

**ĐẠI HỌC QUỐC GIA HÀ NỘI,
TRƯỜNG ĐẠI HỌC CÔNG NGHỆ**



Phạm Thị Quỳnh Trang

**Giảm thiểu quên nghiêm trọng trong học liên tục sử dụng các phương pháp phát lại
bộ nhớ, chuyển giao tri thức và kiến trúc động**

TÓM TẮT LUẬN ÁN NGÀNH HỆ THỐNG THÔNG TIN

Ngành đào tạo: Hệ thống thông tin

Mã số: 9480104

Hà Nội - 2026

**ĐẠI HỌC QUỐC GIA HÀ NỘI,
TRƯỜNG ĐẠI HỌC CÔNG NGHỆ**

Phạm Thị Quỳnh Trang

**Giảm thiểu quên nghiêm trọng trong học liên tục sử dụng các phương pháp phát lại
bộ nhớ, chuyển giao tri thức và kiến trúc động**

TÓM TẮT LUẬN ÁN NGÀNH HỆ THỐNG THÔNG TIN

Major: Information Systems

Mã số: 9480104

Cán bộ Hướng dẫn: PGS.TS Nguyễn Trí Thành

Cán bộ đồng hướng dẫn: TS. Lê Đức Trọng

Hà Nội - 2026

Tóm tắt

Các mô hình học máy truyền thống được huấn luyện trên toàn bộ dữ liệu cùng một lúc. Nhưng trong thực tế, dữ liệu đến liên tục theo thời gian, không thể huấn luyện lại từ đầu mỗi lần có thêm dữ liệu mới. **Học liên tục** ra đời để giải quyết điều này, cho phép mô hình học thêm tác vụ mới mà không “quên” những gì đã học trước đó. Hiện tượng mô hình quên kiến thức cũ khi học tác vụ mới được gọi là **quên nghiêm trọng (catastrophic forgetting)**, đây là thách thức trung tâm của lĩnh vực.

Ba hướng tiếp cận hiện có và hạn chế

Hướng tiếp cận	Ý tưởng	Hạn chế
Phát lại bộ nhớ	Lưu lại một phần dữ liệu cũ để ôn lại	Lo ngại riêng tư, tốn bộ nhớ
Chuyển giao tri thức	Chắt lọc kiến thức từ mô hình cũ sang mới	Bị giới hạn bởi dung lượng tham số
Mở rộng kiến trúc	Thêm thành phần mới cho mỗi tác vụ	Số tham số tăng tuyến tính

Đóng góp của luận án theo ba giai đoạn

Giai đoạn 1: Khi không thể lưu dữ liệu thật

Luận án đề xuất phương pháp **tăng cường bộ nhớ bằng nguyên mẫu**: thay vì lưu dữ liệu thật, các mẫu tổng hợp được tạo ra bằng cách thêm nhiễu Gaussian vào các nguyên mẫu đại diện. Đồng thời, một mô-đun cân chỉnh dựa trên **hàm năng lượng** giúp phát hiện và sửa chữa hiện tượng biểu diễn đặc trưng bị “trôi dạt” qua các tác vụ.

Giai đoạn 2: Khi không có bất kỳ bộ nhớ nào

Luận án phát triển hai phương pháp hoàn toàn không cần lưu dữ liệu cũ. Phương pháp **CL-plugin** thực hiện chuyển giao tri thức có chọn lọc thông qua hai hàm mất mát đặc biệt là CED và CSC, áp dụng cho bài toán phân loại cảm xúc theo khía cạnh. Phương

pháp **ZeroRFF** biến bài toán học tăng dần thành bài toán hồi quy có thể giải chính xác bằng công thức đóng, hoàn toàn không cần tối ưu hóa lặp lại.

Giai đoạn 3: Khi chuỗi tác vụ dài và đa dạng

Luận án đề xuất bốn phương pháp với kiến trúc linh hoạt và tham số được phân bổ động.

Khung **OPE** và **APF** tách biệt tường minh nhiệm vụ nhận dạng tác vụ và phân loại nhãn, từ đó giảm thiểu sự nhầm lẫn giữa các tác vụ khác nhau. Một biến thể mở rộng của khung này tích hợp thêm cơ chế căn chỉnh không gian ẩn, được áp dụng đặc biệt cho bài toán phân loại ảnh y tế.

Phương pháp **DAKTA** kết hợp LoRA với bộ phân loại Kolmogorov-Arnold (KAC) có khả năng mô hình hóa quan hệ đặc trưng linh hoạt hơn so với bộ phân loại tuyến tính truyền thống. Chính quy hóa Fisher Information Matrix được áp dụng đồng thời lên cả backbone lẫn bộ phân loại nhằm bảo vệ các tham số quan trọng qua từng tác vụ.

Phương pháp **SO-LoRA** xây dựng một cơ sở trực giao chung cố định cho toàn bộ tác vụ, trong đó mỗi tác vụ chỉ được phép học các tổ hợp thừa thớt trên hệ tọa độ này. Nhờ đó, dung lượng adapter được duy trì không đổi bất kể chuỗi tác vụ có dài đến đâu.

Kết quả

Các thực nghiệm toàn diện trên các bộ dữ liệu chuẩn về phân loại văn bản và phân loại ảnh cho thấy các phương pháp đề xuất đạt hiệu suất tốt nhất so với các phương pháp hiện có. Các phân tích ablation có hệ thống xác nhận vai trò của từng thành phần, trong khi phân tích lý thuyết chính thức cung cấp nền tảng vững chắc cho phương pháp SO-LoRA. Nhìn chung, luận án cung cấp các giải pháp có cơ sở lý luận rõ ràng và tính ứng dụng cao trong các môi trường có yêu cầu khắt khe về quyền riêng tư, bộ nhớ và tính toán.

MỞ ĐẦU

Động lực nghiên cứu

Học máy tĩnh huấn luyện mô hình trên tập dữ liệu cố định và triển khai mà không có bất kỳ sự thích nghi nào về sau. Mặc dù được áp dụng rộng rãi, mô hình này tồn tại bốn hạn chế có liên hệ mật thiết với nhau: không có khả năng thích nghi với sự thay đổi của phân phối dữ liệu, khả năng tổng quát hóa kém trên dữ liệu chưa từng gặp, thiếu cơ chế cập nhật liên tục, và dễ bị tổn thương trước hiện tượng concept drift. **Học liên tục (Continual Learning - CL)** khắc phục những hạn chế này bằng cách cho phép mô hình tiếp thu tri thức mới một cách tăng dần trong khi vẫn duy trì những gì đã học trước đó, qua đó giảm thiểu nhu cầu huấn luyện lại toàn bộ định kỳ và hiện thực hóa các hệ thống AI thích nghi trong môi trường thực tế phi tĩnh. Hình 1 đối chiếu hai mô hình học này, còn Hình 2 minh họa sự tăng trưởng nhanh chóng của nghiên cứu về CL, từ 243 công bố năm 2018 lên hơn 1.800 công bố vào năm 2024.



Figure 1: (a) Học máy tĩnh huấn luyện trên toàn bộ dữ liệu một lần duy nhất. (b) Học máy liên tục học trên dữ liệu mới khi chúng xuất hiện.

Mặc dù đã đạt được những tiến bộ đáng kể, ba khoảng trống nghiên cứu quan trọng là động lực chính của luận án này. Khoảng trống thứ nhất liên quan đến **quyền riêng tư dữ liệu**: phát lại bộ nhớ là hướng tiếp cận hiệu quả trong học liên tục đòi hỏi lưu trữ các

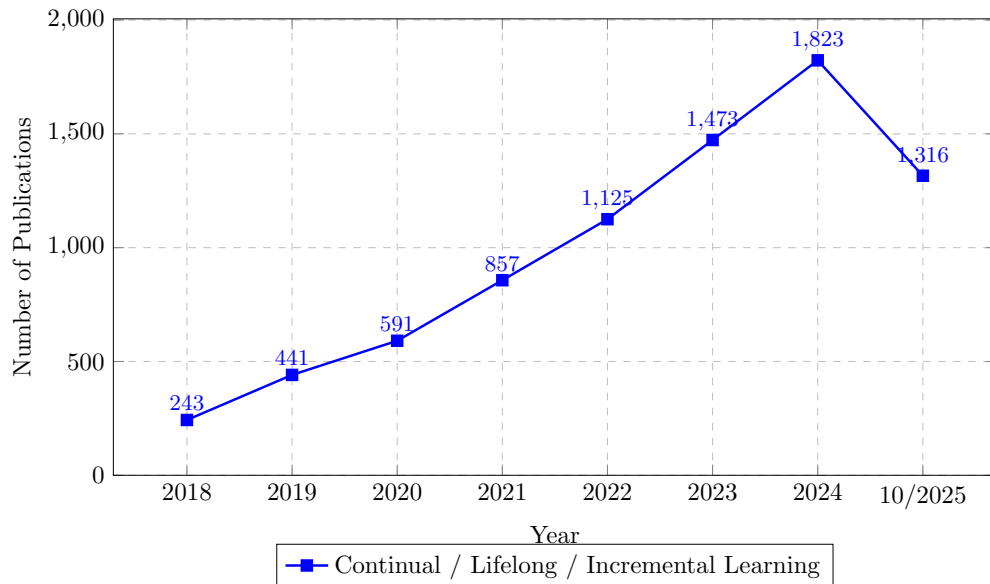


Figure 2: Xu hướng công bố hàng năm của các chủ đề *continual learning*, *lifelong learning* và *incremental learning* trên DBLP (2018–2025).

mẫu thô từ các tác vụ cũ, điều này không khả thi hoặc bị cấm trong các lĩnh vực nhạy cảm như y tế và tài chính. Khoảng trống thứ hai liên quan đến **sự phụ thuộc vào bộ nhớ**: ngay cả các phương pháp xấp xỉ nguyên mẫu giải quyết vấn đề về quyền riêng tư cũng trở nên kém tin cậy khi không gian đặc trưng ngày càng phức tạp, và bất kỳ bộ đệm bộ nhớ nào còn sót lại đều có thể không khả thi trong điều kiện ngân sách tính toán nghiêm ngặt. Khoảng trống thứ ba liên quan đến **dung lượng mô hình**: khi các tác vụ có tính đa dạng cao hoặc chuỗi tác vụ kéo dài, một ngân sách tham số cố định buộc mô hình phải đánh đổi giữa tính ổn định và tính dẻo dai mà cả chính quy hóa lẫn chất lọc tri thức đều không thể giải quyết triệt để.

Ba khoảng trống này xác định logic tiến triển của luận án. Chương 2 giải quyết khoảng trống thứ nhất bằng cách thay thế các mẫu thực được lưu trữ bằng các prototype lớp được nhiễu loạn Gaussian, kết hợp với mô-đun căn chỉnh ẩn dựa trên năng lượng. Chương 3 giải quyết khoảng trống thứ hai thông qua chuyển giao tri thức không cần exemplar: một CL-plugin với chất lọc đối lập cho kịch bản task-incremental, và ZeroRFF, áp dụng random Fourier features trên backbone đồng bộ với nghiệm học phân tích dạng đồng cho bài toán class-incremental learning. Chương 4 giải quyết khoảng trống thứ ba thông qua các chiến lược kiến trúc thích nghi, khung đa đầu phân loại tích hợp Online Prototype Equilibrium và Adaptive Prototypical Feedback, DAKTA kết hợp LoRA với bộ phân loại Kolmogorov-Arnold, và SO-LoRA thực thi tách biệt không gian con thông qua cơ sở trực giao cố định và độ thưa thớt nhóm.

Ba câu hỏi nghiên cứu cấu trúc nên toàn bộ quá trình điều tra là: (Q1) cơ chế tăng cường prototype có thể giảm thiểu hiện tượng quên nghiêm trọng trong khi vẫn bảo toàn quyền riêng tư như thế nào; (Q2) chuyển giao tri thức có chọn lọc có thể giảm thiểu dư thừa thông tin đến mức nào mà không cần lưu trữ exemplar; và (Q3) phân bổ tham số động có thể đồng thời ngăn chặn quên nghiêm trọng, hạn chế chi phí tính toán và cho phép chia sẻ tri thức giữa các tác vụ như thế nào. Các mục tiêu tương ứng là phát triển và kiểm chứng các phương pháp mới trong từng hạng mục nêu trên, với các đóng góp được công bố tại bảy công bố. Hình 3 minh họa cấu trúc tổng thể của luận án.

