiễm trùng cơ hội. Nhóm thứ ba là các chỉ dấu miễn dịch, bao gồm số lượng tế bào CD4, CD8 và tỷ lệ CD4/CD8. Nhóm cuối cùng là các biến liên quan đến điều trị, bao

gồm loại điều trị và các chỉ báo về phác đồ điều trị.

Các đồ thị nhân quả được xây dựng từ hai bộ dữ liệu này, bao gồm cả các đồ thị dựa trên kiến thức chuyên gia và các đồ thị được suy ra từ dữ liệu, sẽ được mô tả chi tiết trong các phần tiếp theo.

4.5.2.2. Tiền xử lý dữ liệu

Hai bộ dữ liệu được sử dụng trong nghiên cứu bao gồm nhiều loại biến khác nhau, bao gồm biến phân loại, biến liên tục, dữ liệu thiếu và sự mất cân bằng giữa các nhóm điều trị. Do đó, một quy trình tiền xử lý dữ liệu có cấu trúc đã được áp dụng trước khi huấn luyện mô hình.

Trước hết, các biến phân loại được chuyển đổi sang dạng số bằng phương pháp label encoding. Việc mã hóa này cho phép các mô hình học máy và học sâu xử lý hiệu quả các biến phân loại trong quá trình huấn luyện.

Đối với các biến liên tục, nghiên cứu áp dụng phương pháp chuẩn hóa robust scaling, trong đó các giá trị được chuẩn hóa dựa trên trung vị và khoảng tứ phân vị (interquartile range). Phương pháp này giúp giảm ảnh hưởng của các giá trị ngoại lai, đồng thời vẫn giữ được ý nghĩa thực tế của các biến lâm sàng.

Khoảng 8,3% dữ liệu trong hai bộ dữ liệu bị thiếu. Đối với các biến liên tục, các giá trị thiếu được thay thế bằng giá trị trung vị của biến tương ứng. Đối với các biến phân loại, các giá trị thiếu được thay thế bằng giá trị xuất hiện nhiều nhất trong dữ liệu. Cách tiếp cận này phản ánh các mẫu thiếu dữ liệu khác nhau thường gặp trong dữ liệu khảo sát cộng đồng và dữ liệu lâm sàng.

Ngoài ra, dữ liệu cũng thể hiện sự mất cân bằng giữa các nhóm điều trị. Thay vì sử dụng các phương pháp sinh dữ liệu tổng hợp như SMOTE [110], vốn có thể tạo ra các hồ sơ bệnh nhân không thực tế về mặt lâm sàng, nghiên cứu áp dụng một chiến lược gán nhãn lại dựa trên phân bố quần thể nhằm giảm mức độ mất cân bằng mà vẫn bảo toàn cấu trúc thực của dữ liệu.

Sau khi hoàn thành các bước tiền xử lý, các bộ dữ liệu được chia ngẫu nhiên thành tập huấn luyện và tập kiểm tra để đánh giá hiệu quả của mô hình. Quy trình tiền xử lý này giúp đảm bảo tính ổn định của mô hình, đồng thời cung cấp nền tảng dữ liệu đáng tin cậy cho các bước phân tích nhân quả và học sâu trong các phần tiếp theo của nghiên cứu.

4.5.2.3. Các giả định trong suy luận nhân quả

Để xác định hiệu ứng nhân quả từ dữ liệu quan sát, nghiên cứu dựa trên các giả định chuẩn của khung kết quả tiềm năng (*potential outcome framework*).

Tính nhất quán (Consistency). Nếu đối tượng i thực sự nhận mức can thiệp t , thì kết quả quan sát được chính là kết quả tiềm năng dưới mức can thiệp đó, được mô tả theo công thức sau:

$$Y_i = Y_i(t), \quad \text{khi } T_i = t \quad (4.45)$$

Tính chồng lấn (Positivity/Overlap). Với mỗi cấu hình biến đặc trưng X_i , mỗi mức

can thiệp đều phải có xác suất xuất hiện khác 0, được mô tả theo công thức sau:

$$0 < P(T = t | X_i) < 1, \quad \forall t \in \mathcal{T} \quad (4.46)$$

Giả định này bảo đảm mô hình có đủ dữ liệu để so sánh hiệu ứng giữa các nhóm can thiệp khác nhau.

Tính không gây nhiễu có điều kiện (Ignorability/Unconfoundedness). Khi đã điều kiện hóa trên các biến quan sát được X_i , việc phân bổ can thiệp độc lập với các kết quả tiềm năng được mô tả theo công thức sau:

$$Y_i(t) \perp\!\!\!\perp T_i | X_i, \quad \forall t \in \mathcal{T} \quad (4.47)$$

Điều này có nghĩa là các biến gây nhiễu quan trọng ảnh hưởng đồng thời đến việc phân bổ can thiệp và kết cục đã được quan sát và đưa vào mô hình.

Giả định SUTVA (Stable Unit Treatment Value Assumption). Kết quả của một đối tượng không bị ảnh hưởng bởi mức can thiệp áp dụng cho đối tượng khác, đồng thời mỗi mức can thiệp được xác định rõ ràng và nhất quán.

Đối với bộ dữ liệu thứ nhất (EDHS-HIV/AIDS), biến can thiệp là S Test (0 = Không, 1 = Có), và biến kết cục là T in LAB (0 = Không, 1 = Có). Trong bối cảnh này, giả định consistency có nghĩa là nếu một cá nhân thực sự tham gia sàng lọc cộng đồng, thì kết quả xác nhận xét nghiệm trong phòng thí nghiệm quan sát được chính là kết quả tiềm năng tương ứng với trạng thái đó. Giả định ignorability ngụ ý rằng sau khi đã điều chỉnh theo các biến nhân khẩu học, hành vi và mức độ hiểu biết liên quan, quyết định tham gia sàng lọc là độc lập với kết cục xét nghiệm tiềm năng. Giả định positivity yêu cầu trong mỗi nhóm đặc trưng quan sát được đều tồn tại cả cá nhân có và không tham gia sàng lọc.

Đối với bộ dữ liệu thứ hai (AIDS Virus Infection Prediction), biến can thiệp là trt (nhóm điều trị: 0 = ZDV, 1 = ZDV+ddI, 2 = ZDV+Zal, 3 = ddI), và biến kết cục là infected (0 = Không, 1 = Có). Trong trường hợp này, giả định consistency có nghĩa là nếu bệnh nhân thực sự nhận một phác đồ điều trị cụ thể thì kết quả nhiễm quan sát được phản ánh đúng kết quả tiềm năng dưới phác đồ đó. Giả định ignorability yêu cầu rằng sau khi đã điều chỉnh theo các biến nhân khẩu học, hành vi, lâm sàng và miễn dịch học, việc phân bổ điều trị là độc lập với kết cục nhiễm tiềm năng. Giả định positivity yêu cầu sự hiện diện của các nhóm điều trị khác nhau trong các phân tầng đặc trưng bệnh nhân.

Trong thực tế, tính hợp lý của các giả định này đã được đánh giá trên cả hai bộ dữ liệu thông qua các kiểm tra thực nghiệm, bao gồm đánh giá mức độ cân bằng giữa các nhóm điều trị và mức độ chồng lấn của phân bố điều trị theo các biến quan sát được. Kết quả không cho thấy dấu hiệu vi phạm nghiêm trọng đối với giả định positivity hoặc sự mất cân bằng lớn ở các biến gây nhiễu quan sát được, cho thấy các giả định này là hợp lý trong bối cảnh nghiên cứu. Tuy nhiên, đối với dữ liệu quan sát, giả định unconfoundedness không thể được kiểm chứng tuyệt đối chỉ dựa trên dữ liệu quan sát, do vẫn có khả năng tồn tại các yếu tố gây nhiễu chưa được đo lường. Nếu giả định này bị vi phạm, các ước lượng hiệu ứng điều trị có thể bị sai lệch và ảnh hưởng đến

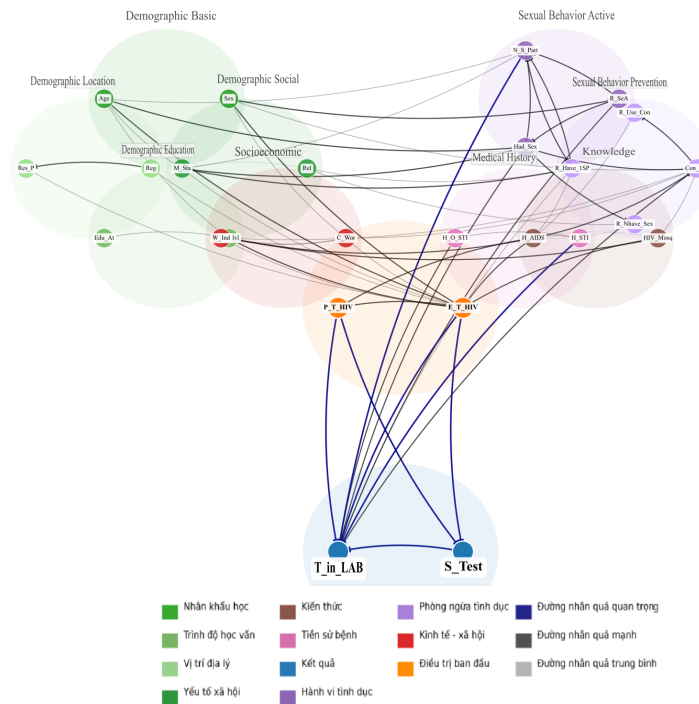
độ tin cậy của kết quả suy luận nhân quả.

4.5.2.4. Xây dựng và lựa chọn đồ thị nhân quả cho hai bộ dữ liệu.

Trong nghiên cứu này, cấu trúc nhân quả được xây dựng riêng biệt cho bộ dữ liệu 1 [37] và bộ dữ liệu 2 [38] theo quy trình hai bước, kết hợp giữa kiến thức chuyên gia và khám phá cấu trúc từ dữ liệu.

