

**NHẬN DIỆN GIAN LẬN BÁO CÁO TÀI CHÍNH:
TỔNG QUAN XU HƯỚNG NGHIÊN CỨU VÀ HƯỚNG ĐI MỚI**
DETECTING FINANCIAL STATEMENT FRAUD: A COMPREHENSIVE REVIEW OF
RESEARCH TRENDS AND FUTURE DIRECTIONS

Ngày nhận bài: 04/03/2025

Ngày nhận bản sửa: 25/05/2025

Ngày chấp nhận đăng: 14/07/2025

Tiêu Thị Thanh Hoa, Hoàng Thanh Hiền[✉]

TÓM TẮT

Bài báo tổng hợp các nghiên cứu trong và ngoài nước về nhận diện GLBCTC của doanh nghiệp, nhằm cung cấp cái nhìn tổng quát về kết quả đạt được. Thông qua so sánh giữa các mô hình dự báo sử dụng chỉ số tài chính, thông tin phi tài chính và kỹ thuật học máy, nghiên cứu nhận diện hiệu quả vượt trội của học máy so với các phương pháp thống kê truyền thống. Đồng thời, bài báo chỉ ra các khoảng trống nghiên cứu tại Việt Nam, như hạn chế tích hợp thông tin phi tài chính, ứng dụng chủ yếu phương pháp truyền thống và phạm vi nghiên cứu giới hạn đối với doanh nghiệp niêm yết trên một sàn giao dịch. Những phát hiện này góp phần củng cố nền tảng lý thuyết và đưa ra hàm ý chính sách cho các bên liên quan.

Từ khóa: Chỉ số tài chính; Gian lận báo cáo tài chính; Thông tin phi tài chính; Học máy.

ABSTRACT

This paper surveys both domestic and international research on detecting financial statement fraud in firms, presenting a concise overview of the main findings. By comparing models that use financial ratios, non-financial information, and machine learning techniques, the study demonstrates that machine learning methods outperform traditional statistical approaches. It also identifies research gaps in Vietnam, such as the limited integration of non-financial data, the predominant reliance on traditional methods, and the narrow focus on companies listed on a single exchange. These findings contribute to the theoretical foundation for developing more effective fraud detection models and offer policy implications for relevant stakeholders.

Keywords: Fraudulent financial statements; Financial indicators; Non-financial information; Machine learning.

1. Giới thiệu

Gian lận báo cáo tài chính (GLBCTC) là một trong những rủi ro nghiêm trọng nhất đối với sự minh bạch và ổn định của thị trường tài chính. Hậu quả của các vụ gian lận tài chính không chỉ giới hạn trong phạm vi doanh nghiệp mà còn gây ra những tác động lan tỏa, làm lung lay niềm tin của nhà đầu tư và ảnh hưởng tiêu

cực đến toàn bộ nền kinh tế. Những vụ bê bối tài chính lớn như Enron, Worldcom, Global Crossing, Adelphia, Lehman Brothers đã để lại bài học sâu sắc về hậu quả nghiêm trọng của việc thao túng báo cáo tài chính, dẫn đến phá sản doanh nghiệp và thiệt hại hàng tỷ USD. Tại Việt Nam, các vụ gian lận tài chính như Công ty Bánh kẹo Biên Hòa (2002), Công ty Đồ hộp

Hạ Long (2002), Công ty Bông Bạch Tuyết (2004-2008), Công ty Dược Viễn Đông (2009), Công ty Gỗ Trường Thành (2016) và gần đây là Tập đoàn Tân Hoàng Minh (2024), Tập đoàn Vạn Thịnh Phát (2024), tiếp tục đặt ra những thách thức đối với hệ thống kiểm soát tài chính và quy trình giám sát doanh nghiệp.

Theo Hiệp hội Các nhà điều tra gian lận (ACFE, 2016-2024), mặc dù GLBCTC không phải là loại gian lận phổ biến nhất so với chiếm đoạt tài sản hay tham nhũng, nhưng lại gây thiệt hại lớn nhất. Nghiên cứu của Kaminski và cộng sự (2004) chỉ ra rằng hầu hết các vụ gian lận chỉ được phát hiện sau nhiều năm, khi hậu quả đã trở nên nghiêm trọng. Báo cáo của ACFE (2022) cũng nhấn mạnh rằng gian lận tài chính trung bình mất 2 năm để phát hiện, và thời gian phát hiện càng lâu thì tổn thất càng lớn. Đáng lo ngại hơn, khi gian lận bị phát hiện, các bằng chứng thường đã bị xóa bỏ hoặc bóp méo, khiến việc xử lý trở nên khó khăn hơn. Điều này đặt ra một câu hỏi cấp thiết: Liệu có thể phát hiện GLBCTC sớm hơn để cung cấp cảnh báo kịp thời cho các bên liên quan hay không?

Trong bối cảnh đó, nghiên cứu về mô hình nhận diện GLBCTC đã và đang trở thành một xu hướng quan trọng trong giới học thuật. Các phương pháp truyền thống dựa trên chỉ số tài chính như Altman Z-score, Beneish M-score đã được sử dụng rộng rãi để phát hiện gian lận. Tuy nhiên, với sự phát triển của công nghệ, các phương pháp mới ứng dụng học máy (machine learning), trí tuệ nhân tạo (AI), và dữ liệu phi tài chính đã chứng minh được hiệu quả vượt trội. Dù vậy, tại Việt Nam, hầu hết các nghiên cứu vẫn chủ yếu áp dụng các mô hình truyền thống, chưa khai thác đầy đủ tiềm năng của các phương pháp tiên tiến.

Bài nghiên cứu này nhằm hệ thống hóa các phương pháp nhận diện GLBCTC, đồng thời xác định những khoảng trống nghiên cứu còn tồn tại, đặc biệt trong bối cảnh Việt Nam. Để

đạt được mục tiêu này, nghiên cứu tập trung trả lời câu hỏi: “Những phương pháp và mô hình nào đã được sử dụng để nhận diện gian lận báo cáo tài chính, và khoảng trống nào còn tồn tại trong nghiên cứu hiện nay, đặc biệt trong bối cảnh Việt Nam?”

Bằng cách tổng hợp các nghiên cứu trong và ngoài nước, nghiên cứu này không chỉ làm rõ xu hướng phát triển trong lĩnh vực phát hiện gian lận tài chính mà còn đưa ra đề xuất về những hướng đi mới, góp phần cải thiện hiệu quả dự báo và hỗ trợ các nhà quản lý, kiểm toán viên trong việc phát hiện sớm các hành vi gian lận. Để đạt được mục tiêu nghiên cứu, bài báo không chỉ dừng ở việc tập trung vào phân tích các phương pháp khai thác dữ liệu được sử dụng cho phát hiện GLBCTC như nghiên cứu tổng quan của Gupta và Mehta (2024). Mà còn xem xét đến các biến/thuộc tính bao gồm chỉ số tài chính và thông tin phi tài chính được đưa vào mô hình dự đoán trong các nghiên cứu như thế nào. Đây là điểm tương đồng với nghiên cứu tổng quan của Shahana và cộng sự (2023). Ngoài ra, theo khuyến nghị của Shahana và cộng sự (2023), các nghiên cứu trong tương lai nên mở rộng tổng quan đối với cả những nghiên cứu tại các quốc gia đang phát triển. Điều này nhằm cung cấp bằng chứng về tình hình nghiên cứu cũng như tính hiệu quả của các phương pháp tự động (học máy, trí tuệ nhân tạo) ở các quốc gia này so với nghiên cứu của các nước phát triển. Như vậy, việc thực hiện tổng quan các nghiên cứu bao gồm cả Việt Nam - một nước đang phát triển có thể góp phần làm phong phú tài liệu về chủ đề nghiên cứu nhận diện GLBCTC. Đây là một số đóng góp cơ bản của nghiên cứu này.