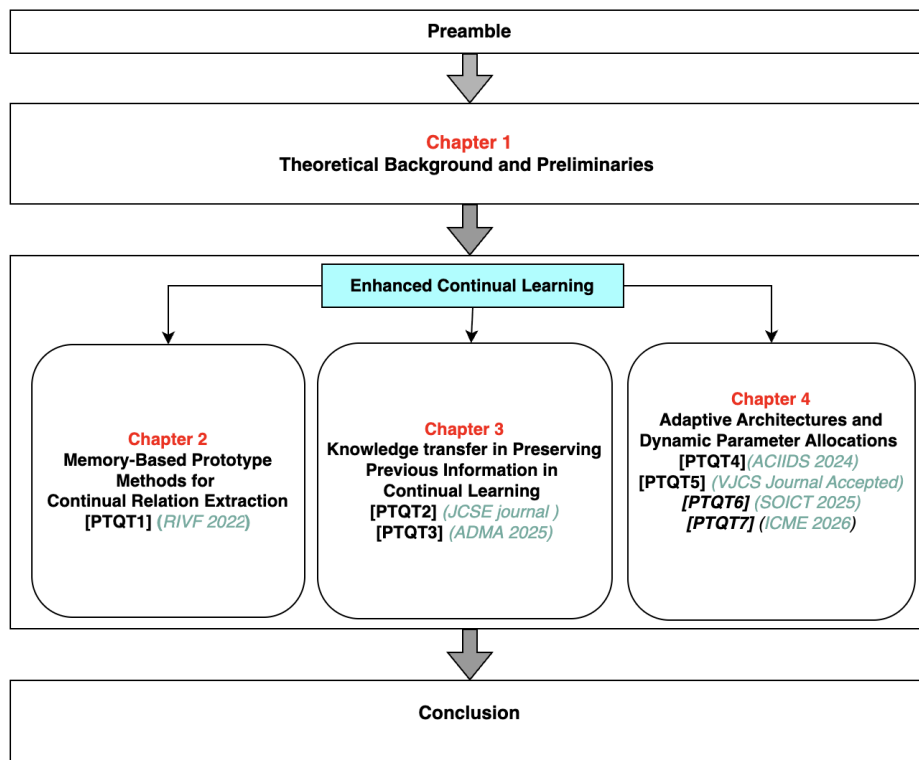


Figure 3: Sơ đồ cấu trúc luận án và các công bố liên quan trong từng chương.

Chapter 1

Nền tảng lý thuyết và các kiến thức tiên đề

Học liên tục (Continual Learning -CL), còn được gọi là học suốt đời hay học tăng dần, là một mô hình học máy trong đó mô hình được huấn luyện tuần tự trên một chuỗi các tác vụ, với mục tiêu đạt hiệu suất tương đương như khi toàn bộ dữ liệu được quan sát đồng thời. Nhìn từ góc độ học sâu, CL đặt nền móng cho các hệ thống thông minh có khả năng thích nghi với môi trường thay đổi, phản ánh năng lực của con người trong việc liên tục tiếp thu, cập nhật và vận dụng tri thức. Trở ngại cốt lõi trong mô hình này là **catastrophic forgetting**: khi mô hình thích nghi với một tác vụ mới, các bước cập nhật tham số chung làm suy giảm hiệu suất trên những tác vụ đã học trước đó. Sự mâu thuẫn giữa việc duy trì tri thức cũ và tiếp thu thông tin mới được gọi chính thức là **khả năng cân bằng tính ổn định-mềm dẻo (stability-plasticity dilemma)**, và đây là động lực chủ đạo của các thách thức nghiên cứu trong lĩnh vực này.

Hai định nghĩa nền tảng hình thành bức tranh lý thuyết của học liên tục. Thrun (1996) lần đầu tiên đề xuất rằng một hệ thống hoàn thành N tác vụ cần tận dụng tri thức tích lũy để mang lại lợi ích cho tác vụ thứ $(N + 1)$, mà không cần lưu giữ dữ liệu thô của các tác vụ trước. Chen và Liu (2018) sau đó mở rộng công thức hóa này, nhấn mạnh ba thuộc tính thiết yếu: học liên tục qua một chuỗi tác vụ, duy trì một Cơ sở Tri thức (Knowledge Base - KB) tích lũy và bảo toàn thông tin đã học, và khả năng khai thác tri thức quá khứ để cải thiện hiệu suất trong tương lai. Khác với định nghĩa trước, Chen và Liu hình thức hóa một cách tường minh việc bảo toàn toàn bộ thông tin tích lũy như một mục tiêu căn bản, thay vì chỉ tối ưu hóa cho tác vụ gần nhất.

Ngoài catastrophic forgetting, CL còn phải đối mặt với một loạt thách thức có liên hệ chặt chẽ với nhau. **Chuyển giao tri thức** đòi hỏi phải cân nhắc kỹ lưỡng tri thức nào cần chia sẻ, thời điểm nào nên chuyển giao, và cách thức thực hiện sao cho không gây ra nhiễu loạn tiêu cực. **Xung đột giữa các tác vụ** nảy sinh khi các biểu diễn được học cho tác vụ này lại làm suy giảm hiệu suất trên tác vụ khác, đòi hỏi sự cân bằng tinh tế giữa các tham số dùng chung và các tham số đặc thù cho từng tác vụ. **Dung lượng bộ nhớ hữu hạn** giới hạn lượng thông tin quá khứ có thể được lưu giữ, trong khi **trôi dạt phân phối** buộc mô hình phải thích nghi với những thay đổi dần dần hoặc đột ngột trong thống kê dữ liệu. Sau cùng, **hiệu quả mẫu** yêu cầu các kỹ năng mới phải được tiếp thu từ số lượng mẫu hạn chế, thay vì thông qua quá trình huấn luyện lại tốn kém.

Tùy theo cách thức nhận dạng tác vụ và cấu trúc không gian nhãn, tám kịch bản chuẩn đã được xác định. **Học tăng dần theo mẫu (Instance-Incremental Learning - IIL)** liên quan đến một tác vụ duy nhất với các mẫu đến theo luồng tuần tự. **Học tăng dần theo miền (Domain-Incremental Learning - DIL)** giữ nguyên không gian nhãn trong khi phân phối đầu vào thay đổi qua các miền, và danh tính tác vụ không cần thiết tại thời điểm kiểm tra. **Học tăng dần theo tác vụ (Task-Incremental Learning - TIL)** đưa vào các tập nhãn rời rạc với danh tính tác vụ được cung cấp tường minh tại thời điểm suy luận, khiến đây trở thành kịch bản đơn giản nhất. **Học tăng dần theo lớp (Class-Incremental Learning - CIL)**, kịch bản được nghiên cứu nhiều nhất, có cùng cấu trúc nhãn rời rạc nhưng ẩn đi danh tính tác vụ tại thời điểm kiểm tra, đòi hỏi phân loại thống nhất trên toàn bộ các lớp đã quan sát. Các kịch bản **không có ranh giới tác vụ (Task-Free - TFCL)** và **học trực tuyến (Online - OCL)** tiếp tục loại bỏ ranh giới tác vụ hoặc giới hạn huấn luyện chỉ một lần duyệt qua mỗi mẫu. Kịch bản **ranh giới mờ (Blurred Boundary - BBCL)** mô hình hóa các chuyển tiếp dần dần với phân phối lớp chồng chéo, trong khi **tiền huấn luyện liên tục (Continual Pre-training - CPT)** hướng đến việc cập nhật tuần tự các mô hình nền tảng lớn bằng các kho ngữ liệu đặc thù theo miền mới.

Năm nhóm phương pháp đã được phát triển để giải quyết các thách thức trên. Các phương pháp **dựa trên chính quy hóa** bổ sung hàm mất mát với một số hạng phạt nhằm làm chậm quá trình học đối với các tham số được coi là quan trọng với các tác vụ trước, tiêu biểu là Elastic Weight Consolidation (EWC), Synaptic Intelligence (SI) và Learning without Forgetting (LwF). Các phương pháp **dựa trên replay** lưu trữ hoặc tổng hợp các mẫu quá khứ để ôn luyện trong các bước huấn luyện tiếp theo (ví dụ: iCaRL, Dark Experience Replay, Gradient Episodic Memory), với chi phí là overhead bộ nhớ và nguy cơ vi phạm quyền riêng tư. Các phương pháp **dựa trên tối ưu hóa** trực tiếp điều chỉnh quy trình huấn luyện thông qua chiều gradient hoặc các chiến lược meta-learning. Các

phương pháp **dựa trên biểu diễn** tập trung vào việc học các không gian đặc trưng có khả năng chuyển giao và kháng trôi dạt thông qua tiền huấn luyện đối lập hoặc tự giám sát, kết hợp với việc duy trì prototype. Cuối cùng, các phương pháp **dựa trên kiến trúc** phân bổ các mô-đun đặc thù cho từng tác vụ trong một backbone dùng chung, như Progressive Neural Networks, PackNet và các thiết kế dựa trên adapter, cho phép thích nghi theo từng tác vụ mà không làm tổng số tham số tăng trưởng quá mức.

Hiệu suất trong CL được đánh giá theo ba chiều bổ sung cho nhau. Hiệu quả tổng thể được đo bằng độ chính xác trung bình (Average Accuracy — AA) và độ chính xác tăng dần trung bình (Average Incremental Accuracy — AIA). Tính ổn định bộ nhớ được phản ánh qua độ đo quên (Forgetting Measure — FM) và chuyển giao ngược (Backward Transfer — BWT), trong đó BWT âm cho thấy hiệu suất trên các tác vụ trước đó đã bị giảm. Tính mềm dẻo trong học tập được đánh giá thông qua chỉ số không nhượng bộ (Intransigence Measure — IM) và chuyển giao xuôi (Forward Transfer — FWT), lần lượt định lượng khả năng học các tác vụ mới và khả năng thụ hưởng từ các biểu diễn đã tích lũy trước đó của mô hình.

Tổng kết

Từ phân tích ba khoảng trống nghiên cứu còn mở, tạo nên động lực cho các đóng góp của luận án này. Thứ nhất, các phương pháp dựa trên phát lại bộ nhớ lưu trữ dữ liệu huấn luyện thực tế đặt ra những *mối lo ngại nghiêm trọng về quyền riêng tư*, đòi hỏi các giải pháp thay thế bảo toàn quyền riêng tư thông qua việc sinh dữ liệu tổng hợp. Thứ hai, các cơ chế chuyển giao tri thức dựa trên chất lọc hiện tại còn *thiếu hiệu quả*, thúc đẩy các chiến lược chất lọc có chọn lọc, tận dụng sức mạnh biểu diễn của các mô hình tiền huấn luyện. Thứ ba, các phương pháp dựa trên kiến trúc thường *thiếu cơ chế chia sẻ tri thức giữa các tác vụ*, chỉ ra hướng thiết kế kiến trúc động cân bằng giữa việc thích nghi đặc thù theo tác vụ và khai thác tri thức chung. Ba khoảng trống này xác định các hướng nghiên cứu cốt lõi được theo đuổi trong các chương tiếp theo của luận án.

Chapter 2

Các phương pháp phát lại bộ nhớ cho bài toán trích xuất quan hệ liên tục

Trích xuất quan hệ (Relation Extraction — RE) là tác vụ nhận dạng và phân loại các quan hệ ngữ nghĩa giữa các cặp thực thể trong văn bản, đóng vai trò là một tác vụ con cốt lõi của trích xuất thông tin trong các ứng dụng như hỏi đáp, phân tích cảm xúc và tìm kiếm có cấu trúc. Các mô hình RE truyền thống hoạt động trên một tập quan hệ được định nghĩa trước cố định, điều này hạn chế khả năng ứng dụng của chúng trong các tình huống thực tế, nơi các loại quan hệ mới liên tục xuất hiện. **Trích xuất quan hệ liên tục (Continual Relation Extraction — CRE)** khắc phục hạn chế này bằng cách huấn luyện mô hình tuần tự trên các tập quan hệ mới trong khi ngăn chặn catastrophic forgetting đối với các quan hệ đã học trước đó. Trong số các nhóm giải pháp chính, dựa trên chính quy hóa, kiến trúc động và dựa trên bộ nhớ, các phương pháp dựa trên bộ nhớ nhất quán đạt hiệu suất mạnh nhất trong các tác vụ NLP nhờ duy trì một ngân hàng bộ nhớ episodic lưu trữ các mẫu quá khứ để ôn luyện.