Bộ dữ liệu 1

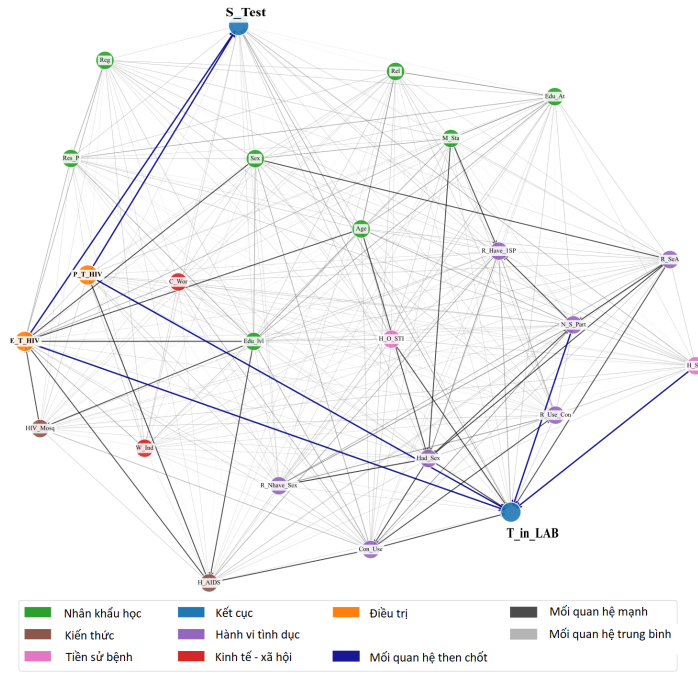
1. *DAG được chuyên gia xác định.* Hình 4.4 minh họa DAG được xây dựng dựa trên kiến thức dịch tễ học HIV. Trong cấu trúc này, các yếu tố nhân khẩu học như *Sex* và *Age* ảnh hưởng đến *Region (Reg)* và *Current Work Status (C_Wor)*, từ đó tác động đến hành vi xét nghiệm HIV. Bối cảnh sống và giáo dục (*Res_P*, *Edu_At*) định hình nhận thức và khả năng tiếp cận dịch vụ. Các yếu tố hôn nhân và kinh tế (*M_STI*, *W_Ind*) ảnh hưởng đến hành vi tình dục và nhận thức phòng ngừa. Các yếu tố sức khỏe như triệu chứng STI (*H_STI*) và hiểu biết về AIDS (*H_AIDS*) liên quan trực tiếp đến quyết định xét nghiệm HIV. Hành vi tình dục (*Had_Sex*, *Con_Use*, *R_Have_ISP*) ảnh hưởng mạnh đến biến *Ever Tested HIV (E_T_HIV)*, sau đó tác động đến *T_HIV_LAB* và *S_Test*.



Hình 4.4: Đồ thị nhân quả do chuyên gia xác định từ bộ dữ liệu 1

2. *DAG suy diễn từ dữ liệu.* Hình 4.5 trình bày đồ thị nhân quả được suy diễn bằng khung khám phá nhân quả REX [100]. Khác với đồ thị nhân quả do chuyên gia xác định, cấu trúc này được học trực tiếp từ dữ liệu. Các biến trung tâm như *E_T_HIV*, *P_T_HIV*, và *S_Test* có nhiều cạnh vào và ra, phản ánh vai trò trung tâm trong mạng phụ thuộc. Một số đường nhân quả mạnh từ hành vi và nhận thức (ví dụ *Con_Use*, *N_S_Part*) đến hành vi xét nghiệm được phát hiện.

3. *So sánh và lựa chọn DAG cuối cùng.* Hai cấu trúc được thực hiện so sánh dựa



Hình 4.5: Đồ thị nhân quả suy diễn từ dữ liệu của bộ dữ liệu 1.

trên các tiêu chí như sau:

1. Tính hợp lý dịch tễ học,
2. Tính ổn định của ước lượng hiệu ứng nhân quả khi sử dụng mỗi đồ thị nhân quả,
3. phân tích độ nhạy của hiệu ứng can thiệp trung bình (ATE) và hiệu ứng can thiệp cá thể (ITE).

DAG do chuyên gia xác định được sử dụng làm cấu trúc tham chiếu trong phân tích độ nhạy. Các cạnh không hợp lý sinh học trong đồ thị nhân quả suy diễn từ dữ liệu được loại bỏ, trong khi các mối liên hệ hợp lý nhưng không có trong DAG do chuyên gia xác định được giữ lại.

DAG cuối cùng cho bộ dữ liệu 1, ký hiệu là $G_{\text{final}}^{\text{dataset1}}$, là phiên bản hiệu chỉnh của đồ thị nhân quả suy diễn từ dữ liệu sau khi được kiểm định bằng kiến thức chuyên gia.

Từ $G_{\text{final}}^{\text{dataset1}}$, nghiên cứu xác định:

biến gây nhiễu (Confounders): Sex, Age, Edu_lvl, Edu_At, M_Sta, C_Wor, W_Ind, H_STI, H_O_STI, H_AIDS, Reg.

biến trung gian (Mediators): R_SeA, Had_Sex, Con_Use, R_Have_1SP.

biến công cụ (Instrumental variables): Rel, Res_P.

Bộ dữ liệu 2

1. *DAG dựa trên chuyên gia xác định.* Hình 4.6 thể hiện đồ thị nhân quả được xây dựng dựa trên kiến thức lâm sàng. Biến *infected* là kết cục trung tâm, chịu ảnh hưởng của các yếu tố hành vi (homo), chỉ số miễn dịch (cd420, cd820), triệu chứng (symptom), lý do thăm khám (z30), và gián đoạn điều trị (offtrt). Các đặc điểm nền như tuổi, giới, chủng tộc, tình trạng chức năng (karnof), điều trị trước đó (oprior, preanti) và thiếu máu (hemo) ảnh hưởng đến nguy cơ nhiễm HIV thông qua các cơ chế trung gian.

2. *DAG suy diễn từ dữ liệu.* Hình 4.7 minh họa DAG được suy diễn bằng phương

$G_{\text{final}}^{\text{dataset2}}$ được sử dụng để xác định tập biến điều chỉnh trong toàn bộ quy trình ước lượng nhân quả của luận án.

4.5.2.5. Chiến lược chia dữ liệu

Đối với cả hai bộ dữ liệu, nghiên cứu chia dữ liệu thành ba phần: 64% cho tập huấn luyện, 16% cho tập xác thực, và 20% cho tập kiểm tra. Tập xác thực được sử dụng để điều chỉnh siêu tham số và lựa chọn mô hình phù hợp, trong khi các kết quả cuối cùng được báo cáo trên tập kiểm tra độc lập. Bên cạnh đó, nghiên cứu cũng áp dụng kỹ thuật xác thực chéo (k-fold cross-validation) trong quá trình huấn luyện và xác thực để đảm bảo rằng kết quả thu được không phụ thuộc vào một cách chia dữ liệu cụ thể.

4.5.2.6. Chỉ số đánh giá

Để đánh giá hiệu năng dự đoán kết cục của mô hình, nghiên cứu sử dụng các chỉ số phân loại phổ biến bao gồm Accuracy, F1-score và AUC-ROC nhằm đánh giá toàn diện khả năng phân biệt kết cục của mô hình. Accuracy được sử dụng để đo lường tỷ lệ dự đoán đúng tổng thể của mô hình đối với cả hai lớp kết cục. F1-score phản ánh sự cân bằng giữa độ chính xác (precision) và độ bao phủ (recall), đặc biệt phù hợp trong các trường hợp dữ liệu mất cân bằng. AUC-ROC được sử dụng để đánh giá khả năng phân biệt giữa các lớp kết cục trên nhiều ngưỡng quyết định khác nhau, qua đó phản ánh năng lực phân loại tổng quát của mô hình.

4.5.2.7. Các mô hình so sánh

Nghiên cứu tiến hành so sánh với các phương pháp nền tiêu biểu đại diện cho nhiều hướng tiếp cận khác nhau. Các phương pháp này bao gồm nhóm meta-learner và ước lượng hiệu quả điều trị, nhóm mô hình dựa trên cây, nhóm học biểu diễn sâu và nhóm kiến trúc dựa trên Transformer. Cụ thể, các phương pháp so sánh gồm: *Double/Debiased Machine Learning (DML)* [111], *Orthogonal Random Forests (ORF)* [112], *Causal Forests* [39], *X-Learner* [113, 114], *CEVAE (Causal Effect Variational Autoencoder)* [115], và *Causality-Aware Transformer (CAT)* [116].

Việc lựa chọn các baseline này nhằm kiểm tra khả năng của mô hình đề xuất trong việc kết hợp học biểu diễn, mô hình hóa kết cục và ước lượng tác động can thiệp. Trong đó, các chỉ số Accuracy, F1-score và AUC-ROC được sử dụng để đánh giá thành phần dự đoán kết cục quan sát được, còn phân tích ATE được sử dụng để diễn giải tác động can thiệp ước lượng ở mức quần thể.

4.5.2.8. Thiết lập thực nghiệm.

Nghiên cứu sử dụng Optuna như một khung tối ưu siêu tham số hiện đại để tự động hiệu chỉnh các tham số quan trọng của mô hình. Quá trình tối ưu được dẫn dắt bởi hiệu năng trên tập xác thực và tích hợp cơ chế dừng sớm (early stopping) nhằm hạn chế quá khớp và giảm thời gian huấn luyện không cần thiết. Chiến lược này cho phép lựa chọn cấu hình mô hình hiệu quả đồng thời đảm bảo khả năng tổng quát hóa trên các bộ dữ liệu khác nhau. Bộ siêu tham số tối ưu thu được được tóm tắt trong Bảng 4.5. Để đảm bảo tính công bằng, các baseline cũng được áp dụng chiến lược tối ưu tương tự.

Bảng 4.5: Các siêu tham số tối ưu được lựa chọn bởi Optuna.