2. Tổng quan nghiên cứu gian lận trong báo cáo tài chính của các doanh nghiệp

Nghiên cứu tổng hợp các tài liệu trong và ngoài nước về nhận diện GLBCTC dựa trên tiêu chí chọn lọc chặt chẽ. Việc tìm kiếm được thực hiện theo tiêu đề, tóm tắt và từ khóa được

sử dụng: “nhận diện gian lận báo cáo tài chính”, “detecting fraudulent financial statements”. Các tài liệu được lựa chọn từ các tạp chí khoa học uy tín thuộc danh mục Scopus, ISI và Hội đồng giáo sư (đối với các bài báo trong nước), xuất bản từ 1995-2024 để đảm bảo tính cập nhật, đồng thời giữ lại các nghiên cứu nền tảng như Altman Z-score. Kết quả tìm kiếm với số lượng 145 bài báo nghiên cứu liên quan đến nhận diện GLBCTC. Những nghiên cứu không có dữ liệu thực nghiệm hoặc không rõ nguồn gốc bị loại bỏ. Đồng thời, chỉ các tài liệu có kết quả nghiên cứu đề cập đến hiệu suất dự đoán gian lận mới được tổng quan và phân tích với số lượng 38 bài (Bảng 1). Bởi vì, một trong những mục đích quan trọng của chủ đề nghiên cứu nhận diện GLBCTC là cung cấp công cụ dự đoán khả năng tồn tại GLBCTC và được đánh giá bằng hiệu suất dự đoán. Các tài liệu sau đó được phân loại theo phương pháp dự đoán gian lận: mô hình truyền thống, học máy, và tích hợp thông tin phi tài chính. Các nghiên cứu ban đầu về dự đoán GLBCTC chủ yếu xây dựng mô hình bằng các chỉ số tài chính. Kết quả nghiên cứu đã cho ra đời các mô hình kinh điển và vẫn được sử dụng cho đến ngày nay chẳng hạn như Altman Z-score và Beneish M-score.

2.1. Mô hình Altman Z-score

Theo Altman (1968) các nghiên cứu đã ngụ ý rằng các chỉ số tài chính có tiềm năng nhất định trong việc dự báo sự phá sản. Do đó, ông sử dụng năm tỷ số tài chính bao gồm: vốn lưu động/tổng tài sản; lợi nhuận giữ lại/tổng tài sản; lợi nhuận trước lãi vay và thuế/tổng tài sản; giá trị thị trường của vốn chủ sở hữu/giá trị sổ sách của tổng nợ; doanh thu/tổng tài sản để phát triển mô hình nghiên cứu. Tác giả vận dụng phương pháp phân tích phân biệt nhiều lần (Multiple discriminant analysis) để so sánh các doanh nghiệp không phá sản và đã phá sản. Kết quả cho thấy mô hình dự đoán chính xác việc phá sản tối đa hai năm trước khi thực tế

xảy ra. Mô hình Z-score ban đầu áp dụng cho các doanh nghiệp sản xuất niêm yết trên thị trường chứng khoán. Sau đó, Altman và cộng sự (1998) đã điều chỉnh loại bỏ chỉ số tài chính “doanh thu/tổng tài sản” và thay đổi các trọng số so với mô hình gốc năm 1968 nhằm áp dụng phù hợp đối với thị trường mới nổi. Mức điểm Z-score càng cao thì sức khỏe tài chính của doanh nghiệp càng tốt.

Nghiên cứu sau này (Lenard và cộng sự, 2009; Võ Văn Nghị và Hoàng Cẩm Trang, 2013) xác nhận rằng Z-score có thể là chỉ báo cho gian lận báo cáo tài chính, do mối quan hệ giữa khó khăn tài chính và thao túng lợi nhuận (DeAngelo và cộng sự, 1994; Rosner, 2003; Charitou và cộng sự, 2007; Chen và cộng sự, 2010; Li và cộng sự, 2014). Các nghiên cứu khác (Beasley và cộng sự, 1999; Kinney và McDaniel, 1989; Mishra và Drtina, 2004) cũng cho thấy các doanh nghiệp có tình trạng tài chính kém có khả năng GLBCTC cao hơn so với doanh nghiệp bình thường, củng cố vai trò của Z-score trong việc phát hiện dấu hiệu bất thường tài chính.

2.2. Mô hình Beneish M-score

Beneish (1999) phát triển M-score, một mô hình thống kê nhằm xác định doanh nghiệp thao túng thu nhập bằng cách sử dụng 8 chỉ số tài chính, bao gồm: tỷ lệ khoản phải thu/doanh thu, lợi nhuận gộp biên, chất lượng tài sản, tăng trưởng doanh thu, khấu hao tài sản, chi phí bán hàng và quản lý, biến đổi tích kế toán/tổng tài sản, và đòn bẩy tài chính. Mô hình này phân loại doanh nghiệp có hoặc không có GLBCTC dựa trên một ngưỡng điểm nhất định.

Nhiều nghiên cứu đã kiểm định độ chính xác của M-score tại các quốc gia khác nhau (Arshad và cộng sự, 2015; Herawati, 2015; Nguyen và Nguyen, 2016; Lotfi và Chadegani, 2018; Holda, 2020; Phạm Thị Mộng Tuyền, 2019; Halilbegovic và cộng sự, 2020; Hoàng Hà Anh và cộng sự, 2022). Kết quả cho thấy mô hình có độ chính xác trung bình khoảng

75%, thậm chí một số nghiên cứu chỉ đạt 40% (Bảng 1). Một số học giả (Bhavani và Amponsah, 2017; Marais và cộng sự, 2023) đã khuyến nghị thận trọng khi sử dụng mô hình này do những hạn chế về độ chính xác.

Để cải thiện hiệu suất dự báo, nhiều nghiên cứu mở rộng M-score bằng cách bổ sung biến tài chính (Kaminski và cộng sự, 2004; Kanapickienė và Grundienė, 2015; Zainudin và Hashim, 2016) hoặc tích hợp thông tin phi tài chính (Dechow và cộng sự, 2011). Dechow và cộng sự (2011) đã phát triển F-score, một biến số kết hợp dữ liệu tài chính và thông tin từ thị trường chứng khoán. Mô hình này đạt độ chính xác khoảng 65,9% khi F-score > 1, cho thấy xác suất cao doanh nghiệp gian lận. Mô hình Dechow F-score đã được kiểm định và mở rộng trong các nghiên cứu về nhận diện GLBCTC tại Việt Nam (Nguyễn Tiến Hùng và Võ Hồng Đức, 2017; Bùi Phương Chi và cộng sự, 2021; Trần Thị Giang Tân và cộng sự, 2015; Dang và cộng sự, 2017; Nguyễn Tiến

Hùng và cộng sự, 2018; Phạm Thị Mộng Tuyền, 2019).

Bên cạnh phương pháp hồi quy logistic truyền thống, nhiều nghiên cứu đã áp dụng học máy để nâng cao độ chính xác dự báo (Green và Choi, 1997; Feroz và cộng sự, 2000; Gaganis, 2009; Lin và cộng sự, 2015; Hajek và Henriques, 2017; Jan, 2018; Hajek, 2019; Đặng Ngọc Hùng và cộng sự, 2022; Nguyen và cộng sự, 2022). Một số nghiên cứu còn kết hợp thông tin phi tài chính trong mô hình học máy để tối ưu hóa khả năng phát hiện gian lận (Ata và Seyrek, 2009; Omar và cộng sự, 2017; Lin và cộng sự, 2003; Kotsiantis và cộng sự, 2006; Kirkos và cộng sự, 2007; Liou, 2008; Ravisankar và cộng sự, 2011; Lokanan và cộng sự, 2019). Tổng hợp từ các nghiên cứu trước đây cho thấy phương pháp học máy có hiệu suất dự báo cao hơn so với phương pháp thống kê truyền thống, đặc biệt khi tích hợp cả chỉ số tài chính và thông tin phi tài chính (Bảng 1).