Công trình liên quan nổi bật nhất là **Học biểu diễn nhất quán (Consistent Representation Learning - CRL)**, kết hợp học đối lập có giám sát với chất lọc tri thức nhằm duy trì các embedding quan hệ ổn định qua các tác vụ tuần tự. Dù đạt hiệu suất tốt, CRL bộc lộ ba hạn chế then chốt: hoàn toàn phụ thuộc vào việc lưu trữ dữ liệu thực, khiến phương pháp này dễ bị tổn thương trước các ràng buộc bộ nhớ và lo ngại về quyền riêng tư; ngân sách bộ nhớ cho mỗi quan hệ ngày càng thu hẹp khi số lượng tác vụ tăng lên,

làm tăng nguy cơ overfitting trên các exemplar thừa thớt; và hiệu suất kém trên các bộ dữ liệu mất cân bằng như TACRED, nơi các lớp chiếm ưu thế kìm hãm việc học các quan hệ thiểu số trong quá trình phát lại bộ nhớ đối kháng.

Để giải quyết thách thức nêu trên, chương này đề xuất một phương pháp tích hợp **Tăng cường prototype (Prototype Augmentation - protoAug)** vào khung CRL, giảm sự phụ thuộc vào dữ liệu thực được lưu trữ. Quy trình huấn luyện gồm bốn giai đoạn. Trong giai đoạn *huấn luyện ban đầu*, một bộ mã hóa dựa trên BERT trích xuất các biểu diễn nhận thức thực thể bằng cách nối các trạng thái ẩn của các token đánh dấu biên thực thể đặc biệt, và bộ mã hóa cùng với đầu chiếu được tối ưu hóa bằng hàm mất mát đối lập có giám sát. Trong giai đoạn *chọn lọc mẫu*, thuật toán phân cụm k-means được áp dụng trên các embedding huấn luyện của tác vụ hiện tại, và mẫu gần nhất với tâm mỗi cụm được chọn để lưu trữ, đảm bảo ngân hàng bộ nhớ bao phủ phân phối đặc trưng của từng quan hệ thay vì dựa vào lấy mẫu ngẫu nhiên.

Điểm mới cốt lõi nằm ở giai đoạn *tăng cường bộ nhớ*. Thay vì chỉ replay các mẫu thực đã lưu trữ, một prototype cho mỗi quan hệ được tính là trung bình của các embedding được lưu trữ của nó, và các mẫu ảo được sinh ra bằng cách nhiễu loạn prototype này với nhiễu Gaussian được chia tỷ lệ theo vết của ma trận hiệp phương sai lớp. Các điểm dữ liệu tổng hợp này làm phong phú ngân hàng bộ nhớ mà không làm tăng lưu trữ dữ liệu thực, thúc đẩy cân bằng lớp và giảm thiểu tính nhạy cảm với overfitting. Trong giai đoạn *replay bộ nhớ* tiếp theo, bộ nhớ tăng cường kết hợp các mẫu thực đã lưu trữ và dữ liệu giả được sinh ra được dùng cho phát lại bộ nhớ đối lập và chất lọc tri thức, làm tăng khả năng phân biệt của các nguyên mẫu của các lớp cũ. Hình 2.3 minh họa kiến trúc tổng thể của mô-đun này.

Một phần mở rộng tiếp theo giới thiệu mô-đun **Căn chỉnh không gian ẩn dựa trên năng lượng (Energy-based Latent Space Alignment - ELI)** nhằm giải quyết hiện tượng trôi dạt biểu diễn giữa các tác vụ. Sau khi huấn luyện mỗi tác vụ, một mô hình dựa trên năng lượng (Energy-Based Model - EBM) được huấn luyện để gán năng lượng thấp cho các biểu diễn ẩn của mô hình trước đó và năng lượng cao cho các biểu diễn từ mô hình hiện tại, sử dụng Langevin dynamics để lấy mẫu xấp xỉ. Tại thời điểm suy luận, thuật toán ELI lặp đi lặp lại di chuyển vector ẩn của từng mẫu kiểm tra dọc theo gradient năng lượng về phía vùng năng lượng thấp, tái căn chỉnh các biểu diễn đã bị lệch trở về đa tạp đặc trưng của tác vụ trước mà không cần lưu trữ thêm dữ liệu. Hình 2.4 thể hiện kiến trúc đầy đủ tích hợp cả tăng cường prototype lẫn ELI.

Các thực nghiệm trên hai bộ benchmark FewRel và TACRED trong kịch bản task-

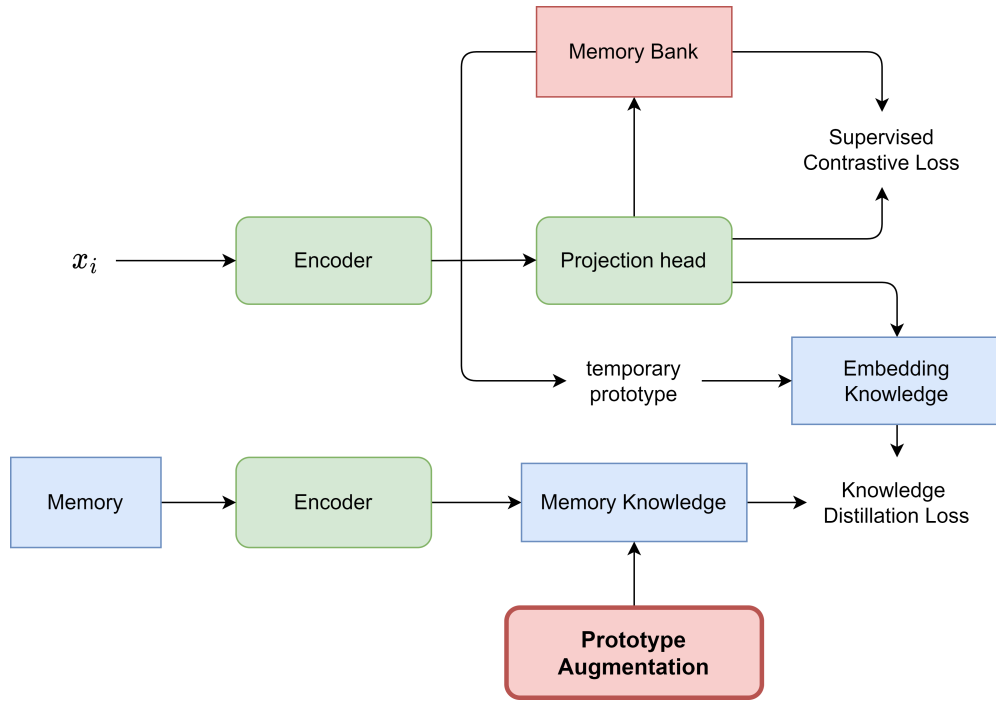


Figure 2.3: Kiến trúc của mô-đun học biểu diễn nhất quán với tăng cường prototype.

incremental với mười tác vụ tuần tự xác nhận hiệu quả của phương pháp được đề xuất. Bảng 2.1 và 2.2 báo cáo độ chính xác của từng phương pháp tại mỗi tác vụ. Trên FewRel, mô hình đề xuất vượt trội hơn CRL ở tám trong mười tác vụ với biên độ 1-2% và vượt qua tất cả các baseline khác trên toàn bộ chuỗi tác vụ. Trên TACRED, bộ dữ liệu khó hơn với tính mất cân bằng cao, mức cải thiện còn rõ rệt hơn, với mô hình vượt CRL tới 4% ở Tác vụ 4 và nhất quán vượt trội EMAR+BERT với biên độ dao động từ 1,1% đến 11,1%. Những cải thiện này càng trở nên nổi bật hơn khi số lượng tác vụ tăng lên, chứng tỏ tính ổn định dài hạn và khả năng giảm thiểu catastrophic forgetting vượt trội.

Tổng kết

Tóm lại, chương này chứng minh rằng việc thay thế các mẫu thực được lưu trữ bằng dữ liệu tổng hợp tăng cường prototype không chỉ là giải pháp đủ hiệu quả mà còn mang lại lợi ích rõ rệt cho bài toán trích xuất quan hệ liên tục. Sự kết hợp giữa chọn lọc mẫu dựa trên k-means, tăng cường prototype Gaussian và căn chỉnh ẩn dựa trên năng lượng tạo ra một giải pháp thay thế bảo toàn quyền riêng tư và tiết kiệm bộ nhớ so với các phương pháp rehearsal hiện có, đặt nền móng cho chương tiếp theo, nơi chúng tôi khảo sát các phương pháp chất lọc tri thức loại bỏ hoàn toàn nhu cầu lưu trữ bất kỳ dữ liệu nào.

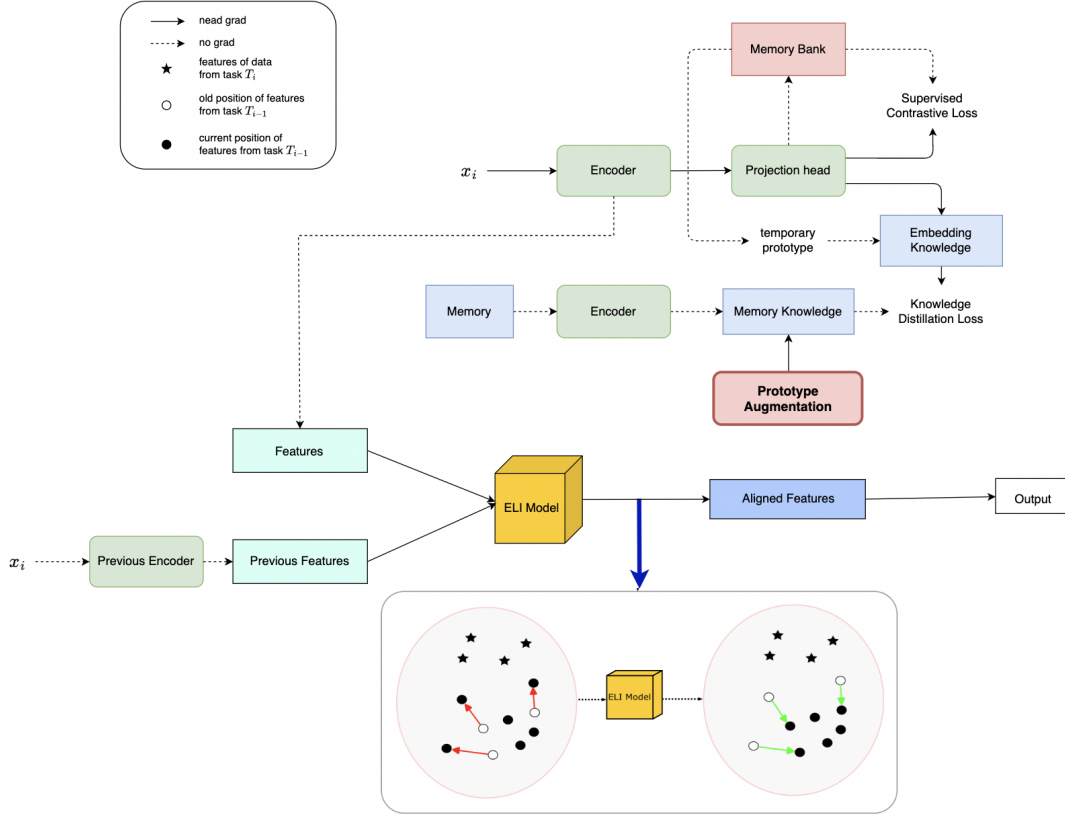


Figure 2.4: Kiến trúc tổng thể của mô hình trích xuất quan hệ liên tục được đề xuất, kết hợp tăng cường prototype và căn chỉnh ẩn dựa trên năng lượng (ELI).

