

การสังเคราะห์งานวิจัยแบบจำลองการพยากรณ์ภาวะความล้มเหลวของธุรกิจ

Research Synthesis on Business Failure Prediction Models

วิชชกานต์ เมธาวิริยะกุล*

บทคัดย่อ

บทความนี้มีวัตถุประสงค์เพื่อสังเคราะห์งานวิจัยแบบจำลองการพยากรณ์ภาวะความล้มเหลวของธุรกิจ ได้แก่ กลุ่มแบบจำลองทางสถิติดั้งเดิม (Traditional Statistical Models) คือ (1) การวิเคราะห์จำแนกประเภทหลายตัวแปร (Multivariate Discriminant Analysis: MDA) เป็นแบบจำลองที่ใช้มากที่สุดในวรรณกรรมเกี่ยวกับภาวะล้มละลาย (Bankruptcy Literature) (2) การวิเคราะห์จำแนกประเภทกำลังสอง (Quadratic Discriminant Analysis: QDA) เป็นแบบจำลองที่ซับซ้อนกว่าในเรื่องสมมติฐานทางสถิติ น้อยกว่า (3) แบบจำลองโลจิส (Logit Model) หรือการวิเคราะห์ความถดถอยโลจิสติก (Logistic Regression Analysis) (4) แบบจำลองโพรบิต (Probit Model) ทั้ง 2 แบบจำลองนำเสนอในรูปแบบของความน่าจะเป็น (Probabilistic Approach) ต่อมาคือกลุ่มแบบจำลองการเรียนรู้ของเครื่องจักร (Machine Learning Models) คือ (5) ต้นไม้การตัดสินใจ (Decision Tree: DT) (6) เครือข่ายประสาทเทียม (Artificial Neural Networks: ANN หรือ NN) (7) ซัพพอร์ตเวกแมชชีน (Support Vector Machine: SVM) ซึ่งเป็นแบบจำลองสมัยใหม่ที่มีประสิทธิภาพในการพยากรณ์ที่ถูกต้องแม่นยำสูง และแบบจำลองดังกล่าวไม่มีข้อจำกัดเกี่ยวกับสมมติฐานทางสถิติเหมือนแบบจำลองทางสถิติดั้งเดิม นอกจากนี้ยังได้นำเสนอข้อมูลสถิติเชิงพรรณนาที่ถูกต้องแม่นยำในการพยากรณ์ของแบบจำลองประเภทต่างๆ ในแต่ละยุคสมัยอีกด้วย

* นักศึกษาระดับปริญญาโท หลักสูตรบัญชีบัณฑิต สำนักวิชาบัญชี มหาวิทยาลัยราชภัฏเชียงใหม่

บทนำ

วัตถุประสงค์ของการดำเนินธุรกิจ (Objectives of Business) สามารถจำแนกได้ 5 ประเภท ได้แก่ (1) ด้านเศรษฐกิจ (Economic Objectives) เช่น รายได้จากการขาย การสร้างลูกค้า การสร้างนวัตกรรม และการใช้ทรัพยากรอย่างมีประสิทธิภาพ (2) ด้านสังคม (Social Objectives) เช่น การผลิตและจัดหาสินค้าและบริการที่มีคุณภาพ การคุ้มครอง การค้าอย่างยุติธรรม และการมีส่วนร่วมในสวัสดิการของสังคม (3) ด้านทรัพยากรมนุษย์ (Human Objectives) เช่น การมีสภาพเศรษฐกิจที่ดีของพนักงาน ความพึงพอใจทางสังคม และจิตวิทยาของพนักงาน การพัฒนาทรัพยากรมนุษย์ และสภาพเศรษฐกิจที่ดีของสังคม และของคนรุ่นต่อไป (4) ด้านประเทศชาติ (National Objectives) เช่น สร้างการจ้างแรงงาน การสร้างความเป็นธรรม การสร้างความมั่นคงให้ประเทศชาติ การสร้างรายได้ให้ประเทศชาติ และเกิดการพึ่งตนเองและส่งเสริมการส่งออก (5) ด้านต่อโลก (Global Objectives) เช่น ยกระดับความเป็นอยู่ที่มีมาตรฐานระดับโลก ช่องว่างหรือความแตกต่างในบรรดาประชาชาติ และเกิดการแข่งขันระดับโลกเกี่ยวกับสินค้าและบริการ (National Institute of Open Schooling, 2016)