Tham số	Giá trị
Batch size	1024
TITAN hidden size	512
TITAN number of layers	8
TITAN dropout	0.1567
TITAN learning rate	1.75e-4
TITAN number of heads	8
MLP hidden size	512
MLP dropout	0.1029
MLP learning rate	1.59e-4
MLP activation	SiLU
RL γ	0.9193
RL τ	0.0139
RL actor learning rate	2.26e-5
RL critic learning rate	5.01e-4
RL hidden dimension	128
RL update frequency	2
TITAN epochs	31
MLP epochs	29
RL epochs	30
Patience	4

4.5.2.9. Các kết quả chính

1. Các kết quả chính.

Bảng 4.6 trình bày hiệu năng của mô hình trên bài toán mô hình hóa kết cục quan sát được, được đánh giá bằng Accuracy, F1-score và AUC-ROC. Các chỉ số này được sử dụng nhằm đánh giá chất lượng của thành phần dự báo kết cục trong mô hình đề xuất, vì khả năng mô hình hóa chính xác kết cục quan sát là điều kiện quan trọng để hỗ trợ ước lượng hiệu quả can thiệp trong suy luận nhân quả. Trên bộ dữ liệu 1, CAT và ORF là hai baseline tốt nhất (Accuracy lần lượt 0.784 và 0.781; AUC-ROC 0.838 và 0.834). Tuy nhiên, mô hình đề xuất vượt trội trên cả ba chỉ số với Accuracy 0.861, F1-score 0.845 và AUC-ROC 0.897. Trên bộ dữ liệu, xu hướng vẫn nhất quán: CAT và ORF tiếp tục dẫn đầu trong nhóm baseline, nhưng vẫn thấp hơn mô hình đề xuất; Mô hình đề xuất đạt Accuracy 0.855, F1-score 0.839 và AUC-ROC 0.892. Kết quả cho thấy mô hình đề xuất cải thiện khả năng mô hình hóa kết cục quan sát so với các phương pháp đối chứng, từ đó tạo nền tảng tốt hơn cho bước ước lượng tác động can thiệp.

Để làm rõ đóng góp của từng thành phần trong kiến trúc, Bảng 4.7 trình bày kết quả của sáu biến thể mô hình được xây dựng từ các tổ hợp khác nhau của TITAN, DRLEARNER, MLP và RL. Kết quả cho thấy mô hình đầy đủ (TITAN + DRLEARNER + RL) đạt hiệu năng cao nhất trên cả hai bộ dữ liệu và nhất quán trên cả ba chỉ số đánh giá.

Bảng 4.6: So sánh mô hình đề xuất với các phương pháp suy luận nhân quả hiện có

Method	Bộ dữ liệu 1			Bộ dữ liệu 2		
	Accuracy	F1-Score	AUC-ROC	Accuracy	F1-Score	AUC-ROC
NH đề xuất	0.861	0.845	0.897	0.855	0.839	0.892
Double/Debiased Machine Learning	0.774	0.749	0.830	0.768	0.744	0.827
Orthogonal Random Forests	0.781	0.758	0.834	0.774	0.753	0.831
CEVAE	0.743	0.720	0.798	0.740	0.753	0.795
Causality-Aware Transformer Networks (CAT)	0.784	0.763	0.838	0.778	0.758	0.835
Causal Forest	0.768	0.747	0.824	0.763	0.742	0.821
X-Learner	0.756	0.738	0.812	0.751	0.733	0.808

Cụ thể, trên Bộ dữ liệu 1, mô hình đề xuất đạt Accuracy 0.861, F1-score 0.845 và AUC-ROC 0.897; trên Bộ dữ liệu 2, các giá trị tương ứng là 0.855, 0.839 và 0.892. So với các thành phần đơn lẻ, TITAN đạt hiệu năng cao hơn DRLearner, cho thấy khả năng học biểu diễn sâu đóng vai trò quan trọng trong việc mô hình hóa các quan hệ phi tuyến phức tạp trong dữ liệu HIV. Tuy nhiên, TITAN đơn thuần chưa tích hợp cơ chế điều chỉnh nhân quả rõ ràng, trong khi DRLearner dù có nền tảng suy luận nhân quả bền vững lại phụ thuộc mạnh vào chất lượng của các mô hình phụ trợ, dẫn đến hiệu năng tổng thể thấp hơn khi hoạt động độc lập.

Biến thể MLP + Causal cho kết quả thấp nhất trên cả hai bộ dữ liệu, cho thấy việc chỉ sử dụng một bộ học biểu diễn đơn giản là chưa đủ để nắm bắt cấu trúc phức tạp của dữ liệu thực tế, ngay cả khi có tích hợp thành phần nhân quả. Đáng chú ý, biến thể TITAN + MLP cho hiệu năng thấp hơn TITAN đơn lẻ, cho thấy việc kết hợp tĩnh giữa các thành phần dự báo không tự động mang lại cải thiện và thậm chí có thể gây nhiễu trong quá trình học.

Khi bổ sung cơ chế RL vào TITAN + MLP, hiệu năng được cải thiện rõ rệt trên cả hai bộ dữ liệu, phản ánh vai trò của cơ chế điều phối thích ứng trong việc lựa chọn và kết hợp động các nguồn ước lượng. Cuối cùng, khi kết hợp đầy đủ học biểu diễn sâu (TITAN), ước lượng nhân quả bền vững (DRLearner) và điều phối thích ứng (RL), mô hình đạt hiệu năng tốt nhất, cho thấy hiệu quả bổ sung lẫn nhau giữa các thành phần trong framework đề xuất.

Bảng 4.7: Hiệu năng các mô hình tổ hợp sử dụng TITAN, DRLearner, MLP và RL trên bộ dữ liệu 1 và bộ dữ liệu 2

Model	Bộ dữ liệu 1			Bộ dữ liệu 2		
	Accuracy	F1-Score	ROC-AUC	Accuracy	F1-Score	ROC-AUC
MH đề xuất	0.861	0.845	0.897	0.855	0.839	0.892
TITAN	0.789	0.767	0.845	0.788	0.765	0.841
DRLearner	0.762	0.729	0.818	0.755	0.722	0.812
MLP + Causal	0.705	0.662	0.729	0.699	0.654	0.723
TITAN + MLP	0.749	0.722	0.799	0.740	0.713	0.793
TITAN + MLP + RL	0.796	0.778	0.843	0.790	0.772	0.837

2. Ước lượng hiệu quả can thiệp trung bình (ATE).

Mặc dù các chỉ số dự báo (Accuracy/F1/AUC) cung cấp góc nhìn về khả năng phân tách lớp, nhưng các chỉ số này không phản ánh trực tiếp khía cạnh suy luận nhân quả. Do đó, nghiên cứu tiếp tục xem xét hiệu quả can thiệp trung bình (Average Treatment Effect - ATE) như một chỉ số bổ sung nhằm hỗ trợ đánh giá tác động can thiệp ở mức quần thể. Trên bộ dữ liệu HIV thực tế thứ nhất (Dataset 1), mô hình đề xuất ước lượng ATE là 0.247; trong khi trên bộ dữ liệu thứ hai (Dataset 2), giá trị này là 0.243. Kết quả này cho thấy mô hình không chỉ hỗ trợ dự báo kết cục mà còn có khả năng định lượng tác động can thiệp, bổ sung khía cạnh hỗ trợ ra quyết định trong bối cảnh HIV.

Bảng 4.8 trình bày giá trị hiệu ứng điều trị trung bình (ATE) được ước lượng bởi các mô hình khác nhau trên hai bộ dữ liệu. Trong khi các chỉ số Accuracy, F1-score và AUC-ROC phản ánh hiệu năng dự báo kết cục, ATE cung cấp góc nhìn bổ sung về mức độ khác biệt giữa các kết quả đầu ra được ước lượng trong hai kịch bản có và không có can thiệp. Tất cả các mô hình đều ước lượng ATE trong khoảng 0.227–0.247 và nhất quán giữa hai bộ dữ liệu, cho thấy sự đồng thuận về mức độ hiệu quả can thiệp trong quần thể HIV. Điều này củng cố độ tin cậy của phát hiện nhân quả, không phụ thuộc vào lựa chọn mô hình. Trong bối cảnh đó, ưu điểm của mô hình đề xuất nằm ở hiệu năng dự báo vượt trội (Accuracy, F1, AUC) khi vẫn duy trì ước lượng ATE phù hợp với các phương pháp nhân quả đã được thiết lập.

Bảng 4.8: So sánh ATE giữa các mô hình

Method	Bộ dữ liệu 1 (ATE)	Bộ dữ liệu 2 (ATE)
MH đề xuất	0.247	0.243
TITAN	0.235	0.231
DRLearner	0.230	0.227
Causal Forest	0.231	0.227
Causality-Aware Transformer (CAT) Networks	0.236	0.233
Orthogonal Random Forests	0.232	0.229
Double/Debiased ML	0.233	0.231

3. Phân tích loại bỏ thành phần (Ablation Study).

Để định lượng đóng góp của từng thành phần trong kiến trúc, nghiên cứu tiến hành ablation study như trình bày trong Bảng 4.9. Mô hình đầy đủ bao gồm TITAN, DRLearner, RL-ensemble, ước lượng bất định và phát hiện sự thay đổi quy luật dữ liệu theo thời gian (*concept drift*). Kết quả cho thấy mô hình đầy đủ đạt hiệu năng dự báo cao nhất trên cả hai bộ dữ liệu (Bộ dữ liệu 1: Accuracy 0.861, F1 0.845; Bộ dữ liệu 2: Accuracy 0.855, F1 0.839), đồng thời ước lượng ATE lần lượt là 0.247 và 0.243. Khi loại bỏ RL-ensemble hoặc TITAN, hiệu năng dự báo suy giảm rõ rệt, phản ánh vai trò then chốt của cơ chế điều phối thích ứng và học biểu diễn sâu. Việc loại bỏ đồ thị nhân quả, ước lượng bất định hoặc phát hiện trôi khái niệm cũng làm suy giảm hiệu năng, cho thấy các mô-đun này góp phần quan trọng vào tính bền vững và ổn định khi áp dụng trên dữ liệu HIV quan sát.