Table 2.1: Độ chính xác (%) của từng mô hình trên tất cả các tác vụ trên bộ dữ liệu **FewRel**. In đậm = tốt nhất; gạch chân = tốt thứ hai.

Mô hình	T1	T2	T3	T4	T5	T6	T7	T8	T9	T10
EA-EMR	89.0	69.0	59.1	54.2	47.8	46.1	43.1	40.7	38.6	35.2
EMAR	88.5	73.2	66.6	63.8	55.8	54.3	52.9	50.9	48.8	46.3
CML	91.2	74.8	68.2	58.2	53.7	50.4	47.8	44.4	43.1	39.7
EMAR+BERT	98.8	89.1	89.5	85.7	83.6	84.8	79.3	80.0	77.1	73.8
RP-CRE	97.9	92.7	91.6	89.2	88.4	86.8	85.1	84.1	82.2	81.5
CRL	<u>98.2</u>	<u>94.6</u>	<u>92.5</u>	<u>90.5</u>	<u>89.4</u>	<u>87.9</u>	<u>86.9</u>	<u>85.6</u>	<u>84.5</u>	<u>83.1</u>
Ours	<u>98.2</u>	94.7	93.9	92.2	90.1	89.6	88.6	87.1	86.1	84.6

Table 2.2: Độ chính xác (%) của từng mô hình trên tất cả các tác vụ trên bộ dữ liệu **TACRED**. In đậm = tốt nhất; gạch chân = tốt thứ hai.

Mô hình	T1	T2	T3	T4	T5	T6	T7	T8	T9	T10
EA-EMR	47.5	40.1	38.3	29.9	24.0	27.3	26.9	25.8	22.9	19.8
EMAR	73.6	57.0	48.3	42.3	37.7	34.0	32.6	30.0	27.6	25.1
CML	57.2	51.4	41.3	39.3	35.9	28.9	27.3	26.9	24.8	23.4
EMAR+BERT	96.6	85.7	81.0	78.6	73.9	72.3	71.7	72.2	72.6	71.0
RP-CRE	97.6	90.6	86.1	82.4	79.8	77.2	75.1	73.7	72.4	72.4
CRL	<u>97.7</u>	<u>93.2</u>	<u>89.8</u>	<u>84.7</u>	<u>84.1</u>	<u>81.3</u>	<u>80.2</u>	<u>79.1</u>	<u>79.0</u>	<u>78.0</u>
Ours	97.8	95.2	90.6	88.2	84.6	81.7	82.8	80.5	80.0	79.6

Chapter 3

Chuyển giao tri thức trong bảo tồn thông tin trước đó trong học liên tục

Chương này khảo sát chuyển giao tri thức như một cơ chế giảm thiểu catastrophic forgetting trong học liên tục mà không cần dựa vào việc lưu trữ các exemplar thực từ các tác vụ trước. Hai chiến lược bổ trợ cho nhau được phát triển: chiến lược thứ nhất chuyển giao tri thức thông qua chất lọc tập hợp đối lập trong khung **Học tăng dần theo tác vụ (Task-Incremental Learning - TIL)** áp dụng cho bài toán Phân loại cảm xúc theo khía cạnh (Aspect Sentiment Classification - ASC); chiến lược thứ hai lưu giữ tri thức hoàn toàn bên trong một backbone tiền huấn luyện đóng băng thông qua học phân tích dạng đóng theo kịch bản **Học tăng dần theo lớp (Class-Incremental Learning - CIL)** / CIL tổng quát (GCIL). Cả hai cùng bao phủ hai hướng đi về mặt bản chất khác nhau để đạt tính ổn định, chất lọc dựa trên gradient và tính toán phân tích không cần gradient, đều đạt được học liên tục không cần exemplar.

Phần 1: Chất lọc tri thức kết hợp học đối lập (TIL).

Hướng tiếp cận thứ nhất giải quyết sự đánh đổi giữa tính ổn định và tính dẻo dai bằng cách kết hợp cô lập kiến trúc với học đối lập. Một **Plugin học liên tục (Continual Learning Plugin, CL-plugin)** được tích hợp vào BERT tại hai vị trí, thay thế adapter chuẩn. CL-plugin chứa hai mô-đun con: **Mô-đun chia sẻ tri thức (Knowledge Sharing Sub-Module - KSM)**, được triển khai dưới dạng các lớp capsule với định tuyến động

nhằm phân cụm tri thức chung giữa các tác vụ, và **Mô-đun đặc thù tác vụ (Task-Specific Sub-Module - TSM)**, được triển khai dưới dạng các lớp kết nối đầy đủ với mặt nạ mềm ở mức nơ-ron nhằm bảo vệ các tham số then chốt của tác vụ khỏi các bước cập nhật gradient trong các tác vụ tiếp theo. Hình 3.3 minh họa kiến trúc tổng thể và Hình 3.4 trình bày chi tiết cấu trúc CL-plugin.

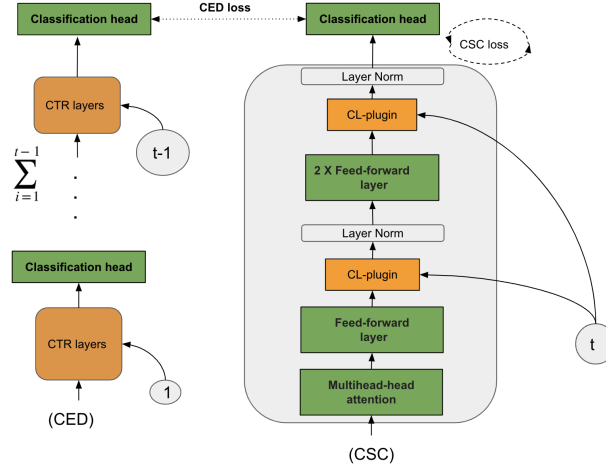


Figure 3.3: Kiến trúc tổng thể của mô hình CL-plugin được đề xuất.

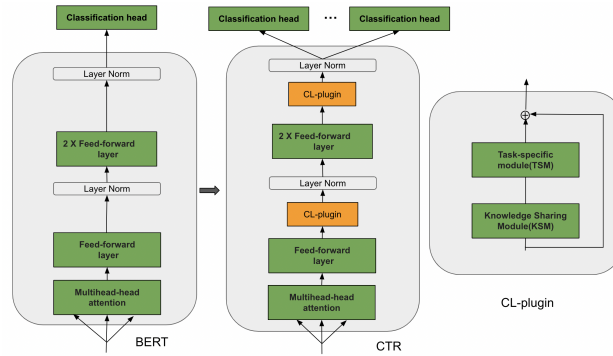


Figure 3.4: Kiến trúc tổng quát của BERT (trái), hệ thống CTR (giữa) và CL-plugin (phải) được tích hợp vào BERT.

Hai hàm mất mát đối lập bổ sung cho hàm mục tiêu cross-entropy chuẩn. **Hàm mất mát Chất lọc tập hợp đối lập (Contrastive Ensemble Distillation -CED)** chất lọc tri thức từ toàn bộ các mô hình tác vụ trước vào tác vụ hiện tại bằng cách coi cặp logit của cùng một đầu vào từ mô hình cũ và mô hình hiện tại là cặp dương, còn tất cả các cặp khác là cặp âm. Tổng hợp trên $t - 1$ mô hình giáo viên trước đó cho ta hàm mất mát CED đầy đủ. **Hàm mất mát Học đối lập có giám sát cho tác vụ hiện tại (Contrastive Supervised Classification - CSC)** tối đa hóa độ tương đồng nội lớp và tối thiểu hóa độ tương đồng liên lớp trong batch của tác vụ hiện tại, hoạt động trên đầu ra TSM đã được

che trước lớp phân loại. Ba hàm mục tiêu được kết hợp theo công thức:

$$\mathcal{L} = \mathcal{L}_{\text{CE}} + \mathcal{L}_{\text{CED}} + \mathcal{L}_{\text{CSC}}. \quad (3.1)$$

Các thực nghiệm được tiến hành trên 19 tác vụ ASC lấy từ bốn bộ benchmark (HL5Domains, Liu3Domains, Ding9Domains, SemEval14), được tóm tắt trong Bảng 3.1. Kết quả trung bình trên năm thứ tự tác vụ ngẫu nhiên được báo cáo trong Bảng 3.2. Mô hình đề xuất đạt Độ chính xác trung bình 88,50% và Macro-F1 cao nhất là 82,35% trong số tất cả các phương pháp học liên tục thuộc nhóm Adapter-BERT được đánh giá. CTR, đối thủ gần nhất, đạt độ chính xác 88,21% và Macro-F1 81,19% trong cùng điều kiện thứ tự ngẫu nhiên, thấp hơn ở cả hai chỉ số. LAMOL báo cáo độ chính xác cao hơn (88,91%) nhưng dựa vào backbone sinh GPT-2 cho pseudo-rehearsal, điều này đòi hỏi chi phí tính toán cao hơn đáng kể và đại diện cho một chiến lược về bản chất hoàn toàn khác.

Nghiên cứu ablation trong Bảng 3.3 xác nhận rằng mỗi thành phần đều đóng góp tích cực: loại bỏ CL-plugin làm giảm độ chính xác 2,67% và Macro-F1 2,60%, trong khi loại bỏ cả CED lẫn CSC gây ra sụt giảm 2,02% độ chính xác và 3,03% Macro-F1.

Phần 2: Lưu giữ tri thức không cần exemplar thông qua học phân tích (GCIL).

Hướng tiếp cận thứ hai loại bỏ hoàn toàn việc huấn luyện dựa trên gradient. Được thúc đẩy bởi những hạn chế của Học liên tục phân tích tổng quát (Generalized Analytic Continual Learning — GACL) — vốn áp dụng bộ phân loại tuyến tính trực tiếp trên vector đặc trưng thô và gặp khó khăn với các phân phối không phân tách tuyến tính — công trình này đề xuất **ZeroRFF**, một phương pháp tích hợp **Đặc trưng Fourier ngẫu nhiên (Random Fourier Features — RFF)** vào khung học liên tục phân tích. Hình 3.8 thể hiện kiến trúc tổng thể.

ZeroRFF hoạt động qua bốn giai đoạn tuần tự cho mỗi tác vụ. Thứ nhất, một backbone DeiT-S/16 đóng băng (tiền huấn luyện trên ImageNet) trích xuất các biểu diễn đặc trưng bậc cao $X_k^{(E)}$ từ ảnh đầu vào, đảm bảo biểu diễn nhất quán trên tất cả các tác vụ mà không cần cập nhật tham số. Thứ hai, các đặc trưng được trích xuất được biến đổi bởi mô-đun RFF sang không gian $D=5000$ chiều thông qua $z(x) = \sqrt{2/D} \cos(\omega^\top x + \beta)$, trong đó $\omega \sim \mathcal{N}(0, \sigma^{-2}I)$ và $\beta \sim \text{Uniform}(0, 2\pi)$, xấp xỉ không gian đặc trưng ẩn của nhân RBF và làm cho dữ liệu xấp xỉ phân tách tuyến tính. Thứ ba, thành phần học phân tích tính toán các tham số tối ưu dạng đóng bằng cách cập nhật đệ quy ma trận tương quan R_k thông qua công thức Sherman-Morrison và suy dẫn ma trận trọng số $\hat{W}_{\text{FCN}}^{(k)}$ theo phương pháp giải tích, hoàn toàn không cần lưu trữ exemplar và không cần tối ưu hóa lặp. Thứ tư, vòng lặp quản lý tác vụ tăng bộ đếm tác vụ và nạp bộ dữ liệu tiếp theo cho đến