Bảng 1. Thống kê kết quả nghiên cứu nhận diện gian lận trong báo cáo tài chính

MÔ HÌNH SỬ DỤNG CHỈ SỐ TÀI CHÍNH					
Phương pháp truyền thống					
STT	Tác giả	Quốc gia	Thuật toán/Phương pháp	Tỷ lệ mẫu (gian lận: không gian lận)	Kết quả dự đoán gian lận tổng thể
1	Persons (1995)	Mỹ	Hồi quy Logistic (Logistic regression)	280:203 (Nhiều: 1)	64%
2	Kaminski và cộng sự (2004)	Mỹ	Phân tích phân biệt (Discriminant analysis)	79:79 (1:1)	42%
3	Kanapickienė và Grundienė (2015)	Lithuania	Hồi quy Logistic (Logistic regression)	40:125 (1:nhiều)	84,80%
4	Arshad và cộng sự (2015)	Malaysia	Hồi quy Logistic (Logistic regression)	24:24 (1:1)	83,30%
5	Tarjo và Herawati (2015)	Indonesia	Hồi quy Logistic (Logistic regression)	35:35 (1:1)	77,10%
6	Nguyen và Nguyen (2016)	Việt Nam	Beneish M-Score	223:0	48,40%
7	Lotfi và Chadegani (2018)	Iran	Hồi quy Logistic (Logistic regression)	137:0	66,03%
8	Hołda (2020)	Ba Lan	Beneish M - score	30:30 (1:1)	100%

9	Halilbegovic và cộng sự (2020)	Liên bang Bosnia và Herzegovina	Beneish M-Score	68	79,41%
10	Hoàng Hà Anh và cộng sự (2022)	Việt Nam	Beneish M-Score Hồi quy Logistic (Logistic regression)	38:65 (1: nhiều)	44,70%
Phương pháp học máy					
11	Green và Choi (1997)	Mỹ	Mạng nơ ron (Neural networks)	86:95 (1: nhiều)	68,64%
12	Lin và cộng sự (2003)	Mỹ	Mạng nơ ron mờ (Neuron - Fuzzy)	20:80 (1: nhiều)	35%
			Hồi quy Logistic (Logistic regression)		97%
13	Kotsiantis và cộng sự (2006)	Hy Lạp	Cây quyết định (Decision tree)	41:123 (1: nhiều)	85,20%
			Mạng nơ ron nhân tạo (Artificial neural networks)		36,60%
			Mạng niềm tin Bayesian (Bayesian Belief network)		51,20%
			Hồi quy Logistic (Logistic regression)		36,60%
			k - láng giềng gần nhất (k-Nearest Neighbour)		56,10%
			Máy vector hỗ trợ (Support vector machine)		48,80%
14	Kirkos và cộng sự (2007)	Hy Lạp	Cây quyết định (Decision tree)	38:38 (1: 1)	75%
			Mạng niềm tin Bayesian (Bayesian Belief network)		91,70%
			Mạng nơ ron nhân tạo (Neural networks)		82,50%
15	Liou (2008)	Đài Loan	Hồi quy Logistic (Logistic regression)	3018:3019 (1: 1)	99,05%
			Mạng nơ ron (Neural networks)		95,82%
			Cây quyết định (Decision tree)		95,59%
16	Ata và Seyrek (2009)	Thổ Nhĩ Kỳ	Cây quyết định (Decision tree)	50:50 (1: 1)	67,92%
17	Ravisankar và cộng sự (2011)	Trung Quốc	Mạng thần kinh truyền thẳng đa lớp (Multiplayer Feed Forward Neural Network)	101:101 (1: 1)	78,36%
			Máy vector hỗ trợ (Support vector machine)		70,41%
			Lập trình di truyền (Genetic programming)		94,14%
			Phương pháp xử lý dữ liệu nhóm (Group Method of Data Handling)		93%
			Mạng nơ ron xác suất (Probabilistic Neural Network)		98,09%

TRƯỜNG ĐẠI HỌC KINH TẾ - ĐẠI HỌC ĐÀ NẴNG

			Hồi quy Logistic (Logistic regression)		66,86%
			Mạng nơ ron (Artificial neural networks)		77,36%
18	Omar và cộng sự (2017)	Malaysia	Mạng nơ ron nhân tạo (Artificial neural networks)	75:475 (1: nhiều)	94,87%
MÔ HÌNH SỬ DỤNG CHỈ SỐ TÀI CHÍNH KẾT HỢP THÔNG TIN PHI TÀI CHÍNH					
Phương pháp thống kê truyền thống					
19	Trần Thị Giang Tân và cộng sự (2015)	Việt Nam	Hồi quy Logistic (Logistic regression)	39:39 (1: 1)	80%
20	Nguyễn Tiến Hùng và Võ Hồng Đức (2017)	Việt Nam	Dechow F - score Hồi quy Logistic (Logistic regression)	44:44 (1: 1)	75%
21	Dang và cộng sự (2017)	Việt Nam	Dechow F- score Hồi quy Logistic (Logistic regression)	151:491 (1: nhiều)	78,21%
22	Nguyễn Tiến Hùng và cộng sự (2018)	Việt Nam	Hồi quy Logistic (Logistic regression)	51:51 (1: 1)	72%
23	Phạm Thị Mộng Tuyền (2019)	Việt Nam	Hồi quy Logistic (Logistic regression)	146:304 (1: nhiều)	40,40%
Phương pháp học máy					
24	Feroz và cộng sự (2000)	Mỹ	Mạng nơ ron nhân tạo (Artificial neural networks)	42:90 (1: nhiều)	81%
25	Gaganis (2009)	Hy Lạp	Phân tích cấu trúc đa nhóm (Multi-group Hierarchical Discrimination)	149:149 (1: 1)	86%
			Hàm tiện ích cộng chấm điểm và phân loại (Utilités Additives Discriminantes)		80,49%
			Máy vector hỗ trợ (Support vector machine)		84,15%
			Phân tích phân biệt (Discriminant analysis)		84,10%
			Mạng nơ ron nhân tạo (Artificial neural networks)		78,05%
			Hồi quy Logistic (Logistic regression)		78%
			k-láng giềng gần nhất (k-Nearest neighbor)		78,05%
			Mạng nơ ron xác suất (Probabilistic neural networks)		78,05%
26	Abbasi và cộng sự (2012)	Mỹ	Máy vector hỗ trợ (Support vector machine)	815:8191 (1: nhiều)	90,4%
			Bộ phân loại Naïve Bayes (Naïve Bayes)		85,1%

			Cây quyết định xen kẽ (Alternating decision tree)		89,6%
			Rừng ngẫu nhiên (Random forest)		85,7%
			Giảm lỗi bằng cách cắt tia cây (Reduced error pruning tree)		88,8%
			k-láng giềng gần nhất (k-Nearest neighbor)		86,5%
			Cắt tia gia tăng lặp lại để giảm lỗi (Repeated incremental pruning to produce error reduction)		87%
			Hồi quy Logistic (Logistic regression)		87,5%
			Mạng nơ ron nhân tạo (Artificial neural networks)		86,5%
27	Lin và cộng sự (2015)	Đài Loan	Mạng nơ ron nhân tạo (Artificial neural networks)	129:447 (1: nhiều)	92,8%
			Hồi quy Logistic (Logistic regression)		90,3%
			Cây quyết định và phân loại (Classification and regression trees)		88,5%
28	Liu và cộng sự (2015)	Trung Quốc	Rừng ngẫu nhiên (Random forest)	138:160 (1: nhiều)	88%
			Máy vector hỗ trợ (Support vector machine)		80,18%
			Cây hồi quy và phân loại (Classification and regression trees)		66,43%
			k-láng giềng gần nhất (k-Nearest neighbor)		60,11%
			Hồi quy Logistic (Logistic regression)		42,91%
29	Hajek và Henriques (2017)	Mỹ	Hồi quy Logistic (Logistic regression)	311:311 (1: 1)	74,53%
			Bộ phân loại Naïve Bayes (Naïve Bayes)		57,83%
			Mạng niềm tin Bayesian (Bayesian Belief Networks)		90,32%
			Bảng quyết định (Decision Table)		89,50%
			Bộ phân loại Naïve Bayes (Naïve Bayes)		
			Máy vector hỗ trợ (Support vector machine)		77,95%
			Cắt tia gia tăng lặp lại để giảm lỗi (Repeated Incremental Pruning to Produce Error Reduction)		87,01%
			Phân mở rộng của cây phân loại vào hồi quy (C4.5)		86,10%