Phân tích bất định (Uncertainty Analysis). Trong các ứng dụng y tế, mức độ tin cậy của ước lượng là yếu tố quan trọng để hỗ trợ ra quyết định. Do đó, nghiên cứu đánh giá bất định dự đoán thông qua Monte Carlo Dropout. Bảng 4.10 cho thấy mô hình

Bảng 4.9: Ablation study: Ảnh hưởng hiệu năng khi loại bỏ từng thành phần hệ thống trên hai bộ dữ liệu

Kỹ thuật	Bộ dữ liệu 1			Bộ dữ liệu 2		
	Accuracy	F1	ATE	Accuracy	F1	ATE
MH đề xuất	0.861	0.845	0.247	0.855	0.839	0.243
Without RL Ensemble	0.789	0.767	0.235	0.788	0.765	0.231
Without TITAN	0.762	0.729	0.230	0.755	0.722	0.227
Without DRLEARNER	0.779	0.753	0.232	0.773	0.748	0.229
Without Causal graph	0.815	0.792	0.239	0.807	0.787	0.235
Without uncertainty estimation	0.850	0.831	0.244	0.843	0.825	0.241
Without concept drift detection	0.857	0.838	0.246	0.849	0.831	0.242

đề xuất có mức bất định thấp nhất trên cả hai bộ dữ liệu (Bộ dữ liệu 1: 0.093; Bộ dữ liệu 2: 0.092), phản ánh độ tin cậy cao hơn so với các cấu hình còn lại. Các mô hình thành phần như TITAN và DRLEARNER có bất định cao hơn đáng kể, trong khi biến thể TITAN + MLP + RL được cải thiện nhưng vẫn không đạt mức ổn định như mô hình đầy đủ. Các kết quả này nhấn mạnh vai trò của việc tích hợp học biểu diễn sâu, ước lượng nhân quả và điều phối thích ứng trong việc giảm bất định, nâng cao khả năng tổng quát hóa và tăng độ tin cậy của đầu ra.

Bảng 4.10: Monte Carlo Dropout: Bất định dự đoán trên hai bộ dữ liệu

Method	Bộ dữ liệu 1	Bộ dữ liệu 2
MH đề xuất	0.093	0.092
TITAN	0.126	0.127
DRLEARNER	0.133	0.134
TITAN + MLP + RL	0.102	0.103

4.5.3. Phân tích chi phí tính toán và khả năng mở rộng

Để đánh giá tính khả thi triển khai của mô hình đề xuất, nghiên cứu tiến hành phân tích chi phí tính toán thông qua hai khía cạnh chính: (i) thời gian huấn luyện và (ii) hiệu năng suy luận.

Bảng 4.11 trình bày thời gian huấn luyện của các phương pháp trên hai tập dữ liệu. Kết quả cho thấy mô hình đề xuất có thời gian huấn luyện cao hơn so với các phương pháp cơ sở, với 67.3 phút trên bộ dữ liệu 1 và 89.7 phút trên bộ dữ liệu 2. Điều này xuất phát từ việc mô hình tích hợp nhiều thành phần có độ phức tạp cao, bao gồm mô hình Transformer có bộ nhớ (TITAN), bộ ước lượng nhân quả DRLEARNER và cơ chế điều phối dựa trên học tăng cường.

Ngược lại, các phương pháp đơn giản hơn như DRLEARNER hoặc MLP kết hợp nhân quả có thời gian huấn luyện thấp hơn đáng kể do cấu trúc mô hình đơn giản và số lượng tham số ít hơn. Tuy nhiên, các phương pháp này bị hạn chế trong việc mô hình hóa các quan hệ phi tuyến và động học phức tạp trong dữ liệu, dẫn đến hiệu năng dự báo thấp hơn trong các bài toán nhân quả.

Bên cạnh chi phí huấn luyện, hiệu năng suy luận đóng vai trò quan trọng đối với các hệ thống hỗ trợ ra quyết định trong y tế. Bảng 4.12 trình bày độ trễ suy luận (latency), thông lượng (throughput), kích thước mô hình và mức độ phức tạp triển khai của các

Bảng 4.11: So sánh thời gian huấn luyện (phút)

Phương pháp	Bộ dữ liệu 1	Bộ dữ liệu 2	Thành phần chính
CAUSALRLSTACK	67.3	89.7	TITAN + DRLearner + RL Agent
TITAN	34.8	46.2	Transformer có bộ nhớ
DRLearner	8.4	12.7	Random Forest + Logistic Regression
MLP + Causal	4.2	6.1	MLP đơn giản
Double ML	12.3	18.9	Meta-learning với cross-fitting

phương pháp trên bộ dữ liệu 1.

Kết quả cho thấy mô hình đề xuất đạt độ trễ suy luận trung bình 4.7 ms trên mỗi mẫu, tương ứng với thông lượng khoảng 213 mẫu mỗi giây, với kích thước mô hình khoảng 156 MB. Mặc dù chậm hơn so với các mô hình đơn giản như MLP (0.3 ms/mẫu), mức độ này vẫn nằm trong phạm vi chấp nhận được đối với các hệ thống xử lý gần thời gian thực.

Bảng 4.12: Hiệu năng suy luận của các phương pháp trên bộ dữ liệu 1

Phương pháp	Latency (ms/mẫu)	Throughput (mẫu/giây)	Kích thước (MB)	Độ phức tạp
MH đề xuất	4.7	213	156.3	Cao
TITAN	2.9	345	89.7	Trung bình
DRLearner	0.8	1250	12.4	Thấp
MLP + Causal	0.3	3333	2.8	Rất thấp
TITAN + MLP	2.1	476	67.2	Trung bình
TITAN + MLP + RL	3.8	263	124.6	Cao
Double ML	1.2	833	18.7	Thấp
Orthogonal RF	1.8	556	45.3	Thấp
CEVAE	2.4	417	78.9	Trung bình
CAT Networks	2.6	385	82.1	Trung bình
Causal Forest	1.5	667	6.8	Thấp
X-Learner	0.9	1111	15.2	Thấp

Về khả năng mở rộng, kết quả thực nghiệm cho thấy độ trễ suy luận tăng theo mức độ phức tạp của mô hình, tuy nhiên vẫn duy trì ở mức thấp và ổn định. Điều này cho thấy mô hình có khả năng áp dụng cho các tập dữ liệu có quy mô lớn hơn mà không làm suy giảm đáng kể hiệu năng suy luận. Đồng thời, việc tách biệt pha huấn luyện và suy luận cho phép mô hình được triển khai hiệu quả trong thực tế, trong đó chi phí huấn luyện cao không ảnh hưởng trực tiếp đến tốc độ phản hồi khi vận hành hệ thống.

Nhìn chung, tồn tại sự đánh đổi giữa chi phí tính toán và hiệu năng dự báo. Mô hình đề xuất chấp nhận chi phí huấn luyện cao hơn để đạt được khả năng mô hình hóa tốt hơn các quan hệ nhân quả phức tạp, trong khi vẫn đảm bảo hiệu năng suy luận ở mức phù hợp cho triển khai thực tế.

4.6. Kết luận chương

Chương này tập trung nghiên cứu bài toán suy luận nhân quả phục vụ ra quyết định trong y tế công cộng, với trọng tâm là các bối cảnh dữ liệu quan sát, dị thể và

có nhiều nguồn bất định. Khác với các hướng tiếp cận dự báo thuần túy, Chương này tiếp cận bài toán ở cấp độ quyết định, nơi mục tiêu không chỉ là dự đoán kết cục, mà là ước lượng hiệu quả can thiệp và hỗ trợ lựa chọn chiến lược tác động phù hợp cho từng cá thể và từng nhóm đối tượng.

Trên cơ sở đó, Chương đã đề xuất và phát triển một mô hình suy luận nhân quả thích ứng, kết hợp có hệ thống giữa học biểu diễn sâu, ước lượng nhân quả và điều phối động dựa trên học tăng cường. Cách tiếp cận này cho phép tách biệt nhưng vẫn phối hợp chặt chẽ các thành phần cốt lõi của suy luận nhân quả, bao gồm: (i) kiểm soát biến và giả định nhân quả thông qua đồ thị nhân quả, (ii) học biểu diễn linh hoạt cho dữ liệu dị thể, (iii) ước lượng hiệu quả can thiệp có tính bền vững thống kê, và (iv) cơ chế điều phối thích ứng theo từng cá thể trong điều kiện bất định.

Các thí nghiệm trên dữ liệu mô phỏng cho thấy mô hình đề xuất có khả năng phục hồi chính xác hiệu quả can thiệp cá thể trong điều kiện kiểm soát, từ đó xác nhận tính đúng đắn về mặt lý thuyết của mô hình. Trên các bộ dữ liệu thực tế, mô hình đề xuất thể hiện hiệu năng vượt trội so với các mô hình cơ sở, không chỉ về các chỉ số dự báo như Accuracy và F1-score, mà quan trọng hơn là về độ ổn định và độ tin cậy của ước lượng hiệu quả can thiệp.

Nghiên cứu loại bỏ cho thấy hiệu quả của mô hình không đến từ một thành phần đơn lẻ, mà là kết quả của sự phối hợp có chủ đích giữa các mô-đun. Đặc biệt, cơ chế điều phối động dựa trên học tăng cường đóng vai trò trung tâm trong việc cá thể hóa suy luận nhân quả, trong khi các mô-đun học biểu diễn sâu và ước lượng doubly robust đảm bảo năng lực biểu diễn và tính hợp lệ thống kê của các ước lượng.

Tổng thể, các kết quả đạt được đã chỉ ra rằng chương đã giải quyết được một hạn chế cốt lõi của các phương pháp suy luận nhân quả hiện nay khi áp dụng vào dữ liệu HIV, đồng thời mở ra một hướng tiếp cận mới cho việc xây dựng các hệ thống hỗ trợ ra quyết định dựa trên bằng chứng nhân quả.