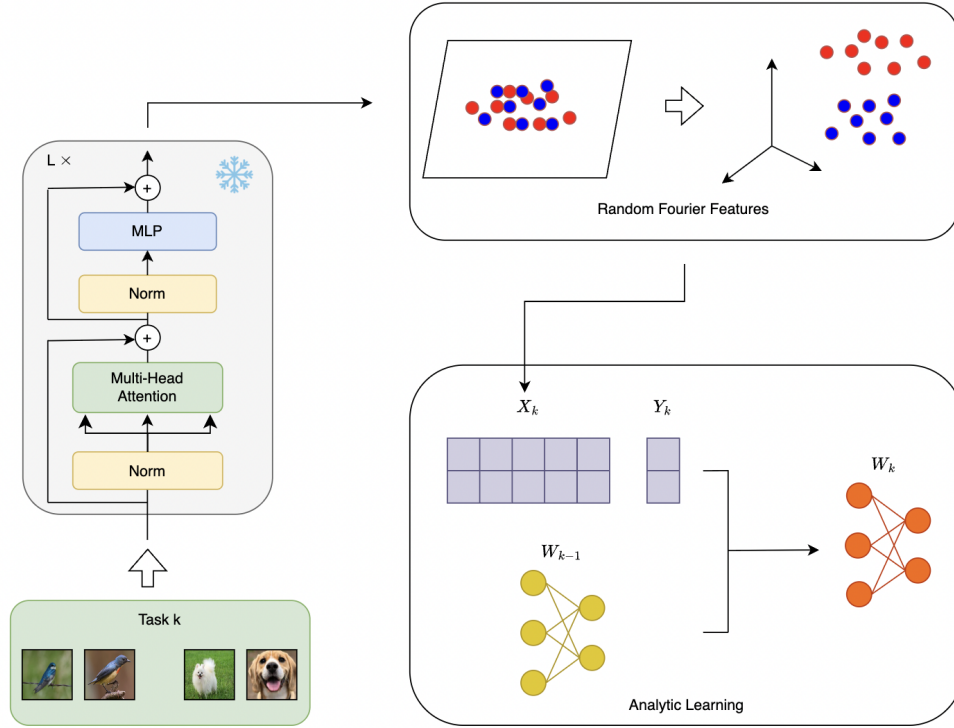


Figure 3.8: Kiến trúc tổng thể của ZeroRFF cho bài toán Học tăng dần theo lớp tổng quát.

khi toàn bộ K tác vụ được xử lý, sau đó trả về ma trận trọng số cuối cùng để suy luận. Thuật toán 1 tóm tắt toàn bộ quy trình.

Các thực nghiệm trên CIFAR-100 theo kịch bản Si-Blurry GCIL (tỷ lệ lớp rời rạc 50%, tỷ lệ mẫu mờ 10%, năm seed) được báo cáo trong Bảng 3.4. ZeroRFF đạt $A_{AUC}=60,48$, $A_{AVG}=58,55$ và $A_{Last}=72,44$, cải thiện so với baseline không cần bộ nhớ tốt nhất là GACL khoảng 2% trên cả ba chỉ số. Đáng chú ý, ZeroRFF còn vượt trội hầu hết các phương pháp dựa trên bộ nhớ sử dụng buffer 500 mẫu, và tiệm cận hiệu suất của MVP-R dựa trên bộ nhớ (kích thước bộ nhớ 2000) trong khi hoàn toàn không cần lưu trữ bất kỳ dữ liệu nào.

Tổng kết

Tóm lại, chương này trình bày hai chiến lược chuyển giao tri thức không cần exemplar mới, chiếm giữ các vị trí khác biệt trong không gian thiết kế của học liên tục. Hướng tiếp cận CL-plugin phù hợp nhất với các kịch bản có sự chồng chéo ngữ nghĩa liên tác vụ phong phú, nơi mà chất lọc đối lập qua các snapshot mô hình đặc thù theo tác vụ mang lại tín hiệu học tập có giá trị. ZeroRFF được ưu tiên khi có sẵn một backbone đóng băng mạnh và khi hiệu quả huấn luyện hoặc quyền riêng tư dữ liệu là yêu cầu hàng đầu, mang lại các bảo đảm dạng đóng và dung lượng bộ nhớ không đổi bất kể số lượng

Algorithm 1 Huấn luyện ZeroRFF

Require: Các tác vụ $\mathcal{D}_1^{\text{train}}, \dots, \mathcal{D}_K^{\text{train}}$; trọng số backbone đóng băng Θ_{backbone}

Ensure: $\hat{W}_{\text{FCN}}^{(K)}$

- 1: **Khởi tạo:** $R_0 \leftarrow \gamma I$, $\hat{W}_{\text{FCN}}^{(0)} \leftarrow 0$
 - 2: **for** $k = 1$ **đến** K **do**
 - 3: $X_k^{(E)} \leftarrow f_{\text{backbone}}(X_k^{\text{train}}, \Theta_{\text{backbone}})$
 - 4: $X_k^{(B)} \leftarrow \text{rff}(X_k^{(E)})$
 - 5: $R_k \leftarrow R_{k-1} - R_{k-1} X_k^{(B)\top} (I + X_k^{(B)} R_{k-1} X_k^{(B)\top})^{-1} X_k^{(B)} R_{k-1}$
 - 6: $\hat{W}^{(k)} \leftarrow \hat{W}^{(k-1)} - R_{k-1} X_k^{(B)\top} X_k^{(B)} \hat{W}^{(k-1)} + R_k X_k^{(B)\top} Y_k^{\text{train}}$
 - 7: **end for**
 - 8: **return** $\hat{W}_{\text{FCN}}^{(K)}$
-

tác vụ. Cả hai phương pháp đều thúc đẩy mục tiêu bảo tồn tri thức đã tích lũy mà không phải gánh chịu các gánh nặng lưu trữ dữ liệu và quyền riêng tư vốn gắn liền với exemplar replay.

Table 3.1: Số câu huấn luyện/kiểm định/kiểm tra cho mỗi tác vụ trên cả bốn bộ dữ liệu.

Nguồn dữ liệu	Tác vụ / miền	Huấn luyện	Kiểm định	Kiểm tra
Liu3Domains	Speaker	352	44	44
	Router	245	31	31
	Computer	283	35	36
HL5Domains	Nokia6610	271	34	34
	Nikon4300	162	20	21
	Creative	677	85	85
	CanonG3	228	29	29
	ApexAD	343	43	43
Ding9Domains	CanonD500	118	15	15
	Canon100	175	22	22
	Diaper	191	24	24
	Hitachi	212	26	27
	Ipod	153	19	20
	Linksys	176	22	23
	MicroMP3	484	61	61
	Nokia6600	362	45	46
	Norton	194	24	25
SemEval14	Restaurant	3452	150	1120
	Laptop	2163	150	638

Table 3.2: Độ chính xác trung bình (Acc.) và Macro-F1 (MF1) trên năm chuỗi ngẫu nhiên gồm 19 tác vụ. In đậm tức là hiệu năng tốt nhất.

Kịch bản	Backbone	Mô hình	Acc.(%)	MF1(%)
Không CL (SDL)	BERT	MTL	91.91	88.11
	BERT	SDL	85.84	76.35
	BERT (Đóng băng)	SDL	78.14	58.13
	Adapter-BERT	SDL	85.96	78.07
	W2V	SDL	77.01	51.89
Học liên tục	BERT (Đóng băng)	NFH	85.51	76.64
	Adapter-BERT	NFH	85.51	76.64
		EWC	86.37	74.52
		OWM	87.02	79.31
		HAT	86.74	78.16
		A-GEM	86.06	78.44
	BERT (Đóng băng)	DER++	84.27	75.08
		KAN	85.49	77.38
		SRK	84.76	78.52
		HAT	86.14	78.52
		BCL	88.29	81.40
	Adapter-BERT	LAMOL	88.91	80.59
		CTR	89.47	83.62
		CTR*	88.21	81.19
		Ours	88.50	82.35

Table 3.3: Kết quả nghiên cứu ablation.

Biến thể mô hình	Acc.(%)	MF1(%)
Ours (mô hình đầy đủ)	88.50	82.35
Không có CL-plugin	85.83	79.75
Không có CSC + CED	86.48	79.32
Không có CSC	87.52	79.93
Không có CED	87.98	80.10

Table 3.4: Kết quả thực nghiệm trên CIFAR-100 theo kịch bản Si-Blurry (%). In đậm = phương pháp không cần bộ nhớ tốt nhất.

Bộ nhớ	Phương pháp	A_{AUC}	A_{AVG}	A_{Last}
2000	EWC++	53.31 ± 1.70	50.95 ± 1.50	52.55 ± 0.71
	ER	56.17 ± 1.84	53.80 ± 1.46	55.60 ± 0.69
	RM	53.22 ± 1.82	52.99 ± 1.65	55.25 ± 0.61
	MVP-R	60.62 ± 1.03	57.58 ± 0.56	64.30 ± 0.29
500	EWC++	48.31 ± 1.81	44.56 ± 0.40	40.52 ± 0.83
	ER	51.59 ± 1.91	48.03 ± 0.80	44.09 ± 0.80
	RM	41.07 ± 1.30	38.10 ± 0.59	32.66 ± 0.34
	MVP-R	56.20 ± 1.44	53.61 ± 0.04	55.35 ± 0.43
0 (không cần bộ nhớ)	LwF	40.71 ± 1.21	38.49 ± 0.56	27.03 ± 2.92
	L2P	42.68 ± 2.70	39.89 ± 0.45	28.59 ± 3.34
	DualPrompt	41.34 ± 2.59	38.59 ± 0.82	22.74 ± 3.40
	MVP	45.07 ± 1.24	44.93 ± 0.54	39.94 ± 0.47
	SLDA	53.00 ± 3.85	50.09 ± 2.77	61.79 ± 3.81
	GACL	57.99 ± 2.46	56.24 ± 3.12	70.31 ± 0.06
	ZeroRFF (Ours)	60.48 ± 4.33	58.55 ± 6.05	72.44 ± 0.13

Chapter 4

Kiến trúc động và phân bổ tham số động

Học liên tục dựa trên kiến trúc giải quyết catastrophic forgetting bằng cách phân bổ động dung lượng mô hình chuyên dụng cho mỗi tác vụ mới, tạo ra các không gian tham số cô lập hoặc chia sẻ một phần nhằm giảm thiểu sự nhiễu loạn giữa các giai đoạn huấn luyện tuần tự. Khác với các phương pháp chính quy hóa vốn ràng buộc quá trình cập nhật hay các phương pháp replay vốn lưu trữ dữ liệu quá khứ, các hướng tiếp cận này trực tiếp thay đổi cấu trúc của mô hình. Chương này trình bày ba đóng góp bổ trợ cho nhau trong kịch bản **Học tăng dần theo lớp (Class-Incremental Learning - CIL)**, được tổ chức theo hai dòng phương pháp luận: backbone ResNet-18 huấn luyện từ đầu cho các benchmark CIL trực tuyến, và backbone ViT-B/16 tiền huấn luyện cho các kịch bản tinh chỉnh hiệu quả tham số (Parameter-Efficient Fine-Tuning - PEFT).

Phần 1: Hướng tiếp cận Multi-Head với Cân bằng Prototype Trực tuyến, Phản hồi Prototype Thích nghi và Cân chỉnh Ấn dựa trên Năng lượng.

Đóng góp thứ nhất mở rộng khung Học Prototype Trực tuyến (OnPro) nhằm khắc phục hai hạn chế: OnPro sử dụng một đầu phân loại chung duy nhất làm lẫn lộn giữa dự đoán trong tác vụ (Within-task Prediction - WP) và dự đoán nhận dạng tác vụ (Task-ID Prediction - TP), đồng thời xử lý đồng đều tất cả các cặp lớp bất kể mức độ dễ bị nhầm lẫn của chúng. Phương pháp được đề xuất, minh họa trong Hình 4.2, giới thiệu hai cơ chế bổ trợ cho nhau. **Cân bằng Prototype Trực tuyến (Online Prototype Equilibrium - OPE)** duy trì các biểu diễn có khả năng phân biệt cao bằng cách đồng thời kéo gần các prototype trực tuyến cùng lớp và đẩy xa các prototype khác lớp, cân bằng giữa việc học

các lớp mới và bảo toàn các lớp đã thấy trước đó thông qua:

$$\mathcal{L}_{\text{OPE}} = \mathcal{L}_{\text{pro}}^{\text{new}}(\mathcal{P}, \hat{\mathcal{P}}) + \mathcal{L}_{\text{pro}}^{\text{seen}}(\mathcal{P}^b, \hat{\mathcal{P}}^b).$$

Phản hồi Prototype Thích nghi (Adaptive Prototypical Feedback — APF) nhận dạng các cặp lớp dễ bị phân loại nhầm bằng cách tính khoảng cách giữa các prototype được lưu trữ thông qua nhân Gaussian đối xứng, sau đó lấy mẫu quá mức các cặp dễ nhầm lẫn trong quá trình memory replay nhằm làm sắc nét ranh giới quyết định giữa chúng. Chiến lược lấy mẫu hai giai đoạn chọn $\alpha \cdot m$ mẫu tỷ lệ với xác suất nhầm lẫn và lấy mẫu ngẫu nhiên phần còn lại để duy trì cân bằng tổng thể.