TRƯỜNG ĐẠI HỌC KINH TẾ - ĐẠI HỌC ĐÀ NẴNG

			Cây hồi quy và phân loại (Classification and regression trees)		86,10%
			Cây mô hình logistic (Logistic model trees)		85,44%
			Mạng thần kinh truyền thẳng đa lớp (Multiplayer Feed Forward Neural Network)		77,93%
			Mạng nơ ron phân loại perceptron (Voted perceptron)		51,16%
			Học tập tổng hợp (đào tạo mô hình học tập độc lập) (Bagging)		87,09%
			Rừng ngẫu nhiên (Random forest)		87,50%
			Tăng cường thích ứng (Adaboost)		77,29%
30	Jan (2018)	Đài Loan	Mạng nơ ron nhân tạo (Artificial neural networks) + Cây hồi quy và phân loại (Classification and regression trees)	40:120 (1: nhiều)	90,21%
			Mạng nơ ron nhân tạo (Artificial neural networks) + Phát hiện tương tác tự động Chi-square (Chi-square automatic interaction detector)		90,06%
			Mạng nơ ron nhân tạo (Artificial neural networks) + Phân mở rộng của cây phân loại vào hồi quy (C5.0)		87,39%
			Mạng nơ ron nhân tạo (Artificial neural networks) + Cây thống kê nhanh, không thiên vị, hiệu quả (Quick unbiased efficient statistical tree)		83,50%
			Máy vector hỗ trợ (Support vector machine) + Cây hồi quy và phân loại (Classification and regression trees)		86,18%
			Máy vector hỗ trợ (Support vector machine) + Phát hiện tương tác tự động Chi-square (Chi-square automatic interaction detector)		79,75%
			Máy vector hỗ trợ (Support vector machine) + Phân mở rộng của cây phân loại vào hồi quy (C5.0)		74,30%
			Máy vector hỗ trợ (Support vector machine) +		77,17%

			Cây thống kê nhanh, không thiên vị, hiệu quả (Quick unbiased efficient statistical tree)		
31	Hajek (2019)	Mỹ	Hệ thống dựa trên quy tắc mờ (Fuzzy rule-based system)	311:311 (1: 1)	86,80%
			Mạng niềm tin Bayesian (Bayesian Belief Networks)		89,80%
			Bảng quyết định (Decision Table)/ Bộ phân loại Naïve Bayes (Naïve Bayes)		87,80%
			Rừng ngẫu nhiên (Random forest)		90,40%
32	Omidi và cộng sự (2019)	Trung quốc	Mạng thần kinh truyền thẳng đa lớp (Multi-layer feed-forward neural network)	910:1749 (1: nhiều)	97,30%
			Mạng nơ ron xác suất (Probabilistic Neural network)		97,76%
			Máy vector hỗ trợ (Support vector machine)		90,70%
			Mô hình logarit tuyến tính đa thức (Multinomial log-linear Model)		89,40%
			Phân tích phân biệt (Discriminant analysis)		90,50%
33	Craja và cộng sự (2020)	Mỹ	Hồi quy Logistic (Logistic regression)	208:7341 (1: nhiều)	84,24%
			Rừng ngẫu nhiên (Random forest)		87,39%
			Máy vector hỗ trợ (Support vector machine)		82,80%
			Tăng cường độ dốc cực đại (Extreme gradient Boosting)		90,83%
			Mạng nơ ron (Artificial neural networks)		90,54%
			Phân tích cấu trúc đa nhóm (Multi-group Hierarchical Discrimination)		65,06%
			Mô hình ngôn ngữ người học đa nhiệm không có giám sát (Language models are unsupervised multitask learners)		69,34%
34		Đài Loan	Cây quyết định (Decision tree)	181:256	93,43%

	Cheng và cộng sự (2021)			(1: nhiều)	
			Rừng ngẫu nhiên (Random forest)		98,92%
			Bộ phân loại cây bổ sung (ExtraTree)		97,56%
			Mạng nơ ron nhân tạo (Artificial neural networks)		85,47%
35	Đặng Ngọc Hùng và cộng sự (2022)	Việt Nam	Rừng ngẫu nhiên (Random forest)	467:1768 (1: nhiều)	91%
36	Nguyễn Anh Phong và cộng sự (2022)	Việt Nam	Mạng nơ ron nhân tạo (Artificial neural networks)	227:637 (1: nhiều)	97%
			Beneish M- Score		
			Mạng nơ ron nhân tạo (Artificial neural networks)		89%
			Atman Z- Score		
37	Cao và cộng sự (2024)	Việt Nam	Máy vector hỗ trợ (Support vector machine)	2235	97%
			Beneish M- Score		
			Máy vector hỗ trợ (Support vector machine)		95%
			Atman Z- Score		
38	Ali và cộng sự (2023)		Rừng ngẫu nhiên (Random forest)	102: 1798 (1: nhiều)	91%
			Mạng nơ ron nhân tạo (Artificial neural networks)		99%
			Hồi quy Logistic (Logistic regression)		73,8%
			Cây quyết định (Decision tree)		82,2%
			Máy vector hỗ trợ (Support vector machine)		88,8%
			Rừng ngẫu nhiên (Random forest)		80,5%
			Tăng cường thích ứng (AdaBoost)		83,3%
			Tăng cường độ dốc cực đại (Extreme gradient Boosting)		93,6%

Nguồn: Tác giả tổng hợp

3. Đánh giá nghiên cứu đã thực hiện và đề xuất hướng nghiên cứu

3.1. Tổng quan kết quả nghiên cứu trước

Các mô hình nhận diện GLBCTC chủ yếu dựa trên chỉ số tài chính hoặc kết hợp với thông tin phi tài chính. Việc sử dụng các chỉ số tài chính xuất phát từ quy định về trách nhiệm của kiểm toán viên trong phát hiện gian lận, thông qua thủ tục phân tích để đánh giá rủi ro sai sót trọng yếu (Thornhill, 1995; Kaminski và cộng sự, 2004; Albrecht và cộng sự, 2008). Thủ tục

này bao gồm phân tích xu hướng, tỷ lệ tài chính, và kiểm tra tính hợp lý của dữ liệu doanh nghiệp (Kaminski và cộng sự, 2004).

Mặc dù phương pháp này hữu ích trong việc phát hiện bất thường, nhiều nghiên cứu thực nghiệm chỉ ra hạn chế của chỉ số tài chính trong dự báo gian lận (Kaminski và cộng sự, 2004; Nia, 2015). Điều này phản ánh sự thay đổi trong cách thức thao túng báo cáo tài chính, làm giảm độ chính xác của các mô hình truyền thống. Bhavani và Amponsah (2017) cũng nhấn mạnh rằng không có một công cụ duy

nhất nào có thể phát hiện gian lận hiệu quả. Do đó, nhiều nghiên cứu đã bổ sung thông tin phi tài chính để cải thiện độ chính xác của mô hình nhận diện gian lận. Thông tin phi tài chính được các học giả định nghĩa là những thông tin bổ sung cho thông tin tài chính và nó có thể bao hàm đến các chủ đề rộng hơn cả trách nhiệm xã hội (CSR), môi trường, xã hội, quản trị doanh nghiệp (ESG) và tính bền vững (Tarquinio, 2020). Thông tin phi tài chính giúp gia tăng sự hiểu biết sâu sắc và toàn diện hơn về đơn vị. Do đó, các nghiên cứu đã khai thác thông tin phi tài chính nhằm tìm kiếm các dấu hiệu giúp nhận diện GLBCTC. Các thông tin phi tài chính đưa vào mô hình dựa trên nền tảng Lý thuyết tam giác gian lận được đề xuất bởi Cressey (1953). Lý thuyết này giải thích nguyên nhân gian lận xuất phát từ 3 yếu tố áp lực, cơ hội và hợp lý hóa. Nhiều học giả đã sử dụng các thông tin phi tài chính đại diện cho các yếu tố của tam giác gian lận để nghiên cứu dự đoán GLBCTC. Ngoài ra, Lý thuyết đại diện của Jensen và Meckling (1976) và Lý thuyết bất cân xứng thông tin của Akerlof (1970) đề cập việc tồn tại sự khác biệt về lợi ích và sự bất cân xứng thông tin giữa người ủy quyền (chủ sở hữu) và người đại diện (nhà quản lý) có thể tạo ra một môi trường trong đó gian lận có nhiều khả năng phát sinh hơn (Ali và cộng sự, 2020). Các lý thuyết này góp phần làm rõ bản chất của GLBCTC và được các nghiên cứu sử dụng làm cơ sở cho việc bổ sung và lựa chọn các biến đưa vào mô hình dự đoán đặc biệt là thông tin phi tài chính liên quan đến yếu tố về quản trị doanh nghiệp.