KẾT LUẬN CHUNG

Luận án này được thực hiện với mục tiêu phát triển các phương pháp học sâu phục vụ dự báo và hỗ trợ ra quyết định trong giám sát và điều trị HIV. Luận án xem xét nhu cầu phân tích dữ liệu xuyên suốt từ giám sát dịch tễ ở cấp cộng đồng, theo dõi tiến triển bệnh ở cấp cá nhân, đến hỗ trợ đánh giá các chiến lược can thiệp. Trên cơ sở đó, luận án tập trung phát triển các phương pháp học sâu cho ba bài toán cốt lõi trong giám sát và điều trị HIV, bao gồm dự báo dịch tễ, phân tích sống sót và suy luận nhân quả.

Đóng góp thứ nhất của luận án là đề xuất một mô hình học sâu dựa trên đồ thị nhằm dự báo ngắn hạn số ca nhiễm HIV ở cấp quận/huyện từ dữ liệu giám sát ca bệnh. Phương pháp được xây dựng trên cơ sở khai thác các mối quan hệ không gian, thời gian và dịch tễ giữa các ca bệnh để học biểu diễn từ cấu trúc đồ thị ở quy mô toàn thành phố, đồng thời tổng hợp thông tin theo từng khu vực nhằm phục vụ dự báo ở cấp địa phương. Cách tiếp cận này cho phép kết hợp đồng thời thông tin toàn cục và cục bộ, qua đó cải thiện khả năng dự báo so với các phương pháp dựa trên dữ liệu tổng hợp hoặc mô hình chuỗi thời gian truyền thống.

Đóng góp thứ hai là phát triển một mô hình phân tích sống sót dựa trên mô hình không gian trạng thái chọn lọc nhằm dự báo tiến triển bệnh và nguy cơ xảy ra biến cố trong điều trị HIV. Phương pháp được đề xuất cho phép mô hình hóa động học tiềm ẩn của trạng thái bệnh theo thời gian, đồng thời tăng khả năng chọn lọc các tín hiệu quan trọng từ dữ liệu lâm sàng dọc. So với các phương pháp phân tích sống sót truyền thống hoặc các mô hình học sâu chỉ khai thác trực tiếp dữ liệu quan sát, cách tiếp cận này cung cấp khả năng biểu diễn linh hoạt hơn đối với các quá trình tiến triển bệnh phức tạp.

Đóng góp thứ ba của luận án là đề xuất một khung suy luận nhân quả nhằm hỗ trợ đánh giá tác động của các can thiệp trên dữ liệu quan sát. Phương pháp kết hợp học biểu diễn sâu với các kỹ thuật ước lượng nhân quả và cơ chế tối ưu thích ứng nhằm nâng cao độ chính xác và độ ổn định trong ước lượng hiệu ứng điều trị. Đóng góp này mở rộng phạm vi ứng dụng của học sâu từ dự báo kết quả sang hỗ trợ ra quyết định can thiệp, góp phần tăng giá trị ứng dụng thực tiễn của các phương pháp phân tích dữ liệu trong quản lý HIV.

Tổng thể, luận án đóng góp cả về phương diện phương pháp luận và ứng dụng. Về phương diện phương pháp, nghiên cứu mở rộng khả năng ứng dụng của các kiến trúc học sâu hiện đại vào các bài toán phân tích dữ liệu HIV có cấu trúc phức tạp, bao gồm dữ liệu đồ thị, dữ liệu theo thời gian và dữ liệu quan sát phục vụ suy luận nhân quả. Về phương diện ứng dụng, các phương pháp được đề xuất góp phần xây dựng nền tảng phân tích dữ liệu phục vụ giám sát dịch tễ, tiên lượng điều trị và hỗ trợ hoạch định can thiệp trong bối cảnh y tế công cộng.

Bên cạnh các kết quả đạt được, luận án cũng mở ra một số hướng nghiên cứu tiếp theo. Thứ nhất, việc mở rộng các mô hình với nhiều nguồn dữ liệu đa dạng hơn, chẳng hạn như dữ liệu hành vi, dữ liệu can thiệp, dữ liệu mạng lưới xã hội hoặc dữ liệu cập nhật theo thời gian thực, có thể giúp phản ánh toàn diện hơn diễn tiến của dịch HIV

và nâng cao hiệu quả phân tích. Thứ hai, các hướng nghiên cứu tiếp theo có thể tập trung vào việc tích hợp chặt chẽ hơn giữa dự báo, tiên lượng và suy luận nhân quả để xây dựng các hệ thống hỗ trợ ra quyết định thống nhất, không chỉ dự đoán diễn biến bệnh mà còn đánh giá tác động của các chiến lược can thiệp khác nhau. Cuối cùng, việc triển khai và đánh giá các mô hình trong môi trường vận hành thực tế, tích hợp với các hệ thống giám sát và điều trị HIV, sẽ là bước quan trọng để chuyển hóa các đóng góp học thuật của luận án thành các công cụ hỗ trợ ra quyết định có giá trị ứng dụng trong thực tiễn.

DANH MỤC CÁC CÔNG TRÌNH CỦA TÁC GIẢ

- CT1. Phạm Thành Đạt**, Khai Quang Tran, Viet Anh Nguyen, “CAUSALRLSTACK: Adaptive balancing of deep representation and causal effect estimation with application to HIV-related health data,” *BioData Mining*, vol. 18, p. 77, 2025, ISSN 1756-0381, <https://doi.org/10.1186/s13040-025-00492-3> (SCIE Q1)
- CT2. Phạm Thành Đạt**, Duong Van Nguyen, Thanh Tan Tran, Viet Anh Nguyen, “Casenet: A novel HIV forecasting model using attention-based spatio-temporal graph network leveraging city-wide infection correlations,” *Neural Computing and Applications*, vol. 38, p. 119, 2026, ISSN 0941-0643, <https://doi.org/10.1007/s00521-025-11804-3> (SCOPUS Q1)
- CT3. Phạm Thành Đạt**, Duong Nguyen Van, Thanh Tran Tan, Viet Anh Nguyen, “Spatio-temporal graph learning with epidemiological factors for HIV epidemic short-term prediction,” *Journal of Computer Science and Cybernetics*, 2024, ISSN 1813-9663, <https://doi.org/10.15625/1813-9663/21185>
- CT4. Phạm Thành Đạt**, Duong Van Nguyen, Thanh Tan Tran, Thang Truong Nguyen, “ViT-Survive: Temporal patch-based survival analysis of HIV treatment dropout using longitudinal data,” in *Advances in Information and Communication Technology*, Proc. ICTA 2025, LNNS 1833, Springer, 2026, pp. 1–10, https://doi.org/10.1007/978-3-032-18162-6_14
- CT5. Phạm Thành Đạt**, Nguyen Van Duong, Tran Tan Thanh, Vu Tuan Anh, Nguyen Truong Thang, “Improving survival prediction of HIV treatment dropout with group-wise self-attention network,” *Proceedings of the 18th National Conference on Fundamental and Applied Information Technology Research (FAIR 2025)*, Dec. 2025, <https://doi.org/10.15625/vap.2025.0322>
- CT6. Phạm Thành Đạt**, Nguyen Van Duong, and Nguyen Viet Anh, “Harnessing spatio-temporal learning with 2-layer attention-driven graph neural networks for HIV epidemic forecasting in Ho Chi Minh City,” *Proceedings of the 27th National Conference on Information and Communication Technology (ICTA 2024)*, Nha Trang, Vietnam, Oct. 2024
- CT7. Phạm Thành Đạt**, Nguyen Van Duong, Nguyen Tuan Cuong, and Ha Thi Hong Van, “Apply probabilistic approach with unsupervised learning for patient matching in health information exchange,” *Proceedings of the 16th National Conference on Fundamental and Applied Information Technology Research (FAIR 2023)*, Da Nang, Vietnam, Sep. 2023, <https://doi.org/10.15625/vap.2023.0031>
- CT8. Phạm Thanh Đạt**, Tran Quang Khai, and Nguyen Viet Anh, “A Unified Memory-Guided Framework for Disease State Learning in Personalized HIV Survival Analysis, submitted to Discover Artificial Intelligent, 2026.

TÀI LIỆU THAM KHẢO

1. UNAIDS, “Unaid fact sheet 2024: Global hiv & aids statistics,” 2024, accessed June 2025. [Online]. Available: <https://www.unaids.org/en/resources/fact-sheet>
2. World Health Organization, “Hiv data and statistics,” 2023, accessed June 2025. [Online]. Available: <https://www.who.int/data/gho/data/themes/hiv-aids>
3. UNAIDS Asia Pacific, “Asia and the pacific: Hiv epidemic and response 2023,” 2024, accessed June 2025. [Online]. Available: <https://www.unaids.org/en/resources/documents/2024/asia-pacific-hiv-data-2023>
4. Thủ tướng Chính phủ, “Quyết định số 1246/qĐ-ttg ngày 14 tháng 8 năm 2020 của thủ tướng chính phủ về việc phê duyệt chiến lược quốc gia chấm dứt dịch bệnh aids vào năm 2030,” 2020. [Online]. Available: <https://chinhphu.vn/default.aspx?docid=200774&pageid=27160>
5. A. S. Fauci, R. R. Redfield, G. Sigounas, M. D. Weahkee, and B. P. Giroir, “Ending the hiv epidemic: A plan for the united states,” *JAMA*, vol. 321, no. 9, pp. 844–845, 2019, doi: 10.1001/jama.2019.1343.
6. H. C. Kim, W. Yoo, T. Yoo, and E. H. Lee, “Forecasting the number of human immunodeficiency virus infections in the korean population using the autoregressive integrated moving average model,” *Osong Public Health and Research Perspectives*, vol. 4, no. 6, pp. 358–362, 2013, doi: 10.1016/j.phrp.2013.10.009.
7. X. Zhang *et al.*, “Epidemiological and time series analysis on the incidence and death of aids and hiv in china,” *BMC Public Health*, vol. 20, p. 1912, 2020, doi: 10.1186/s12889-020-09977-8.
8. G. Wang, W. Wei, J. Jiang, C. Ning, H. Chen, J. Huang, B. Liang, N. Zang, Y. Liao, R. Chen, J. Lai, O. Zhou, J. Han, H. Liang, and L. Ye, “Application of a long short-term memory neural network: a burgeoning method of deep learning in forecasting hiv incidence in guangxi, china,” *Epidemiology and Infection*, vol. 147, p. e194, Jan 2019, doi: 10.1017/S095026881900075X.
9. X. Guo, Y. Yang, S.-P. Tseng, Z. Yin, and C. Wang, “Research on hiv/aids epidemic trend prediction model based on arima-lstm,” in *2022 10th International Conference on Orange Technology (ICOT)*, 2022, pp. 1–4, doi: 10.1109/ICOT56925.2022.10008185.
10. P. T. Dat, D. N. Van, T. T. Tan, and V. A. Nguyen, “Spatio-temporal graph learning with epidemiological factors for hiv epidemic short-term prediction,” *Journal of Computer Science and Cybernetics*, vol. 40, no. 4, pp. 363–380, 2024, doi: 10.15625/1813-9663/21185.