ดังนั้นจะเห็นได้ว่าหน่วยธุรกิจ (Business Unit) จึงมีความสำคัญเป็นอย่างมากต่อระบบเศรษฐกิจ ในการเป็นหน่วยที่ก่อให้เกิดการผลิตสินค้าและบริการที่มีคุณภาพ ก่อให้เกิดการจ้างแรงงาน การกระจายรายได้อย่างเป็นธรรม การส่งออกสินค้าและบริการเพื่อให้เกิดดุลการค้ากับต่างประเทศ รวมไปถึงการเพิ่มมูลค่าของผลิตภัณฑ์มวลรวมในประเทศ (GDP) แต่จากวิกฤตการณ์ทางเศรษฐกิจและการเงินที่เกิดขึ้นอย่างต่อเนื่องทั้งในและต่างประเทศ ภาวะเศรษฐกิจในประเทศที่ถดถอย ปัญหาเสถียรภาพทางการเงินภายในประเทศ ปัญหากิจการดำเนินธุรกิจและการกำกับดูแลธุรกิจที่ไม่มีประสิทธิภาพที่เกิดจากสาเหตุที่ควบคุมได้ (Controllable Factors) หรือควบคุมไม่ได้ (Uncontrollable Factors) ทำให้ส่งผลกระทบต่อความอยู่รอดและความล้มเหลวของหน่วยธุรกิจ หากเกิดภาวะล้มเหลวของหน่วยธุรกิจ จะส่งผลกระทบอย่างกว้างขวาง ไม่ว่าจะเป็นผู้ลงทุน เจ้าหนี้ พนักงาน ผู้สอบบัญชี ประชาชนทั่วไป รวมไปถึงการสร้าง ความเสียหายต่อระบบเศรษฐกิจของทั้งในและต่างประเทศอย่างมหาศาล ถึงแม้ว่าภาวะความล้มเหลวอาจจะเป็นเหตุการณ์ที่ไม่เกิดขึ้นบ่อยครั้ง แต่เป็นเหตุการณ์ที่มีต้นทุนสูงสำหรับเจ้าของเงินทุน

เนื่องจากการปรับโครงสร้างองค์กรหรือต้นทุนในการเลิกกิจการ จะไปลดมูลค่าส่วนใหญ่ขององค์กร (Beaver, 1968) จากผลกระทบดังกล่าวทำให้สังคมเริ่มตระหนักถึงความสำคัญและความจำเป็นของเครื่องมือหรือแบบจำลองที่จะส่งสัญญาณเตือนภัยล่วงหน้าทางการเงิน (Financial Early Warning Models) จึงทำให้นักวิจัยเกิดแนวคิดในการที่จะพยายามคิดค้นหาและพัฒนาเครื่องมือหรือแบบจำลองที่จะใช้ในการพยากรณ์หรือทำนายภาวะล้มเหลวของบริษัท (Business Failures Prediction Models) เพื่อให้องค์กรสามารถประเมินสถานการณ์หรือคาดคะเนถึงสาเหตุ และสามารถเตรียมรับมือ หรือปรับเปลี่ยนวิธีการดำเนินงานของบริษัทเพื่อให้รอดพ้นจากภาวะความล้มเหลวที่อาจจะเกิดขึ้นได้ทันทั่วทั้งที่และที่สำคัญอีกประการหนึ่งคือประโยชน์ที่ได้จากการพยากรณ์ภาวะล้มเหลวทางการเงินของธุรกิจมีความสำคัญอย่างมากต่อการเจริญเติบโตอย่างยั่งยืนของบริษัท (ประเสริฐสิทธิ์ หาวาสน์ & มนวิภา ผดุงสิทธิ์, 2552; ปานรดา พิลาศรี, 2553; สิทธิพล สมชม & ณัฐวุฒิ คุ้มพัฒนเชียรชัย, 2557)