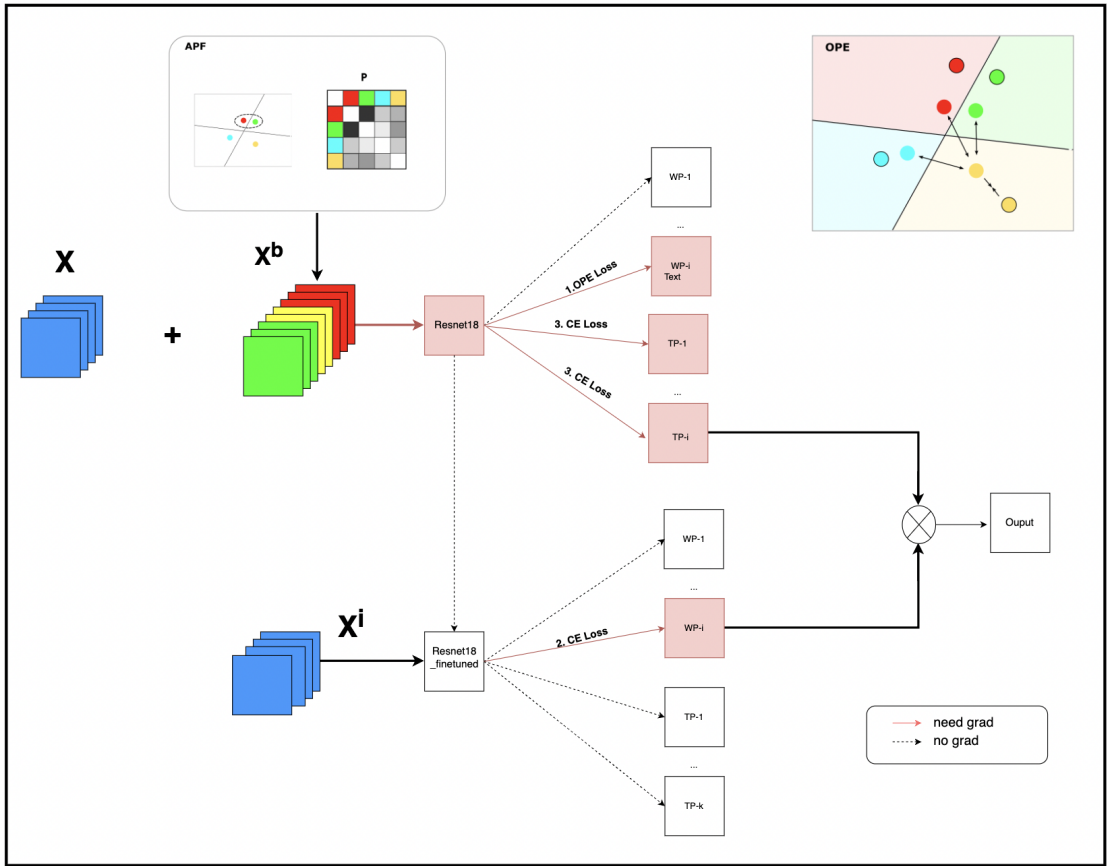


Figure 4.2: Kiến trúc tổng thể của mô hình multi-head với OPE, APF và hai đầu phân loại (WP và OOD).

Lớp phân loại được tái cấu trúc thành hai đầu chuyên biệt: một **đầu OOD** ước lượng xác suất tác vụ bằng cách phân biệt các mẫu thuộc phân phối của tác vụ hiện tại với các mẫu ngoài phân phối trong bộ đệm bộ nhớ, và một **đầu WP** để dự đoán chi tiết trong nội bộ tác vụ. Dự đoán CIL cuối cùng được phân rã theo công thức $P(X_{k,j}|x) = P(X_{k,j}|x, k) P(X_k|x)$, kết hợp đầu ra của cả hai đầu. Quá trình huấn luyện diễn ra theo ba bước: (1) huấn luyện chung bộ trích xuất đặc trưng và đầu OOD; (2) tinh

chỉnh đầu WP trong khi giữ nguyên bộ trích xuất đặc trưng; (3) tinh chỉnh tất cả các đầu OOD trước đó trên dữ liệu replay nhằm giảm thiểu thiên lệch đối với các tác vụ đầu tiên.

Mô-đun **Căn chỉnh Ẩn dựa trên Năng lượng (Energy-based Latent Alignment — ELI)** được tích hợp thêm nhằm chống lại hiện tượng trôi dạt biểu diễn. Sau mỗi tác vụ, một Mô hình dựa trên Năng lượng (Energy-Based Model — EBM) được huấn luyện để gán năng lượng thấp cho các biểu diễn ẩn từ snapshot bộ mã hóa trước đó và năng lượng cao cho các biểu diễn từ bộ mã hóa hiện tại, sử dụng Langevin dynamics để lấy mẫu. Tại thời điểm suy luận, ELI lặp đi lặp lại dịch chuyển vector ẩn của từng mẫu kiểm tra dọc theo gradient năng lượng trở về đa tạp đặc trưng của tác vụ trước, hiệu chỉnh sự dịch chuyển biểu diễn mà không cần lưu trữ thêm bất kỳ dữ liệu nào.

Bảng 4.1 và 4.2 báo cáo độ chính xác trung bình và mức độ quên trên CIFAR-10 và CIFAR-100. Phương pháp đề xuất (ModifiedOnPro) nhất quán vượt trội hơn OnPro về độ chính xác ở tất cả các kích thước bộ nhớ đệm lên tới 11,2 điểm trên CIFAR-10 ($\mathcal{M}=100$, từ 52,4% lên 63,6%) và lên tới 6,4 điểm trên CIFAR-100 ($\mathcal{M}=500$, từ 20,7% lên 27,1%). Trên benchmark hình ảnh y tế CCH5000 (Bảng 4.5), phương pháp đạt điểm quên lãng tốt nhất là 2,9 trong kịch bản hai lớp mỗi tác vụ, cải thiện 27,5% so với LO2LN.

Table 4.1: Độ chính xác trung bình (%) với các kích thước bộ nhớ đệm $\mathcal{M}$ khác nhau. In đậm = tốt nhất.

	CIFAR-10			CIFAR-100		
	$\mathcal{M}=100$	$\mathcal{M}=200$	$\mathcal{M}=500$	$\mathcal{M}=500$	$\mathcal{M}=1000$	$\mathcal{M}=2000$
Phương pháp						
iCaRL	31.0	33.9	42.0	12.8	16.5	17.6
DER++	31.5	39.7	50.9	16.0	21.4	23.9
SCR	40.2	48.5	59.1	19.3	26.5	32.7
OCM	47.5	59.6	70.1	19.7	27.5	34.4
OnPro	52.4	63.0	67.7	20.7	27.2	34.9
Ours	63.6	69.3	72.8	27.1	30.3	36.3

Phần 2: DAKTA — Bộ phân loại Kolmogorov-Arnold có định hướng cho Số học tác vụ.

Đóng góp thứ hai, DAKTA (Hình 4.6), được xây dựng dựa trên khung Số học tác vụ tăng dần với LoRA (Incremental Task Arithmetic with LoRA — ITA-LoRA) và giải

Table 4.2: Mức độ quên lãng trung bình (%) với các kích thước buffer bộ nhớ $\mathcal{M}$ khác nhau. Càng thấp càng tốt.

	CIFAR-10			CIFAR-100		
Phương pháp	$\mathcal{M}=100$	$\mathcal{M}=200$	$\mathcal{M}=500$	$\mathcal{M}=500$	$\mathcal{M}=1000$	$\mathcal{M}=2000$
iCaRL	52.7	49.3	38.3	16.5	11.2	10.4
DER++	57.8	46.7	33.6	41.0	34.8	33.2
SCR	43.2	35.5	24.1	29.3	20.4	11.5
OCM	35.5	23.9	13.5	18.3	15.2	10.8
OnPro	24.2	16.5	13.7	17.8	14.9	10.0
Ours	23.2	18.3	13.5	19.2	14.5	11.2

quyết hạn chế then chốt của khung này: đầu phân loại tuyến tính chuẩn áp đặt ranh giới quyết định đồng nhất trên tất cả các tác vụ. DAKTA thay thế đầu tuyến tính bằng một **Bộ phân loại Kolmogorov-Arnold (Kolmogorov-Arnold Classifier - KAC)**, tính toán các kích hoạt lớp dưới dạng tổng có trọng số của các Hàm cơ sở hướng tâm Gaussian (Radial Basis Functions - RBFs) có thể học được, cung cấp các ranh giới nhạy cảm cục bộ và có tính chọn lọc kênh giúp giảm đáng kể nhiễu loạn giữa các tác vụ. Cơ chế **Hợp nhất vector có định hướng** tiếp tục kết hợp logit dựa trên RBF với số hạng độ tương đồng cosine:

$$l_c(x) = l_c^{\text{init}}(x) + \gamma_c \cdot \cos(x, W_c),$$

nắm bắt đồng thời các mẫu tương đồng cục bộ và sự can chỉnh định hướng toàn cục để phân biệt lớp một cách bền vững hơn. **Chính quy hóa Ma trận thông tin Fisher (Fisher Information Matrix - FIM) bậc hai** được áp dụng cho cả backbone ViT lẫn đầu KAC, phạt các độ lệch của các tham số quan trọng so với giá trị tiền huấn luyện:

$$\mathcal{L} = \mathcal{L}_{\text{CE}} + \mathcal{L}_{\text{reg}}^{\text{backbone}} + \mathcal{L}_{\text{reg}}^{\text{cls}}.$$

Sau khi huấn luyện mỗi tác vụ, các thay đổi tham số được lưu trữ dưới dạng vector tác vụ và tổng hợp tại thời điểm suy luận bằng phép lấy trung bình, kế thừa các bảo đảm ổn định của khung số học tác vụ.

Bảng 4.6 báo cáo độ chính xác cuối cùng trên sáu benchmark. DAKTA đạt kết quả tốt nhất trong năm trong sáu bộ dữ liệu, vượt ITA-LoRA tới 2,04 điểm (CIFAR-100) và 2,04 điểm trên CUB-200, benchmark nhận dạng chi tiết nơi độ nhạy cảm cục bộ của KAC

Table 4.5: Kết quả so sánh trên CCH5000 (hình ảnh y tế) với một và hai lớp mỗi tác vụ. In đậm = tốt nhất.

Phương pháp	1 lớp/tác vụ		2 lớp/tác vụ	
	Acc	Fgt	Acc	Fgt
iCaRL	91.1	9.0	93.0	6.8
UCIR	92.0	5.5	93.9	4.4
PODNet	89.2	6.0	92.0	5.2
DER	91.0	5.6	93.0	6.4
LO2LN	94.5	3.9	94.6	4.0
OnPro	89.5	3.6	92.8	3.8
Ours	91.6	4.6	94.6	2.9

phát huy tác dụng rõ rệt nhất. So sánh về mức độ quên tiếp tục khẳng định sự vượt trội của DAKTA: phương pháp đạt điểm quên thấp nhất trên cả sáu benchmark, giảm mức độ quên tới 0,94 điểm so với ITA-LoRA trên CIFAR-100. Nghiên cứu ablation (Bảng 4.7) xác nhận rằng cả đầu KAC lẫn chính quy hóa FIM đều là thành phần thiết yếu, với mô hình đầy đủ nhất quán vượt trội hơn các biến thể loại bỏ một trong hai thành phần.