11. P. T. Dat, N. V. Duong, and N. V. Anh, “Harnessing spatio-temporal learning with 2-layer attention-driven graph neural networks for hiv epidemic forecasting in ho chi minh city,” in *Proceedings of the 27th National Conference on Information and Communication Technology (ICTA 2024)*, Nha Trang, Vietnam, Oct. 2024.
12. G. Panagopoulos, G. Nikolentzos, and M. Vazirgiannis, “Transfer graph neural networks for pandemic forecasting,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, 2021, pp. 4838–4845, doi: 10.1609/aaai.v35i6.16616.
13. S. Yu, F. Xia, S. Li, M. Hou, and Q. Z. Sheng, “Spatio-temporal graph learning for epidemic prediction,” *ACM Trans. Intell. Syst. Technol.*, vol. 14, no. 2, Feb. 2023, doi: 10.1145/3579815. [Online]. Available: <https://doi.org/10.1145/3579815>
14. G. Wang, W. Wei, J. Jiang, C. Ning, H. Chen, J. Huang, B. Liang, N. Zang, Y. Liao, R. Chen, J. Lai, O. Zhou, J. Han, H. Liang, and L. Ye, “Application of a long short-term memory neural network: A burgeoning method of deep learning in forecasting hiv incidence in guangxi, china,” *Epidemiology and Infection*, vol. 147, p. e194, 2019.
15. Z. Wang, H. Khaleghy, Y. Yao, and J. Chen, “Graph-aware spatio-temporal attention for forecasting hiv/aids case counts in public health surveillance,” *Frontiers in Public Health*, vol. 14, p. 1771586, 2026.
16. D. Tang, Y. Jin, X. Hu, D. Lin, A. Kapar, Y. Wang, F. Yang, and H. Li, “Study on the prediction performance of aids monthly incidence in xinjiang based on time series and deep learning models,” *BMC Public Health*, vol. 25, p. 780, 2025.
17. N. N. Mchunu, H. G. Mwambi, T. Reddy, N. Yende-Zuma, and K. Naidoo, “Joint modelling of longitudinal and time-to-event data: an illustration using cd4 count and mortality in a cohort of patients initiated on antiretroviral therapy,” *BMC Infectious Diseases*, vol. 20, no. 1, p. 256, 2020.
18. F. K. Muhammed, D. B. Belay, A. M. Presanis, and A. T. Sebu, “Dynamic predictions from longitudinal cd4 count measures and time to death of hiv/aids patients using a bayesian joint model,” *Scientific African*, vol. 19, p. e01519, 2023.
19. S. S. Pilangorgi, S. Khodakarim, Z. Shayan, and M. Nejat, “Evaluation of factors related to longitudinal cd4 count and the risk of death among hiv-infected patients using bayesian joint models,” *BMC Public Health*, vol. 25, no. 1, p. 979, 2025.
20. E. L. Kaplan and P. Meier, “Nonparametric estimation from incomplete observations,” *Journal of the American Statistical Association*, vol. 53, no. 282, pp. 457–481, 1958, doi: 10.1080/01621459.1958.10501452.
21. D. R. Cox, “Regression models and life-tables,” *Journal of the Royal Statistical Society: Series B (Methodological)*, vol. 34, no. 2, pp. 187–202, 1972.

22. J. L. Katzman, U. Shaham, A. Cloninger, J. Bates, T. Jiang, and Y. Kluger, “DeepSurv: Personalized treatment recommender system using a cox proportional hazards deep neural network,” *BMC Medical Research Methodology*, vol. 18, p. 24, 2018, doi: 10.1186/s12874-018-0482-1.
23. C. Lee, W. Zame, J. Yoon, and M. van der Schaar, “DeepHit: A deep learning approach to survival analysis with competing risks,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 32, no. 1, 2018, doi: 10.1609/aaai.v32i1.11842. [Online]. Available: <https://doi.org/10.1609/aaai.v32i1.11842>
24. H. Wang and J. Sun, “SurvTrace: Transformer-based survival analysis for electronic health records,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, 2021, pp. 626–633.
25. T. D. Pham, V. D. Nguyen, T. T. Tran, and T. T. Nguyen, “Vit-survive: Temporal patch-based survival analysis of hiv treatment dropout using longitudinal data,” in *Advances in Information and Communication Technology*, ser. Lecture Notes in Networks and Systems, vol. 1833. Springer, 2026.
26. T. D. Pham, V. D. Nguyen, T. T. Tran, T. A. Vu, and T. T. Nguyen, “Improving survival prediction of hiv treatment dropout with group-wise self-attention network,” in *Proceedings of the 18th National Conference on Fundamental and Applied Information Technology Research (FAIR 2025)*, Ha Noi, Vietnam, August 2025, august 21–22, 2025.
27. M. A. Hernán, B. Brumback, and J. M. Robins, “Marginal structural models to estimate the causal effect of zidovudine on the survival of hiv-positive men,” *Epidemiology*, vol. 11, no. 5, pp. 561–570, 2000, doi: 10.1097/00001648-200009000-00012.
28. J. A. C. Sterne, M. May, D. Costagliola, F. de Wolf, A. N. Phillips, R. Harris, M. J. Funk, R. B. Geskus, M. J. Gill, F. Dabis *et al.*, “Timing of initiation of antiretroviral therapy in aids-free hiv-1-infected patients: A collaborative analysis of 18 hiv cohort studies,” *New England Journal of Medicine*, vol. 360, no. 18, pp. 1815–1826, 2009, doi: 10.1056/NEJMoa0904184.
29. U. Shalit, F. D. Johansson, and D. Sontag, “Estimating individual treatment effect: Generalization bounds and algorithms,” in *Proceedings of the 34th International Conference on Machine Learning*, vol. 70. PMLR, 2017, pp. 3076–3085, doi: 10.48550/arXiv.1606.03976.
30. C. Shi, D. Blei, and V. Veitch, “Adapting neural networks for the estimation of treatment effects,” in *Advances in Neural Information Processing Systems*, vol. 32, 2019, pp. 14 017–14 030, doi: 10.48550/arXiv.1906.02120.
31. E. H. Kennedy, “Towards optimal doubly robust estimation of heterogeneous causal effects,” arXiv preprint, Tech. Rep., 2020, doi: 10.48550/arXiv.2004.14497.

32. M. A. Hernán, B. Brumback, and J. M. Robins, “Marginal structural models to estimate the causal effect of zidovudine on the survival of hiv-positive men,” *Epidemiology*, vol. 11, no. 5, pp. 561–570, 2000.
33. H. Ko, J. W. Hogan, and K. H. Mayer, “Estimating causal treatment effects from longitudinal hiv natural history studies using marginal structural models,” *Biometrics*, vol. 59, no. 1, pp. 152–162, 2003.
34. S. R. Cole, M. G. Hudgens, P. C. Tien, K. Anastos, L. Kingsley, J. S. Chmiel, and L. P. Jacobson, “Marginal structural models for case-cohort study designs to estimate the association of antiretroviral therapy initiation with incident aids or death,” *American Journal of Epidemiology*, vol. 175, no. 5, pp. 381–390, 2012.
35. UCI Machine Learning Repository, “AIDS Clinical Trials Group Study 175 (ACTG 175),” <https://archive.ics.uci.edu/dataset/890/aids+clinical+trials+group+study+175>, 2022, university of California, Irvine, School of Information and Computer Sciences.
36. “Actg 388 longitudinal cd4 t-cell dataset,” <https://sph.uth.edu/dept/bads/faculty-home/hulinwu/datasets/actg388longitudinaldatacd4tcell>, accessed: 2025-03-20.
37. D. Mesafint, “Edhs hiv/aids dataset,” <https://www.kaggle.com/datasets/danielmesafint1985/edhs-HIVAIDS-dataset>, 2022, accessed: 2025-05-07.
38. A. Velu, “AIDS Virus Infection Prediction,” <https://www.kaggle.com/datasets/aadarshvelu/AIDS-virus-infection-prediction>, 2022, accessed: 2025-05-07.
39. S. Zhong and L. Bian, “What drives disease flows between locations?” *Transactions in GIS*, vol. 24, no. 6, pp. 1740–1755, December 2020, epub 2020 Aug 5.
40. A. Sherstinsky, “Fundamentals of recurrent neural network (rnn) and long short-term memory (lstm) network,” *Physica D: Nonlinear Phenomena*, vol. 404, p. 132306, 2020, doi: 10.1016/j.physd.2019.132306.
41. R. K. Pathan, M. Biswas, and M. U. Khandaker, “Time series prediction of covid-19 by mutation rate analysis using recurrent neural network-based lstm model,” *Chaos, Solitons & Fractals*, vol. 138, p. 110018, Sep 2020, doi: 10.1016/j.chaos.2020.110018.
42. Y. Rizk, D. A. dos Santos, A. A. Moustafa, and A. E. Hassanien, “Predictions for covid-19 with deep learning models of lstm, gru and bi-lstm,” *Chaos, Solitons & Fractals*, vol. 140, p. 110212, 2020, doi: 10.1016/j.chaos.2020.110212.
43. K. E. ArunKumar, D. V. Kalaga, C. M. S. Kumar, M. Kawaji, and T. M. Brenza, “Forecasting of covid-19 using deep layer recurrent neural networks (rnns) with gated recurrent units (grus) and long short-term memory (lstm) cells,” *Chaos, Solitons & Fractals*, vol. 146, p. 110861, May 2021, doi: 10.1016/j.chaos.2021.110861.