Bên cạnh đó, phương pháp học máy ngày càng được ứng dụng rộng rãi trong phát hiện gian lận. Ban đầu, các nghiên cứu sử dụng hồi quy logistic, nhưng với sự phát triển của công nghệ, các kỹ thuật học máy như Mạng nơ-ron nhân tạo (Artificial neural networks), Máy vector hỗ trợ (Support vector machine), Cây quyết định (Decision tree), Rừng ngẫu nhiên (Random forest), Mạng niềm tin Bayesian

(Bayesian Belief Networks), và Lập trình di truyền (Genetic programming) đã chứng minh độ chính xác cao hơn trong phân loại gian lận. Sự ưu việt của học máy trong dự đoán gian lận được giải thích qua hai lý do chính. Thứ nhất, các mô hình thống kê truyền thống dựa trên các giả định về tính tuyến tính, phân phối chuẩn và độc lập giữa các biến, nhưng với số lượng biến ngày càng gia tăng, điều này dễ dẫn đến vi phạm các giả định, ảnh hưởng đến hiệu suất dự báo (Breiman, 2001). Trong khi đó, học máy không yêu cầu cấu trúc mô hình cố định, giúp xử lý dữ liệu có tương quan cao và không cần giả định về phân phối (Carmichael và Marron, 2018). Thứ hai, học máy dựa trên Lý thuyết học thống kê, cho phép xây dựng các hàm dự đoán tối ưu từ dữ liệu lớn và phức tạp (Bennett và cộng sự, 2022). Các thuật toán học máy có thể phát hiện các khuôn mẫu ẩn trong dữ liệu mà phương pháp truyền thống không nhận diện được (Carmichael và Marron, 2018), từ đó nâng cao độ chính xác của mô hình dự báo.

Dựa trên những phát hiện trên, các hướng nghiên cứu tiềm năng bao gồm việc tích hợp dữ liệu phi tài chính vào mô hình học máy nhằm cải thiện hiệu suất dự báo, phát triển mô hình kết hợp giữa các kỹ thuật học máy tiên tiến như Deep Learning và Ensemble Learning để tối ưu hóa khả năng nhận diện gian lận, cũng như mở rộng ứng dụng học máy trong bối cảnh thị trường mới nổi, đặc biệt là Việt Nam, nơi gian lận tài chính vẫn là một thách thức lớn. Nhìn chung, học máy đang trở thành xu hướng tất yếu trong nghiên cứu nhận diện gian lận báo cáo tài chính, giúp nâng cao hiệu suất dự đoán và hỗ trợ các bên liên quan trong việc phát hiện gian lận sớm hơn.

3.2. Đề xuất hướng nghiên cứu

3.2.1. Tính cấp thiết và những thách thức trong nghiên cứu nhận diện GLBCTC tại Việt Nam

Học máy là một trong số các thành tựu của sự phát triển khoa học công nghệ trên thế giới. Dù chỉ là một nhánh của trí tuệ nhân tạo nhưng

đã cung cấp một công cụ tính toán phục vụ khai thác và phân tích dữ liệu, hữu ích trong việc trích xuất thông tin, nhận dạng khuôn mẫu dữ liệu và dự đoán (Ge và cộng sự, 2017). Nó cho thấy tiềm năng mang lại lợi ích to lớn của việc phát triển khoa học công nghệ. Nhận thức được điều này, Đảng và nhà nước ta luôn đề cao tầm quan trọng của việc phát triển khoa học công nghệ và xem đây là quốc sách hàng đầu. Chủ trương này được thể hiện trong các Nghị quyết tại các kỳ đại hội Đảng. Chẳng hạn như Chiến lược phát triển kinh tế - xã hội 10 năm 2021 - 2030 được thông qua tại Đại hội XIII nhấn mạnh *“Phát triển mạnh mẽ khoa học, công nghệ, đổi mới sáng tạo và chuyển đổi số là động lực chính để tăng trưởng kinh tế”*. Nhờ có sự định hướng của nhà nước cùng với sự năng động của các tổ chức, doanh nghiệp, bước đầu Việt Nam đã đạt được những thành tựu đáng ghi nhận về phát triển khoa học công nghệ trong nhiều lĩnh vực bao gồm cả kinh tế, tài chính. Đặc biệt là việc ứng dụng trí tuệ nhân tạo (học máy) đã được chú trọng trong hoạt động của tổ chức, doanh nghiệp cũng như hoạt động nghiên cứu dự báo những vấn đề liên quan đến kinh tế vĩ mô, tài chính. Tuy nhiên, đối với việc dự đoán GLBCTC tại Việt Nam, số lượng nghiên cứu nói chung và số lượng sử dụng phương pháp học máy nói riêng còn hạn chế (Bảng 1). Trong khi đó, hành vi GLBCTC đặc biệt là ở các doanh nghiệp niêm yết Việt Nam vẫn còn tồn tại với các thủ thuật thực hiện và che giấu ngày càng tinh vi. Các vụ việc GLBCTC này thường liên quan đến các sai phạm trong hoạt động quản trị doanh nghiệp và kiểm toán, dẫn đến những thiệt hại đặc biệt nghiêm trọng. Hậu quả của GLBCTC không chỉ ảnh hưởng đến thị trường tài chính mà còn tác động tiêu cực đến tình hình kinh tế - xã hội điển hình như vụ bê bối tài chính của tập đoàn Vạn Thịnh Phát. Do đó, việc nghiên cứu nhận diện GLBCTC, đặc biệt là vận dụng các phương pháp học máy hiện đại theo xu hướng của thế giới là điều cần thiết nhằm cung cấp

công cụ hữu ích giúp cảnh báo sớm cho các bên liên quan.

Việc nghiên cứu chủ đề này trong bối cảnh Việt Nam có thể gặp phải những thách thức liên quan đến dữ liệu, nguồn nhân lực và cơ sở hạ tầng công nghệ. Hiện nay, Việt Nam chưa có cơ sở dữ liệu công khai về GLBCTC. Bên cạnh đó, dữ liệu về thông tin tài chính và phi tài chính phục vụ cho việc lựa chọn các biến/thuộc tính đưa vào mô hình dự đoán còn rời rạc, chưa đầy đủ. Trong khi đó, đối với phương pháp học máy, chất lượng dữ liệu đầu vào là yếu tố quan trọng nhất quyết định kết quả đầu ra do dựa trên nguyên lý học tập từ dữ liệu (Mahesh, 2020). Ngoài ra, nước ta còn thiếu hụt nguồn nhân lực liên ngành vừa am hiểu chuyên môn về kinh tế, tài chính, kế toán và trí tuệ nhân tạo (học máy). Đồng thời, cơ sở hạ tầng phục vụ khai thác dữ liệu chưa đủ mạnh để xử lý dữ liệu lớn. Đây là những trở ngại trong việc ứng dụng trí tuệ nhân tạo (học máy) trong hoạt động nghiên cứu kinh tế nói chung và phát hiện GLBCTC nói riêng.