Phần 3: LoRA Trực giao Thừa thốt cho Học liên tục (SO-LoRA).

Đóng góp thứ ba, SO-LoRA (Hình 4.8), giải quyết sự mâu thuẫn căn bản giữa cô lập adapter theo từng tác vụ (ổn định nhưng tốn $\mathcal{O}(T)$ bộ nhớ) và adapter dùng chung (bộ nhớ không đổi nhưng dễ bị nhiễu loạn). SO-LoRA tái khái niệm hóa bài toán theo hướng hình học: phương pháp duy trì một cơ sở trực chuẩn cố định $A \in \mathbb{R}^{d \times r}$ (khởi tạo một lần thông qua phân tích QR và đóng băng) và hạn chế mỗi tác vụ chỉ được học một ma trận hệ số thừa thốt $B^{(t)} \in \mathbb{R}^{r \times k}$:

$$W^{(t)} = W_0 + AB^{(t)}.$$

Vì $A^\top A = I_r$, nhiễu loạn giữa các tác vụ i và j được chứng minh là bị chặn bởi $|\langle \Delta W_i, \Delta W_j \rangle_F| \leq \|\Delta B_i\|_{2,1} \|\Delta B_j\|_{2,1}$, và được kiểm soát trực tiếp bằng cách tối thiểu hóa chuẩn $\ell_{2,1}$ (group lasso) trên $B^{(t)}$, buộc mỗi tác vụ chỉ kích hoạt một tập con thừa thốt các chiều cơ sở. Một **mặt nạ mềm nhận biết gradient** $M \in \mathbb{R}^r$ tiếp tục bảo vệ các chiều rank có tầm quan trọng tích lũy: sau mỗi tác vụ, chuẩn gradient hàng được tổng

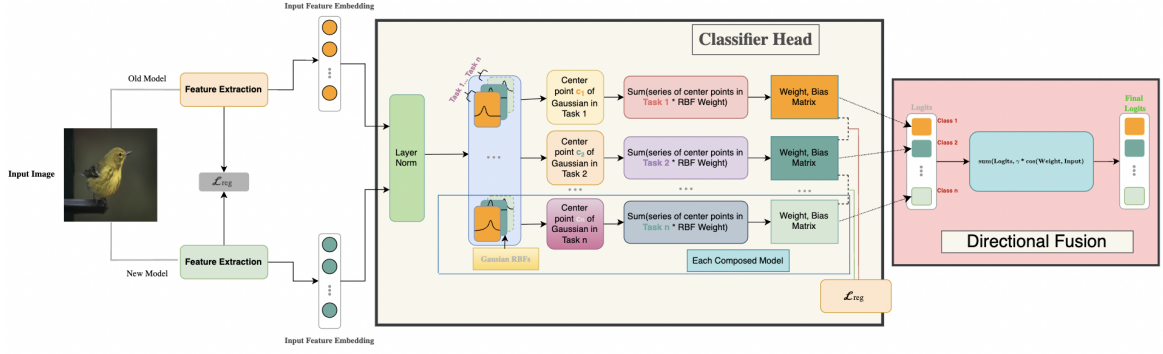


Figure 4.6: Pipeline của khung DAKTA tích hợp backbone ViT+LoRA, Bộ phân loại Kolmogorov-Arnold (KAC), Số học tác vụ tăng dần (ITA) và chính quy hóa dựa trên Fisher.

hợp thông qua trung bình động theo hàm mũ, và trong quá trình huấn luyện tác vụ tiếp theo, gradient của các hàng được bảo vệ mạnh nhưng hiện không liên quan bị triệt tiêu theo hệ số tỷ lệ $S_j = 1 - M_{t-1,j}(1 - \phi_j)$.

Bảng 4.9 và 4.10 tóm tắt các kết quả chính. Trên ImageNet-R, SO-LoRA xếp hạng nhất ở tất cả các mức độ chi tiết tác vụ (5/10/20 tác vụ), với mức giảm độ chính xác rất nhỏ từ 79,78% xuống 78,15% khi chuỗi tăng từ 5 lên 20 tác vụ, trong khi O-LoRA suy giảm mạnh từ 72,76% xuống 60,98%. Trên ImageNet-A (dịch chuyển miền nghiêm trọng), SO-LoRA đạt độ chính xác cuối cùng 60,27%, cao hơn gần 5 điểm so với SD-LoRA, nêu bật tầm quan trọng của mặt nạ nhận biết gradient khi dữ liệu hiện tại khác biệt xa so với phân phối tiền huấn luyện. Trên CUB-200 (nhận dạng chi tiết), SO-LoRA đạt độ chính xác cao nhất (81,58%), xác nhận rằng các tổ hợp trực giao thưa thớt nắm bắt được các đặc điểm tinh tế mà không khuếch đại nhiễu loạn. Hình 4.9 trực quan hóa quỹ đạo độ chính xác theo các độ dài chuỗi khác nhau.

Nghiên cứu ablation (Bảng 4.11) xác nhận rằng mặt nạ nhận biết gradient mang lại những cải thiện nhất quán nhất, đặc biệt trong điều kiện dịch chuyển miền (ImageNet-A) và chuỗi dài (20 tác vụ trên ImageNet-R), trong khi độ thưa thớt nhóm bổ sung thêm lợi ích đặc biệt trong các kịch bản ngoài miền.

Tổng kết

Tóm lại, chương này trình bày ba chiến lược dựa trên kiến trúc có tính hỗ trợ cho nhau, cùng nhau giải quyết các hạn chế then chốt của các phương pháp học liên tục hiện có. Hướng tiếp cận multi-head OPE/APF/ELI cải thiện cả biểu diễn đặc trưng lẫn phân

Table 4.6: So sánh độ chính xác cuối cùng (%) với các phương pháp tiên tiến nhất. In đậm = tốt nhất.

Mô hình	C-100	CUB	Caltech	MIT	RESISC	CropDis.
CODA	86.48	78.54	90.57	77.73	69.50	74.65
SEED	83.39	85.35	90.04	86.34	74.81	92.77
InfLoRA	87.17	79.14	90.53	79.14	79.92	89.05
DER++	83.02	76.10	86.11	68.88	67.23	98.77
ITA-LoRA	89.96	85.55	92.65	86.60	82.00	95.85
ITA-(IA) ³	90.66	85.67	92.67	84.74	83.73	95.41
DAKTA (Ours)	91.21	87.59	93.61	87.97	85.82	97.75

Table 4.7: Nghiên cứu ablation về chính quy hóa bậc hai và KAC trong DAKTA. In đậm = tốt nhất.

Biến thể	C-100	CUB	Cal.	MIT	RES.	Crop.
Không có Reg.	85.24	80.57	89.10	82.25	79.38	90.84
Không có KAC	89.91	84.12	92.35	85.24	82.25	94.82
Mô hình đầy đủ	91.21	87.59	93.61	87.97	85.82	97.75

loại trong CIL trực tuyến bằng cách tách biệt dự đoán trong tác vụ khỏi dự đoán nhận dạng tác vụ và hiệu chỉnh sự trôi dạt ẩn. DAKTA thay thế đầu tuyến tính thông thường bằng bộ phân loại Kolmogorov-Arnold nhạy cảm cục bộ, mở rộng chính quy hóa dựa trên Fisher sang lớp phân loại và đạt độ chính xác vượt trội cùng mức độ quên thấp hơn trên các benchmark nhận dạng chi tiết đa dạng. SO-LoRA cung cấp một giải pháp có cơ sở lý thuyết vững chắc và dung lượng tham số không đổi, cô lập hình học các biểu diễn tác vụ trong một cơ sở trực giao cố định và bảo vệ động các chiều quan trọng thông qua mặt nạ nhận biết gradient, thể hiện khả năng mở rộng mạnh mẽ trên các chuỗi tác vụ dài và các điều kiện dịch chuyển phân phối nghiêm trọng.

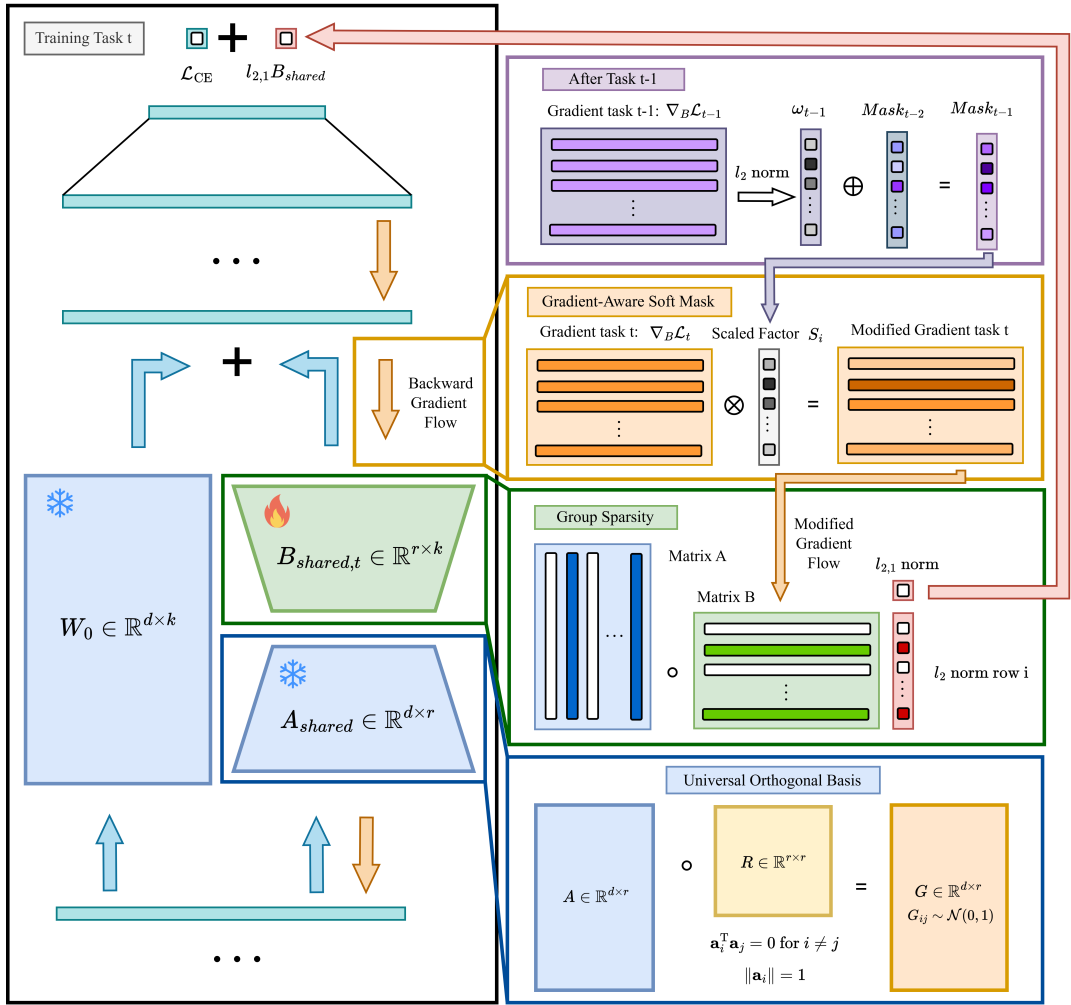


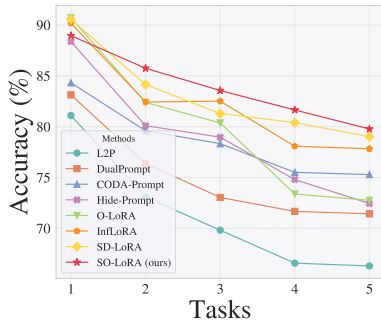
Figure 4.8: SO-LoRA biểu diễn sự thích nghi dưới dạng các hệ số thưa thớt $B^{(t)}$ trên cơ sở trực chuẩn đồng bằng A . Độ thưa thớt nhóm giới hạn sự chồng chéo; mặt nạ nhận biết gradient M bảo vệ các chiều quan trọng về mặt lịch sử.

Table 4.9: Hiệu suất trên ImageNet-R với các mức độ chi tiết tác vụ khác nhau. AAA = Độ chính xác trung bình. In đậm = tốt nhất.