44. R. R. Maaliw, Z. P. Mabunga, and F. T. Villa, "Time-series forecasting of covid-19 cases using stacked long short-term memory networks," in *2021 International Conference on Innovation and Intelligence for Informatics, Computing, and Technologies (3ICT)*, 2021, pp. 435–441, doi: 10.1109/3ICT53449.2021.9581688.
45. R. Ma, X. Zheng, P. Wang, H. Liu, and C. Zhang, "The prediction and analysis of covid-19 epidemic trend by combining lstm and markov method," *Scientific Reports*, vol. 11, no. 1, p. 17421, Aug 31 2021, doi: 10.1038/s41598-021-97037-5.
46. Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner, "Gradient-based learning applied to document recognition," *Proceedings of the IEEE*, vol. 86, no. 11, pp. 2278–2324, 1998.
47. M. Diqi, S. H. Mulyani, and R. Pradila, "Deepcov: Effective prediction model of covid-19 using cnn algorithm," *SN Computer Science*, vol. 4, no. 4, p. 396, 2023, doi: 10.1007/s42979-023-01834-w.
48. Z. Muhammad Zain and N. Alturki, "Covid-19 pandemic forecasting using cnn-lstm: A hybrid approach," *Journal of Control Science and Engineering*, vol. 2021, pp. 1–23, 07 2021, doi: 10.1155/2021/8785636.
49. A. Abade, L. F. Porto, A. R. Scholze, D. Kuntath, N. da Silva Barros, T. Z. Berra, A. C. V. Ramos, R. A. Arcêncio, and J. D. Alves, "A comparative analysis of classical and machine learning methods for forecasting tb/hiv co-infection," *Scientific Reports*, vol. 14, no. 1, p. 18991, 2024, doi: 10.1038/s41598-024-69580-4.
50. A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, "Attention is all you need," in *Proceedings of the 31st International Conference on Neural Information Processing Systems*, ser. NIPS'17. Red Hook, NY, USA: Curran Associates Inc., 2017, p. 6000–6010.
51. N. Wu, B. Green, X. Ben, and S. O'Banion, "Deep transformer models for time series forecasting: The influenza prevalence case," *arXiv preprint arXiv:2001.08317*, 2020, doi: 10.48550/arXiv.2001.08317.
52. L. Kalahasti, H. Youssef, Z. Miao *et al.*, "Foundation models for time series forecasting and policy evaluation in infectious disease epidemics," *Nature Communications*, vol. 16, no. 1, p. 2548, 2025, doi: 10.1038/s41467-025-57880-3.
53. C. Liu, Y. Zhan, J. Wu, C. Li, B. Du, W. Hu, T. Liu, and D. Tao, "Graph pooling for graph neural networks: progress, challenges, and opportunities," in *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence*, ser. IJCAI '23, 2023. [Online]. Available: <https://doi.org/10.24963/ijcai.2023/752>

54. R. Ying, J. You, C. Morris, X. Ren, W. L. Hamilton, and J. Leskovec, "Hierarchical graph representation learning with differentiable pooling," in *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, ser. NIPS'18. Red Hook, NY, USA: Curran Associates Inc., 2018, p. 4805–4815.
55. F. M. Bianchi, D. Grattarola, and C. Alippi, "Mincut pooling in graph neural networks," in *International Conference on Learning Representations (ICLR)*, 2020, doi: 10.48550/arXiv.1907.00481. [Online]. Available: <https://arxiv.org/abs/1907.00481>
56. H. Yuan and S. Ji, "Structpool: Structured graph pooling via conditional random fields," in *International Conference on Learning Representations*, 2020.
57. A. Ahmadi and P. Moradi, "Memory-based graph networks," *arXiv preprint arXiv:2002.09518*, 2020, doi: 10.48550/arXiv.2002.09518.
58. J. Lee, I. Lee, and J. Kang, "Self-attention graph pooling," in *Proceedings of the 36th International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, K. Chaudhuri and R. Salakhutdinov, Eds., vol. 97. Long Beach, California, USA: PMLR, 2019, pp. 3734–3743. [Online]. Available: <https://proceedings.mlr.press/v97/lee19c.html>
59. H. Gao and S. Ji, "Graph u-nets," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 44, no. 9, pp. 4948–4960, 2022, doi: 10.1109/TPAMI.2021.3081010.
60. A. Ahmadi and P. Moradi, "Memory-based graph networks," in *Proceedings of the International Conference on Learning Representations (ICLR)*, 06 2020, doi: 10.48550/arXiv.2002.09518.
61. L. Zhang, X. Wang, H. Li, G. Zhu, P. Shen, P. Li, X. Lu, S. A. A. Shah, and M. Bennamoun, "Structure-feature based graph self-adaptive pooling," in *Proceedings of The Web Conference 2020*, ser. WWW '20. New York, NY, USA: Association for Computing Machinery, 2020, p. 3098–3104, doi: 10.1145/3366423.3380083.
62. J. Zhou, G. Cui, S. Hu, Z. Zhang, C. Yang, Z. Liu, L. Wang, C. Li, and M. Sun, "Graph neural networks: A review of methods and applications," *AI Open*, vol. 1, pp. 57–81, 2020, doi: 10.1016/j.aiopen.2021.01.001.
63. T. N. Kipf and M. Welling, "Semi-supervised classification with graph convolutional networks," in *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*, 2017. [Online]. Available: <https://arxiv.org/abs/1609.02907>
64. P. Veličković, G. Cucurull, A. Casanova, A. Romero, P. Liò, and Y. Bengio, "Graph attention networks," *6th International Conference on Learning Representations*, 2017.

65. N. Li, L. Hu, Z.-L. Deng, T. Su, and J.-W. Liu, "Research on gru neural network satellite traffic prediction based on transfer learning," *Wireless Personal Communications*, vol. 118, pp. 1–13, 05 2021, doi: 10.1007/s11277-020-08045-z.
66. Z. Wu, M. Huang, A. Zhao, and Z. lan, "Traffic prediction based on gcn-lstm model," *Journal of Physics: Conference Series*, vol. 1972, p. 012107, 07 2021, doi: 10.1088/1742-6596/1972/1/012107.
67. T. Zhu, M. Boada, and M. J. Boada, "Adaptive graph attention and long short-term memory-based networks for traffic prediction," *Mathematics*, vol. 12, p. 255, 01 2024, doi: 10.3390/math12020255.
68. A. Pareja, G. Domeniconi, J. Chen, T. Ma, T. Suzumura, H. Kanezashi, T. Kaler, and C. Leiserson, "Evolvegcnn: Evolving graph convolutional networks for dynamic graphs," in *Proceedings of the AAAI Conference on Artificial Intelligence*, 2020.
69. M. Qiu, Z. Tan, and B.-K. Bao, "Msgnn: A multi-scale spatio-temporal graph neural network for epidemic forecasting," *Data Mining and Knowledge Discovery*, vol. 38, no. 4, pp. 2348–2376, 2024, doi: 10.1007/s10618-024-01035-w.
70. C. Lee, W. R. Zame, J. Yoon, and M. van der Schaar, "Dynamic-deephit: A deep learning approach for dynamic survival analysis with competing risks based on longitudinal data," *IEEE Transactions on Biomedical Engineering*, vol. 67, no. 1, pp. 122–133, 2020, doi: 10.1109/TBME.2019.2909027.
71. M. Poli, G. Pagano, D. Luan, T. Dao *et al.*, "Hyena: Efficient long-range sequence modeling with structured state spaces," *arXiv preprint arXiv:2302.10866*, 2023.
72. Z. Wu, A. Gu *et al.*, "State space models for sequence modeling," in *Advances in Neural Information Processing Systems*, vol. 34, 2021, pp. 14 042–14 054.
73. M. Poli, G. Pagano, D. Luan, T. Dao *et al.*, "Hyena: Efficient long-range sequence modeling with structured state spaces," *arXiv preprint arXiv:2302.10866*, 2023.
74. M. F. Gensheimer and B. Narasimhan, "A scalable discrete-time survival model for neural networks," *Journal of Machine Learning Research*, vol. 20, no. 22, pp. 1–24, 2019.
75. J. Durbin and S. J. Koopman, *Time Series Analysis by State Space Methods*, 2nd ed. Oxford, UK: Oxford University Press, 2012.
76. E. Begoli, T. Bhattacharya, and D. Kusnezov, "The need for uncertainty quantification in machine-assisted medical decision making," *Nature Machine Intelligence*, vol. 1, no. 1, pp. 20–23, 2019, doi: 10.1038/s42256-018-0004-1.