3.2.2. Đề xuất hướng nghiên cứu

a) Mô hình nghiên cứu

Mặc dù nhiều nghiên cứu về nhận diện GLBCTC đã sử dụng chỉ số tài chính và thông tin phi tài chính, nhưng phần lớn đều kế thừa từ các mô hình truyền thống. Tại Việt Nam, các nghiên cứu chủ yếu áp dụng Altman Z-score, Beneish M-score, Dechow F-score, trong khi số lượng nghiên cứu kết hợp mô hình hoặc bổ sung biến số mới còn hạn chế. Do đó, việc mở rộng và tích hợp thêm chỉ số tài chính chưa được khai thác cùng với các yếu tố phi tài chính có thể cải thiện độ chính xác và độ tin cậy của mô hình dự đoán gian lận.

b) Phương pháp nghiên cứu

Biến phụ thuộc trong hầu hết các mô hình nhận diện gian lận là biến nhị phân “Gian lận báo cáo tài chính”, với hai giá trị: có gian lận và không gian lận. Các nghiên cứu quốc tế

thường dựa vào dữ liệu công khai từ cơ quan chức năng, trong khi tại Việt Nam, gian lận chủ yếu được xác định qua chênh lệch lợi nhuận trước và sau kiểm toán do thiếu nguồn dữ liệu chính thức. Tuy nhiên, gian lận tài chính mang bản chất cố ý, tạo ra thách thức trong việc đo lường và xác định thang đo phù hợp. Do đó, cần có nghiên cứu bổ sung để đảm bảo tính chặt chẽ và chính xác trong việc phân biệt sai sót kế toán và gian lận có chủ đích.

Hầu hết các mô hình định lượng tại Việt Nam vẫn dựa trên hồi quy logistic, trong khi các nghiên cứu áp dụng học máy còn hạn chế. Do đó, việc ứng dụng phương pháp học máy trong mô hình nhận diện gian lận tài chính là cần thiết nhằm nâng cao hiệu suất dự đoán và tối ưu hóa khả năng phát hiện gian lận trong bối cảnh doanh nghiệp Việt Nam.

c) Phạm vi nghiên cứu

Các nghiên cứu về GLBCTC tại Việt Nam chủ yếu tập trung vào doanh nghiệp niêm yết trên một sàn giao dịch chứng khoán (HOSE hoặc HNX), trong khi rất ít nghiên cứu mở rộng trên cả ba sàn giao dịch (HOSE, HNX, UpCom). Do đó, việc mở rộng phạm vi nghiên cứu để bao quát toàn bộ doanh nghiệp niêm yết có thể mang lại đóng góp quan trọng cả về mặt

lý thuyết lẫn thực tiễn, giúp xây dựng một mô hình nhận diện gian lận phù hợp hơn với thị trường tài chính Việt Nam.

4. Kết luận

GLBCTC là một chủ đề thu hút sự quan tâm rộng rãi trong giới học thuật, dẫn đến sự ra đời của nhiều mô hình định lượng nhằm dự báo và phát hiện gian lận. Các mô hình truyền thống như Altman Z-score, Beneish M-score, Dechow F-score vẫn được sử dụng và mở rộng, trong khi các phương pháp hiện đại, đặc biệt là học máy, ngày càng được ứng dụng để nâng cao độ chính xác và hiệu suất dự đoán.

Tại Việt Nam, dù đã có một số nghiên cứu về chủ đề này, nhưng số lượng còn hạn chế và chưa khai thác đầy đủ tiềm năng của dữ liệu phi tài chính và công nghệ học máy. Tổng quan nghiên cứu cho thấy đa số các mô hình hiện tại vẫn tập trung vào chỉ số tài chính, trong khi việc tích hợp thêm các yếu tố khác có thể giúp tăng cường khả năng phát hiện gian lận. Do đó, việc mở rộng nghiên cứu theo hướng ứng dụng học máy và kết hợp dữ liệu phi tài chính là cần thiết để hoàn thiện nền tảng lý thuyết và cung cấp giải pháp thực tiễn cho thị trường tài chính Việt Nam.

TÀI LIỆU THAM KHẢO

- Abbasi, A., Albrecht, C., Vance, A., & Hansen, J. (2012). Metafraud: a meta-learning framework for detecting financial fraud. *Mis Quarterly*, 36(4), 1293-1327. <https://doi.org/10.2307/41703508>
- ACFE. (2016). *Report to the Nations on Occupational Fraud and Abuse: 2016 Global Fraud Study*. <https://www.acfe.com/-/media/files/acfe/pdfs/rtnn/2016/2016-report-to-the-nations.pdf>
- ACFE. (2018). *Report to the Nations: 2018 Global Study on Occupational Fraud and Abuse*. <https://s3-us-west-2.amazonaws.com/acfe-public/2018-report-to-the-nations.pdf>
- ACFE. (2020). *Report to the Nations: 2020 Global Study on Occupational Fraud and Abuse*. <https://acfe-public.s3-us-west-2.amazonaws.com/2020-Report-to-the-Nations.pdf>
- ACFE. (2022). *Occupational Fraud 2022: A Report to the Nations*. <https://acfe-public.s3.us-west-2.amazonaws.com/2022+Report+to+the+Nations.pdf>

- ACFE. (2024). *Occupational Fraud 2024: A Report to the Nations*. <https://www.acfe.com/-/media/files/acfe/pdfs/rtnn/2024/2024-report-to-the-nations.pdf>
- Akerlof, G. A. (1970). The Market for "Lemons": Quality Uncertainty and the Market Mechanism. *The Quarterly Journal of Economics*, 84(3), 488-500.
- Albrecht, C., Kranacher, M. J., & Albrecht, S. (2008). Asset misappropriation research white paper for the Institute for Fraud Prevention. *Institute for Fraud Prevention, Research studies*.
- Ali, A. A., Khedr, A. M., El-Bannany, M., & Kanakkayil, S. (2023). A powerful predicting model for financial statement fraud based on optimized XGBoost ensemble learning technique. *Applied Sciences*, 13(4), 2272. <https://doi.org/10.3390/app13042272>
- Ali, C. B., Boubaker, S., & Magnan, M. (2020). Auditors and the principal-principal agency conflict in family controlled firms. *Auditing: A Journal of Practice & Theory*, 39(4), 31-55. <https://doi.org/10.2308/AJPT-17-147>
- Altman, E. I. (1968). Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. *The Journal of Finance*, 23(4), 589-609. <https://doi.org/10.2307/2978933>
- Altman, E. I., Hartzell, J., & Peck, M. (1998). Emerging market corporate bonds—a scoring system. In *Emerging Market Capital Flows: Proceedings of a Conference held at the Stern School of Business, New York University on May 23-24, 1996* (pp. 391-400). Boston, MA: Springer US. https://doi.org/10.1007/978-1-4615-6197-2_25
- Arshad, R., Iqbal, S. M., & Omar, N. (2015). Prediction of business failure and fraudulent financial reporting: Evidence from Malaysia. *Indian Journal of Corporate Governance*, 8(1), 34-53. <https://doi.org/10.1177/0974686215574424>
- Ata, H. A., & Seyrek, I. H. (2009). The use of data mining techniques in detecting fraudulent financial statements: an application on manufacturing firms. *Suleyman Demirel University Journal of Faculty of Economics & Administrative Sciences*, 14(2).
- Beasley, M. S., Carcello, J. V., Hermanson, D. R., & Committee of Sponsoring Organizations of the Treadway Commission. (1999). *Fraudulent financial reporting: 1987-1997: an analysis of US public companies*.
- Beneish, M. D. (1999). The detection of earnings manipulation. *Financial Analysts Journal*, 55(5), 24-36. <https://doi.org/10.2469/faj.v55.n5.2296>
- Bennett, M., Hayes, K., Kleczyk, E. J., & Mehta, R. (2022). Similarities and differences between machine learning and traditional advanced statistical modeling in healthcare analytics. *arXiv preprint arXiv:2201.02469*. <https://doi.org/10.48550/arXiv.2201.02469>
- Bhavani, G., & Amponsah, C. T. (2017). M-Score and Z-Score for detection of accounting fraud. *Accountancy Business and the Public Interest*, 1(1), 68-86.
- Breiman, L. (2001). Statistical modeling: The two cultures (with comments and a rejoinder by the author). *Statistical Science*, 16(3), 199-231. <https://doi.org/10.1214/ss/1009213726>
- Bùi Phương Chi, Nguyễn Thị Hồng Thuý, & Lăng Trịnh Mai Hương. (2021). Dự báo rủi ro gian lận báo cáo tài chính của các công ty niêm yết trên thị trường chứng khoán Việt Nam. *VNU Journal of Economics and Business*, 37(1), 40-49. <https://doi.org/10.25073/2588-1108/vnueab.4494>