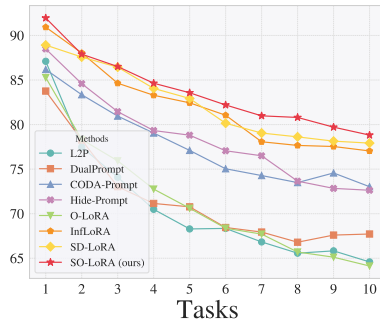
Phương pháp	IN-R (5)		IN-R (10)		IN-R (20)	
	AAA↑	Acc↑	AAA↑	Acc↑	AAA↑	Acc↑
CODA	78.68	75.29	77.11	73.03	73.71	68.53
HiDe	77.42	72.45	76.84	72.63	75.88	71.56
O-LoRA	76.71	72.76	71.30	64.13	68.04	60.98
InfLoRA	81.63	77.83	80.24	77.04	76.78	70.94
SD-LoRA	82.85	79.02	81.05	77.92	80.62	75.43
SO-LoRA	83.94	79.78	83.71	78.83	83.18	78.15

Table 4.10: Hiệu suất trên ImageNet-A, CIFAR-100 và CUB-200 (10 tác vụ). In đậm = tốt nhất.

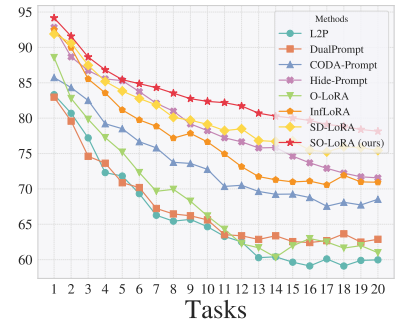
Phương pháp	IN-A		C100		CUB	
	AAA↑	Acc↑	AAA↑	Acc↑	AAA↑	Acc↑
CODA	56.30	45.94	91.37	86.82	82.36	72.41
HiDe	55.72	44.67	91.17	86.53	81.71	74.15
O-LoRA	54.17	44.46	89.96	82.51	79.66	70.12
InfLoRA	61.63	51.02	91.44	86.09	81.43	71.13
SD-LoRA	63.24	55.35	91.71	87.28	85.15	77.38
SO-LoRA	67.89	60.27	90.88	86.61	87.71	81.58



(5 tác vụ)



(10 tác vụ)



(20 tác vụ)

Figure 4.9: Độ chính xác trung bình trên ImageNet-R theo các độ dài chuỗi khác nhau (5, 10, 20 tác vụ).

Table 4.11: Nghiên cứu ablation trên các benchmark 10 tác vụ. Tác động của độ thưa thớt nhóm và mặt nạ nhận biết gradient. In đậm = tốt nhất.

Biến thể	ImageNet-A		ImageNet-R (10)	
	AAA	Acc	AAA	Acc
SO-LoRA (đầy đủ)	67.89	60.27	83.71	78.83
Không có độ thưa thớt	66.86	59.30	83.52	78.77
Không có mặt nạ	66.31	58.15	83.40	78.37

Kết luận và Hướng phát triển

Luận án này đã giải quyết bài toán catastrophic forgetting trong học liên tục thông qua một tiến trình ba giai đoạn, trong đó mỗi giai đoạn được thúc đẩy bởi những hạn chế thực tiễn và lý thuyết nảy sinh ở các tầng bậc kế tiếp của bài toán.

Giai đoạn thứ nhất (**Chương 2**) giới thiệu cơ chế replay prototype bảo toàn quyền riêng tư. Thay vì lưu trữ các mẫu huấn luyện thực tế, phân phối của các lớp trong quá khứ được xấp xỉ thông qua tăng cường prototype với nhiễu Gaussian, bảo toàn cấu trúc phân phối cần thiết cho replay đối lập. Mô-đun Căn chỉnh Ấn dựa trên Năng lượng (ELI) hỗ trợ cho cơ chế này bằng cách phát hiện và hiệu chỉnh sự trôi dạt biểu diễn sau mỗi tác vụ thông qua Langevin dynamics. Các thực nghiệm trên FewRel và TACRED chứng minh rằng sự kết hợp này đạt hiệu suất ngang bằng hoặc vượt trội so với các phương pháp lưu trữ dữ liệu thực, khẳng định rằng replay bảo toàn quyền riêng tư vừa khả thi vừa hiệu quả.

Giai đoạn thứ hai (**Chương 3**) loại bỏ hoàn toàn việc lưu trữ dữ liệu. Khung **CL-plugin** kết hợp Mô-đun chia sẻ tri thức và Mô-đun đặc thù tác vụ với mặt nạ ở mức nơ-ron, được tăng cường bởi các hàm mất mát Chắt lọc tập hợp đối lập và Học đối lập có giám sát, đạt Macro-F1 tiên tiến nhất trên benchmark phân loại cảm xúc theo khía cạnh gồm 19 tác vụ. **ZeroRFF** mở rộng nguyên lý này sang các bài toán thị giác class-incremental bằng cách ánh xạ các đặc trưng từ backbone đóng băng qua random Fourier features và tính toán các bộ phân loại phân tích dạng đóng, đạt kết quả cạnh tranh trong kịch bản Si-Blurry mà không cần lưu trữ bất kỳ exemplar nào. Cả hai phương pháp cùng nhau khẳng định rằng học liên tục không cần exemplar là khả thi trên cả hai lĩnh vực NLP và thị giác máy tính.

Giai đoạn thứ ba (**Chương 4**) giải quyết ràng buộc dung lượng vốn là giới hạn cuối cùng của cả hai hướng tiếp cận replay lẫn chắt lọc. Bốn phương pháp được đóng góp. Khung **multi-head OPE/APF/ELI** tách biệt dự đoán trong tác vụ khỏi dự đoán nhận dạng tác vụ thông qua hai đầu WP và OOD riêng biệt, đạt điểm quên lãng tốt nhất trên benchmark hình ảnh y tế CCH5000. **DAKTA** thay thế bộ phân loại tuyến tính thông

thường bằng Bộ phân loại Kolmogorov-Arnold được hỗ trợ bởi chính quy hóa Ma trận thông tin Fisher trên cả backbone lẫn đầu phân loại, đạt mức cải thiện độ chính xác nhất quán so với ITA-LoRA trên sáu benchmark. **SO-LoRA** cố định một cơ sở adapter trực giao phổ quát và hạn chế mỗi tác vụ chỉ học các tổ hợp hệ số thừa thớt, chặn nhiễu loạn giữa các tác vụ bằng tích của các chuẩn độ thừa thớt và duy trì dung lượng tham số không đổi bất kể độ dài chuỗi, đạt độ chính xác cao nhất trên ImageNet-A, CUB-200 và kịch bản 20 tác vụ trên ImageNet-R.

Nhìn nhận một cách tổng thể, ba giai đoạn hợp thành một câu trả lời có hệ thống và nhạy cảm với ngữ cảnh cho bài toán catastrophic forgetting: replay prototype khi cho phép lưu trữ hạn chế; chất lọc tri thức có chọn lọc khi việc lưu giữ dữ liệu là không khả thi; và các chiến lược kiến trúc thích nghi khi sự đa dạng của tác vụ hoặc độ dài chuỗi khiến ngân sách tham số cố định trở thành ràng buộc cứng. Cơ chế ELI, được phát triển ở Chương 2, được chuyển giao trực tiếp sang kiến trúc multi-head của Chương 4, và tính ổn định dạng đóng của ZeroRFF có chung nền tảng khái niệm với SO-LoRA, điều này cho thấy các giai đoạn không tồn tại độc lập mà có tính tương hỗ và củng cố lẫn nhau.

Tuy nhiên, một số hạn chế vẫn còn tồn tại. Các phương pháp đều giả định ranh giới tác vụ được biết trước, điều kiện này bị vi phạm trong các kịch bản không có ranh giới tác vụ và ranh giới mờ. Khả năng mở rộng lên các chuỗi hàng trăm tác vụ chưa được đánh giá một cách có hệ thống. Phạm vi thực nghiệm còn giới hạn ở phân loại văn bản và ảnh, và sự không nhất quán về backbone giữa hai dòng ResNet-18 và ViT-B/16 trong Chương 4 hạn chế khả năng so sánh trực tiếp giữa các phương pháp. Hướng nghiên cứu tương lai nên theo đuổi một meta-khung thích nghi theo ngữ cảnh, có khả năng tự động lựa chọn giữa replay, chất lọc và mở rộng kiến trúc dựa trên bộ nhớ khả dụng, độ tương đồng giữa các tác vụ và ngân sách tính toán. Mở rộng số học tác vụ dựa trên LoRA sang các mô hình ngôn ngữ lớn, học liên tục không có ranh giới tác vụ thông qua phát hiện ranh giới điều khiển bởi OPE, và các kịch bản đa phương thức nơi các phương thức khác nhau tự nhiên chiếm giữ các không gian con tham số riêng biệt là những hướng phát triển đặc biệt đầy triển vọng.

Catastrophic forgetting vẫn là một trở ngại căn bản trên con đường hướng tới trí tuệ nhân tạo thực sự thích nghi. Ba đóng góp của luận án này cung cấp những bước tiến cụ thể và có cơ sở nguyên tắc vững chắc hướng tới một hệ thống học tập có khả năng tích lũy tri thức suốt vòng đời hoạt động, mở rộng được trên các chuỗi tác vụ dài, và vận hành trong các ràng buộc tài nguyên thực tiễn.

Danh sách công bố

- [P1] **Quynh-Trang Pham Thi**, Anh-Cuong Pham, Ngoc-Huyen Ngo, and Duc-Trong Le. "Memory-Based Method using Prototype Augmentation for Continual Relation Extraction." In 2022 RIVF International Conference on Computing and Communication Technologies (RIVF), pp. 1-6. IEEE, 2022. Scopus.
- [P2] Thanh-Hai Dang, **Quynh-Trang Pham Thi**, Duc-Trong Le, and Tri-Thanh Nguyen. Contrastive Learning for Boosting Knowledge Transfer in Task-Incremental Continual Learning of Aspect Sentiment Classification Tasks. VNU Journal of Science: Computer Science and Communication Engineering 40, no. 1 (2024).
- [P3] Duc-Hung Nguyen, Tri-Thanh Nguyen, Thanh-Hai Dang and **Quynh-Trang Pham Thi**. ZeroRFF: A Random Fourier Features and Analytic Learning Method for Generalized Class Incremental Learning. The 21st International Conference on Advanced Data Mining and Applications 2025. Scopus Rank B.
- [P4] **Quynh-Trang Pham Thi**, Duc-Hung Nguyen, Thanh Hai Dang, Duc-Trong Le, Tri-Thanh Nguyen, and Quang-Thuy Ha. "Towards Robust Continual Learning: A Multi-Head Approach with Online Prototype Equilibrium and Adaptive Prototypical Feedback." In Asian Conference on Intelligent Information and Database Systems, pp. 275-286. Singapore: Springer Nature Singapore, 2024. Scopus Rank B.
- [P5] **Quynh-Trang Pham Thi**, Duc-Hung Nguyen, Anh-Duc Nguyen-Tran, Hung-Dung Nguyen, Tri-Thanh Nguyen, Quang-Thuy Ha and Thanh Hai Dang. "Addressing Catastrophic Forgetting in Continual Learning using Multihead Approach and Latent Alignment via Energy-Based Model." Vietnam journal of computer science. (Accepted). WoS Q3.
- [P6] **Quynh-Trang Pham Thi**, Van-Truong Le, Thanh-Huyen Pham, Thanh-Nga Hoang Thi and Duc-Trong Le. DAKTA: Directional Kolmogorov-Arnold Classifier for Task Arithmetic in Continual Learning. SOICT 2025. Scopus.

[P7] **Quynh-Trang Pham Thi**, Quang-Dat Mac, Thanh-Ha Nguyen and Duc-Trong Le.
SO-LoRA: Sparse Orthogonal LoRA for Parameter-Efficient Continual Learning.
In 2026 IEEE International Conference on Multimedia and Expo (ICME 2026).
Scopus Rank A.

Danh sách có 7 công bố.