77. A. Kendall and Y. Gal, “What uncertainties do we need in bayesian deep learning for computer vision?” in *Advances in Neural Information Processing Systems*, vol. 30, 2017.
78. C. M. Bishop, *Pattern Recognition and Machine Learning*. Springer, 2006.
79. Y. Gal and Z. Ghahramani, “Dropout as a bayesian approximation: Representing model uncertainty in deep learning,” in *Proceedings of the 33rd International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, vol. 48. PMLR, 2016, pp. 1050–1059.
80. D. B. Rubin, “Estimating causal effects of treatments in randomized and nonrandomized studies,” *Journal of Educational Psychology*, vol. 66, no. 5, p. 688, 1974, doi: 10.1037/h0037350.
81. J. Pearl, *Causality: Models, Reasoning and Inference*. Cambridge University Press, 2009.
82. P. R. Rosenbaum and D. B. Rubin, “The central role of the propensity score in observational studies for causal effects,” *Biometrika*, vol. 70, no. 1, pp. 41–55, 1983, doi: 10.1093/biomet/70.1.41.
83. J. M. Robins, M. Á. Hernán, and B. Brumback, “Marginal structural models and causal inference in epidemiology,” *Epidemiology*, vol. 11, no. 5, pp. 550–560, 2000, doi: 10.1097/00001648-200009000-00011.
84. R. Wang, Y. Liu, Y. Cao, and L. Yao, “Causality-aware transformer networks for robotic navigation,” *arXiv preprint arXiv:2409.02669*, 2024, doi: 10.48550/arXiv.2409.02669.
85. S. R. Künzel, J. S. Sekhon, P. J. Bickel, and B. Yu, “Meta-learners for estimating heterogeneous treatment effects using machine learning,” *Proceedings of the National Academy of Sciences*, vol. 116, no. 10, pp. 4156–4165, 2019, doi: 10.1073/pnas.1804597116.
86. T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*, 2nd ed., ser. Springer Series in Statistics. New York, NY: Springer, 2009.
87. A. Géron, *Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow*, 3rd ed. Sebastopol, CA: O’Reilly Media, 2022.
88. C. M. Bishop, *Pattern Recognition and Machine Learning*. New York: Springer, 2006.
89. W. Wei, G. Wang, X. Tao, Q. Luo, L. Chen, X. Bao, Y. Liu, J. Jiang, H. Liang, and L. Ye, “Time series prediction for the epidemic trends of monkeypox using the arima, exponential smoothing, gm (1, 1) and lstm deep learning methods,” *Journal of General Virology*, vol. 104, no. 4, Apr 2023, doi: 10.1099/jgv.0.001839.

90. Y. Li, R. Yu, C. Shahabi, and Y. Liu, "Diffusion convolutional recurrent neural network: Data-driven traffic forecasting," in *6th International Conference on Learning Representations (ICLR)*, 2018. [Online]. Available: <https://arxiv.org/abs/1707.01926>
91. G. Panagopoulos, G. Nikolentzos, and M. Vazirgiannis, "Transfer graph neural networks for pandemic forecasting," in *Proceedings of the AAAI Conference on Artificial Intelligence*, 2021, pp. 4838–4845, doi: 10.1609/aaai.v35i6.16616.
92. T. Zhu, M. Boada, and M. J. Boada, "Adaptive graph attention and long short-term memory-based networks for traffic prediction," *Mathematics*, vol. 12, p. 255, 01 2024, doi: 10.3390/math12020255.
93. Z. Wu, M. Huang, A. Zhao, and Z. lan, "Traffic prediction based on gcn-lstm model," *Journal of Physics: Conference Series*, vol. 1972, p. 012107, 07 2021, doi: 10.1056/NEJM199610103351501.
94. N. Li, L. Hu, Z.-L. Deng, T. Su, and J.-W. Liu, "Research on gru neural network satellite traffic prediction based on transfer learning," *Wireless Personal Communications*, vol. 118, pp. 1–13, 05 2021.
95. A. Gu and T. Dao, "Mamba: Linear-time sequence modeling with selective state spaces," *arXiv preprint arXiv:2312.00752*, 2023, doi: 10.48550/arXiv.2312.00752.
96. S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Advances in Neural Information Processing Systems*, vol. 30. Curran Associates, Inc., 2017, doi: 10.5555/3295222.3295230.
97. S. Shastri, K. Singh, S. Kumar, P. Kour, and V. Mansotra, "Time series forecasting of covid-19 using deep learning models: India-usa comparative case study," *Chaos, Solitons & Fractals*, vol. 140, p. 110227, Nov 2020, doi: 10.1016/j.chaos.2020.110227.
98. A. Behrouz, P. Zhong, and V. Mirrokni, "Titans: Learning to memorize at test time," *arXiv preprint arXiv:2501.00663*, 2024, doi: 10.48550/arXiv.2501.00663.
99. R. S. Sutton and A. G. Barto, *Reinforcement Learning: An Introduction*. MIT Press, 2018.
100. J. Renero, I. Ochoa, and R. Maestre, "Rex: Causal discovery based on machine learning and explainability techniques," *arXiv preprint arXiv:2501.12706*, 2025, doi: 10.48550/arXiv.2501.12706. [Online]. Available: <https://doi.org/10.48550/arXiv.2501.12706>
101. J. M. Mooij, J. Peters, D. Janzing, J. Zscheischler, and B. Schölkopf, "Distinguishing cause from effect using observational data: Methods and benchmarks," *Journal of Machine Learning Research*, vol. 17, no. 32, pp. 1–102, 2016.

102. A. Gretton, O. Bousquet, A. Smola, and B. Schölkopf, “Measuring statistical dependence with hilbert-schmidt norms,” in *International Conference on Algorithmic Learning Theory*. Springer, 2005, pp. 63–77.
103. P. Young and S. Shellswell, “Time series analysis, forecasting and control,” *IEEE Transactions on Automatic Control*, vol. 17, no. 2, pp. 281–283, April 1972, doi: 10.1109/TAC.1972.1099963.
104. S. Aribé Jr, B. Gerardo, and R. Medina, “Time series forecasting of hiv/aids in the philippines using deep learning: Does covid-19 epidemic matter?” *International Journal of Emerging Technology and Advanced Engineering*, vol. 12, pp. 144–157, 2022, doi: 10.46338/ijetae0922_15.
105. V. R. Konda and J. N. Tsitsiklis, “Actor-critic algorithms,” in *Advances in Neural Information Processing Systems*, vol. 12. MIT Press, 1999, pp. 1008–1014.
106. Y. Gal and Z. Ghahramani, “Dropout as a bayesian approximation: Representing model uncertainty in deep learning,” in *Proceedings of the 33rd International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, vol. 48. PMLR, 2016, pp. 1050–1059, doi: 10.48550/arXiv.1506.02142.
107. D. M. Belete and M. D. Huchaiah, “Performance evaluation of classification models for HIV/AIDS dataset,” in *Data Management, Analytics and Innovation*, ser. Lecture Notes on Data Engineering and Communications Technologies. Pune, India: Springer, Singapore, 2021, vol. 70, pp. 109–125, doi: 10.1007/978-981-16-2934-1_7.
108. —, “Wrapper based feature selection techniques on EDHS-HIV/AIDS dataset,” *International Journal of Advanced Computer Science and Applications (IJACSA)*, vol. 12, no. 7, pp. 368–374, 2021, doi: 10.14569/IJACSA.2021.0120745.
109. S. M. Hammer, D. A. Katzenstein, M. D. Hughes, H. Gundacker, R. T. Schooley, R. Haubrich, W. K. Henry, M. M. Lederman, J. P. Phair, M. Niu, M. S. Hirsch, and T. C. Merigan, “A trial comparing nucleoside monotherapy with combination therapy in hiv-infected adults with cd4 cell counts from 200 to 500 per cubic millimeter,” *New England Journal of Medicine*, vol. 335, no. 15, pp. 1081–1090, 1996, doi: 10.1056/NEJM199610103351501.
110. Y. Zheng, Z. Li, J. Xin, and G. Zhou, “A spatial-temporal graph based hybrid infectious disease model with application to covid-19,” in *Proceedings of the 13th International Joint Conference on Biomedical Engineering Systems and Technologies (BIOSTEC)*, 2021, pp. 357–364, doi: 10.5220/0010349003570364.
111. M. Daw, A. Daw, N. Sifennasr, A. Draha, A. Daw, M. Ahmed, E. Mokhtar, A. El-Bouzedi, I. Daw, S. Adam, and S. W. I. association with Libyan Study Group of Hepatitis & HIV, “Spatiotemporal analysis and epidemiological characterization of the human immunodeficiency virus (hiv) in libya within a twenty-five-year

- period: 1993-2017,” *AIDS Research and Therapy*, vol. 16, no. 1, p. 14, 2019, doi: 10.1186/s12981-019-0228-0.
112. N. Admasu, A. Lomboro, E. Kebede, and et al., “Recent hiv infection and associated factors among newly diagnosed hiv cases in the southwest ethiopia regional state: Hiv case-based surveillance analysis (2019–2022),” *BMC Infectious Diseases*, vol. 24, p. 609, 2024, doi: 10.1186/s12879-024-09481-z.
 113. R. Harklerode, S. Schwarcz, J. Hargreaves, A. Boulle, J. Todd, S. Xueref, and B. Rice, “Feasibility of establishing hiv case-based surveillance to measure progress along the health sector cascade: Situational assessments in tanzania, south africa, and kenya,” *JMIR Public Health and Surveillance*, vol. 3, no. 3, p. e44, July 2017, doi: 10.2196/publichealth.7610.
 114. J. P. Vaughan, C. Victora, and A. M. R Chowdhury, “Outbreaks and Epidemics,” in *Practical Epidemiology: Using Epidemiology to Support Primary Health Care*. Oxford, UK: Oxford University Press, 10 2021, doi: 10.1093/med/9780192848741.003.0006.
 115. B. Xu, J. Li, and M. Wang, “Epidemiological and time series analysis on the incidence and death of aids and hiv in china,” *BMC Public Health*, vol. 20, p. 1906, 2020, doi: 10.1186/s12889-020-09977-8.
 116. P. Ngangue, E. Bedard, H. T. Zomahoun, J. Payne-Gagnon, C. Fournier, J. Afounde, and M.-P. Gagnon, “Returning for hiv test results: A systematic review of barriers and facilitators,” *International Scholarly Research Notices*, vol. 2016, p. 6304820, December 2016, doi: 10.1155/2016/6304820.