- Carmichael, I., & Marron, J. S. (2018). Data science vs. statistics: two cultures?. *Japanese Journal of Statistics and Data Science*, 1(1), 117-138. <https://doi.org/10.1007/s42081-018-0009-3>
- Charitou, A., Lambertides, N., & Trigeorgis, L. (2007). Earnings behaviour of financially distressed firms: The role of institutional ownership. *Abacus*, 43(3), 271-296. <https://doi.org/10.1111/j.1467-6281.2007.00230.x>
- Chen, Y., Chen, C. H., & Huang, S. L. (2010). An appraisal of financially distressed companies' earnings management: Evidence from listed companies in China. *Pacific Accounting Review*, 22(1), 22-41. <https://doi.org/10.1108/01140581011034209>
- Cheng, C. H., Kao, Y. F., & Lin, H. P. (2021). A financial statement fraud model based on synthesized attribute selection and a dataset with missing values and imbalanced classes. *Applied Soft Computing*, 108, 107487. <https://doi.org/10.1016/j.asoc.2021.107487>
- Cao, T. N., Dang, N. H., & Vu, T. T. B. (2024). Using random forest and artificial neural network to detect fraudulent financial reporting: Data from listed companies in Vietnam. *Calitatea*, 25(202), 160-173. <https://doi.org/10.47750/QAS/25.202.17>
- Craja, P., Kim, A., & Lessmann, S. (2020). Deep learning for detecting financial statement fraud. *Decision Support Systems*, 139, 113421. <https://doi.org/10.1016/j.dss.2020.113421>
- Cressey, D. R. (1953). *Other people's money; a study of the social psychology of embezzlement*. New York, US: Free Press.
- Đặng Ngọc Hùng, Phạm Thị Hồng Diệp, & Cao Thị Nhiên. (2022). Nghiên cứu gian lận báo cáo tài chính tiếp cận theo thuật toán rừng ngẫu nhiên. *Tạp chí Khoa học và Công nghệ*, 58(2), 155-161.
- Dang, H. N., Hoang, H. T. V., & Dang, B. T. (2017). Application of F-score in predicting fraud, errors: Experimental research in Vietnam. *International Journal of Accounting and Financial Reporting*, 7(2), 303-322. <https://doi.org/10.5296/ijaf.v7i2.12174>
- DeAngelo, H., DeAngelo, L., & Skinner, D. J. (1994). Accounting choice in troubled companies. *Journal of Accounting and Economics*, 17(1-2), 113-143. [https://doi.org/10.1016/0165-4101\(94\)90007-8](https://doi.org/10.1016/0165-4101(94)90007-8)
- Dechow, P. M., Ge, W., Larson, C. R., & Sloan, R. G. (2011). Predicting material accounting misstatements. *Contemporary Accounting Research*, 28(1), 17-82. <https://doi.org/10.1111/j.1911-3846.2010.01041.x>
- Feroz, E. H., Kwon, T. M., Pastena, V. S., & Park, K. (2000). The efficacy of red flags in predicting the SEC's targets: an artificial neural networks approach. *International Journal of Intelligent Systems in Accounting, Finance and Management*, 9(3), 145-157. [https://doi.org/10.1002/1099-1174\(200009\)9:3<145::AID-ISAF185>3.0.CO;2-G](https://doi.org/10.1002/1099-1174(200009)9:3<145::AID-ISAF185>3.0.CO;2-G)
- Gaganis, C. (2009). Classification techniques for the identification of falsified financial statements: a comparative analysis. *Intelligent Systems in Accounting, Finance and Management: International Journal*, 16(3), 207-229. <https://doi.org/10.1002/isaf.303>
- Ge, Z., Song, Z., Ding, S. X., & Huang, B. (2017). Data mining and analytics in the process industry: The role of machine learning. *Ieee Access*, 5, 20590-20616. <https://doi.org/10.1109/ACCESS.2017.2756872>
- Green, B. P., & Choi, J. H. (1997). Assessing the risk of management fraud through neural network technology. *Auditing*, 16, 14-28.

- Gupta, S., & Mehta, S. K. (2024). Data mining-based financial statement fraud detection: Systematic literature review and meta-analysis to estimate data sample mapping of fraudulent companies against non-fraudulent companies. *Global Business Review*, 25(5), 1290-1313. <https://doi.org/10.1177/0972150920984857>
- Hajek, P. (2019). Interpretable fuzzy rule-based systems for detecting financial statement fraud. In *IFIP international conference on artificial intelligence applications and innovations* (pp. 425-436). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-030-19823-7_36
- Hajek, P., & Henriques, R. (2017). Mining corporate annual reports for intelligent detection of financial statement fraud-A comparative study of machine learning methods. *Knowledge-Based Systems*, 128, 139-152. <https://doi.org/10.1016/j.knosys.2017.05.001>
- Halilbegovic, S., Celebic, N., Cero, E., Buljubasic, E., & Mekic, A. (2020). Application of Beneish M-score model on small and medium enterprises in Federation of Bosnia and Herzegovina. *Eastern Journal of European Studies*, 11(1), 146-163.
- Herawati, N. (2015). Application of Beneish M-Score models and data mining to detect financial fraud. *Procedia-Social and Behavioral Sciences*, 211, 924-930. <https://doi.org/10.1016/j.sbspro.2015.11.122>
- Hoàng Hà Anh, Trần Minh Dạ Hạnh, Lê Na, & Nguyễn Ngọc Thuỳ. (2022). Đánh giá việc thao túng báo cáo tài chính của các doanh nghiệp trên sàn chứng khoán Thành phố Hồ Chí Minh. *Tạp chí Khoa học Đại học Mở Thành phố Hồ Chí Minh - Kinh tế và Quản trị Kinh doanh*, 18(2), 119-132. <https://doi.org/10.46223/HCMCOUJS.econ.vi.18.2.2203.2023>
- Hołda, A. (2020). Using the Beneish M-score model: Evidence from non-financial companies listed on the Warsaw Stock Exchange. *Investment Management & Financial Innovations*, 17(4), 389-401. [https://doi.org/10.21511/imfi.17\(4\).2020.33](https://doi.org/10.21511/imfi.17(4).2020.33)
- Jan, C. L. (2018). An effective financial statements fraud detection model for the sustainable development of financial markets: Evidence from Taiwan. *Sustainability*, 10(2), 513. <https://doi.org/10.3390/su10020513>
- Jensen, M. C., & Meckling, W. H., (1976). Theory of the firm: Managerial behavior, agency costs and ownership structure. *Journal of Financial Economics*, 3(4), 305-360. [https://doi.org/10.1016/0304-405X\(76\)90026-X](https://doi.org/10.1016/0304-405X(76)90026-X)
- Kaminski, K. A., Sterling Wetzel, T., & Guan, L. (2004). Can financial ratios detect fraudulent financial reporting?. *Managerial Auditing Journal*, 19(1), 15-28. <https://doi.org/10.1108/02686900410509802>
- Kanapickienė, R., & Grundienė, Ž. (2015). The model of fraud detection in financial statements by means of financial ratios. *Procedia-Social and Behavioral Sciences*, 213, 321-327. <https://doi.org/10.1016/j.sbspro.2015.11.545>
- Kinney Jr, W. R., & McDaniel, L. S. (1989). Characteristics of firms correcting previously reported quarterly earnings. *Journal of Accounting and Economics*, 11(1), 71-93. [https://doi.org/10.1016/0165-4101\(89\)90014-1](https://doi.org/10.1016/0165-4101(89)90014-1)
- Kirkos, E., Spathis, C., & Manolopoulos, Y. (2007). Data mining techniques for the detection of fraudulent financial statements. *Expert Systems With Applications*, 32(4), 995-1003. <https://doi.org/10.1016/j.eswa.2006.02.016>

- Kotsiantis, S., Koumanakos, E., Tzelepis, D., & Tampakas, V. (2006). Forecasting fraudulent financial statements using data mining. *International Journal of Computational Intelligence*, 3(2), 104-110.
- Lin, C. C., Chiu, A. A., Huang, S. Y., & Yen, D. C. (2015). Detecting the financial statement fraud: The analysis of the differences between data mining techniques and experts' judgments. *Knowledge-Based Systems*, 89, 459-470. <https://doi.org/10.1016/j.knosys.2015.08.011>
- Lin, J. W., Hwang, M. I., & Becker, J. D. (2003). A fuzzy neural network for assessing the risk of fraudulent financial reporting. *Managerial Auditing Journal*, 18(8), 657-665. <https://doi.org/10.1108/02686900310495151>
- Liou, F. M. (2008). Fraudulent financial reporting detection and business failure prediction models: a comparison. *Managerial Auditing Journal*, 23(7), 650-662. <https://doi.org/10.1108/02686900810890625>
- Liu, C., Chan, Y., Kazmi, S. H. A., & Fu, H. (2015). Financial fraud detection model: Based on random forest. *International Journal of Economics and Finance*, 7(7), 178-188. <http://dx.doi.org/10.5539/ijef.v7n7p178>
- Lokanan, M., Tran, V., & Vuong, N. H. (2019). Detecting anomalies in financial statements using machine learning algorithm: The case of Vietnamese listed firms. *Asian Journal of Accounting Research*, 4(2), 181-201. <https://doi.org/10.1108/AJAR-09-2018-0032>
- Lotfi, N., & Chadegani, A. A. (2018). Detecting corporate financial fraud using Beneish M-score model. *International Journal of Finance & Managerial Accounting*, 2(8), 29-34.
- Mahesh, B. (2020) Machine Learning Algorithms—A Review. *International Journal of Science and Research*, 9(1), 381-386.
- Marais, A., Vermaak, C., & Shewell, P. (2023). Predicting financial statement manipulation in South Africa: A comparison of the Beneish and Dechow models. *Cogent Economics & Finance*, 11(1), 2190215. <https://doi.org/10.1080/23322039.2023.2190215>
- Mishra, C., & Drtina, R. (2004). Accounting manipulations and business failures: the case for effective financial disclosure and corporate governance. *The Journal of Private Equity*, 7(4), 27-35.
- Nguyen, H. A., & Nguyen, H. L. (2016). Using the M-score model in detecting earnings management: Evidence from non-financial Vietnamese listed companies. *VNU Journal of Economics and Business*, 32(2), 14-23.
- Nguyen, A. P., Phan, H. T., & Ngo, P. T. (2022). Fraud identification of financial statements by machine learning technology: case of listed companies in Vietnam. In *International Econometric Conference of Vietnam* (pp. 425-436). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-030-98689-6_28
- Nguyễn Tiến Hùng, Huỳnh Văn Sáu, & Nguyễn Trí Dũng. (2018). Gian lận báo cáo tài chính tại các doanh nghiệp niêm yết trên Sở Giao dịch Chứng khoán Thành phố Hồ Chí Minh. *Tạp chí Khoa học ĐHQGHN: Kinh tế và Kinh doanh*, 34(4), 45-55. <https://doi.org/10.25073/2588-1108/vnueab.4129>
- Nguyễn Tiến Hùng, & Võ Hồng Đức. (2017). Nhận diện gian lận báo cáo tài chính: Bằng chứng thực nghiệm tại các doanh nghiệp niêm yết ở Việt Nam. *Tạp chí Công nghệ Ngân hàng*, 132(5), 58-72. <https://doi.org/10.63065/ajeb.vn.2017.132.33261>

- Nia, S. H. (2015). Financial ratios between fraudulent and non-fraudulent firms: Evidence from Tehran Stock Exchange. *Journal of Accounting and Taxation*, 7(3), 38-44. <https://doi.org/10.5897/JAT2014.0166>
- Omar, N., Johari, Z. A., & Smith, M. (2017). Predicting fraudulent financial reporting using artificial neural network. *Journal of Financial Crime*, 24(2), 362-387. <https://doi.org/10.1108/JFC-11-2015-0061>
- Omidi, M., Min, Q., Moradinaftchali, V., & Piri, M. (2019). The efficacy of predictive methods in financial statement fraud. *Discrete Dynamics in Nature and Society*, 2019(1), 4989140. <https://doi.org/10.1155/2019/4989140>
- Persons, O. S. (1995). Using financial statement data to identify factors associated with fraudulent financial reporting. *Journal of Applied Business Research*, 11(3), 38. <https://doi.org/10.19030/jabr.v11i3.5858>
- Phạm Thị Mộng Tuyền. (2019). Kết hợp mô hình M-Score Beneish và chỉ số Z-Score để nhận diện khả năng gian lận báo cáo tài chính. *Tạp chí Kế toán & Kiểm toán*, 57-61.
- Ravisankar, P., Ravi, V., Rao, G. R., & Bose, I. (2011). Detection of financial statement fraud and feature selection using data mining techniques. *Decision Support Systems*, 50(2), 491-500. <https://doi.org/10.1016/j.dss.2010.11.006>
- Rosner, R. L. (2003). Earnings manipulation in failing firms. *Contemporary Accounting Research*, 20(2), 361-408. <https://doi.org/10.1506/8EVN-9KRB-3AE4-EE81>
- Shahana, T., Lavanya, V., & Bhat, A. R. (2023). State of the art in financial statement fraud detection: A systematic review. *Technological Forecasting and Social Change*, 192, 122527. <https://doi.org/10.1016/j.techfore.2023.122527>
- Tarjo, & Herawati, N. (2015). Application of Beneish M-Score models and data mining to detect financial fraud. *Procedia-Social and Behavioral Sciences*, 211, 924-930. <https://doi.org/10.1016/j.sbspro.2015.11.122>
- Tarquinio, L., & Posadas, S. C. (2020). Exploring the term “non-financial information”: an academics’ view. *Meditari Accountancy Research*, 28(5), 727-749. <https://doi.org/10.1108/MEDAR-11-2019-0602>
- Thornhill, W.T. (1995). *Forensic Accounting: How to Investigate Financial Fraud*. New York: Irwin Professional Publishing
- Trần Thị Giang Tân, Nguyễn Trí Tri, Đinh Ngọc Tú, Hoàng Trọng Hiệp, & Nguyễn Đình Hoàng Uyên. (2015). Đánh giá rủi ro gian lận báo cáo tài chính của các công ty niêm yết tại Việt Nam. *Tạp chí Phát triển Kinh tế*, 26(1), 74-94.
- Võ Văn Nghi, & Hoàng Cẩm Trang. (2013). Hành vi điều chỉnh lợi nhuận & nguy cơ phá sản của các doanh nghiệp niêm yết trên sở giao dịch chứng khoán TP. HCM. *Tạp chí Phát triển Kinh tế*, 276S, 48-57.
- Zainudin, E. F., & Hashim, H. A. (2016). Detecting fraudulent financial reporting using financial ratio. *Journal of Financial Reporting and Accounting*, 14(2), 266-278. <https://doi.org/10.1108/JFRA-05-2015-0053>