แบบจำลองการพยากรณ์ภาวะความล้มเหลว

นักวิจัยในทศวรรษที่ผ่านมาได้ตระหนักว่าภาวะความล้มเหลว (Failure) ไม่ได้เกิดขึ้นอย่างกะทันหัน แต่จะมีการส่งสัญญาณเตือนล่วงหน้า โดยทั่วไปแล้วภาวะความล้มเหลวมักจะใช้เวลาหลายปี ก่อนที่จะมีการล้มละลายอย่างเป็นทางการ (Official Bankruptcy) ในรูปแบบของการปรับโครงสร้างองค์กร (Re-organization) หรือการชำระบัญชี (Liquidation) ดังนั้นจึงมีความจำเป็นอย่างยิ่งที่จะต้องมีการพัฒนาแบบจำลองสัญญาณเตือนภัยล่วงหน้า เพื่อจะประเมินสุขภาพทางการเงินของบริษัท (Financial Health of Corporations) ซึ่งได้แก่ จุดแข็ง (Strengths) และจุดอ่อน (Weaknesses) ทางการเงินของบริษัท ทำให้มีการพัฒนาแบบจำลองการพยากรณ์ภาวะความล้มเหลว (Failure Prediction Models) ขึ้นอย่างมากมายโดยใช้เทคนิค (Techniques) และวิธีการ (Methods) ที่แตกต่างกันไป การวิเคราะห์อัตราส่วนทางการเงิน (Financial Ratio Analysis) เป็นวิธีการที่ใช้โดยทั่วไปในการที่จะวิเคราะห์ความเข้มแข็งและความอ่อนแอทางการเงินของบริษัท (Bal, Cheung, & Wu, 2013; Kpodoh, 2009)

ตามที่เชื่อกันโดยทั่วไปว่าข้อมูลทางการเงินถือว่าเป็นองค์ประกอบที่สำคัญและสามารถเข้าถึงได้มากที่สุดในการตรวจสอบประสิทธิภาพการทำงานของบริษัท และในการเป็นข้อมูลเพื่อการทำนายแนวโน้มของการล้มเหลว แบบจำลองการพยากรณ์ภาวะความล้มเหลวของธุรกิจส่วนใหญ่อยู่บนพื้นฐานของข้อมูลทางการเงิน โดยข้อมูลทางการเงินถือเป็นอาการ (Symptoms) ของภาวะความล้มเหลวมากกว่าสาเหตุ (Causes) (Bal, Chemsirakul & We, 2013 อ้างอิงจาก Argenti, 1976) ฉะนั้นหมายความว่าตัวเลขทางการเงินและการถนำมาพิจารณาเป็นตัวบ่งชี้ในการพยากรณ์หรือทำนายความเป็นไปได้ของภาวะความล้มเหลว มีงานศึกษาจำนวนมากที่ใช้เทคนิคทางสถิติ (Statistical Techniques) ในการพัฒนาร่วมกับอัตราส่วนทางการเงินที่จะพยากรณ์หรือคาดการณ์ภาวะความล้มเหลวของธุรกิจ โดยเทคนิคทางสถิติที่ได้รับความนิยมมากที่สุดคือแบบจำลองการวิเคราะห์จำแนกประเภทหลายตัวแปร (Multivariate Discriminant Models) และแบบจำลอง โลจิสติก (Logistic Models) (Bal et al., 2013) ดังแสดงในตารางที่ 1

ตารางที่ 1 สรุปวิธีการของการพยากรณ์ภาวะความล้มเหลว (Summary of Failure Prediction Methods)

	นักวิจัย (Author)	Year	วิธีการ (Method)
การศึกษารุ่นก่อน	Tam and Tesser	1966	Index of Risk
Early Studies	Beaver	1966	Univariate Analysis
	Altman	1968	Multivariate Discriminant Analysis
	Taffler and Tisshaw	1977	Multivariate Discriminant Analysis
ตัวแปรเชิงเฉพาะเจาะจง	Mason and Harris	1979	Multivariate
			Discriminant Analysis

ตารางที่ 1 สรุปวิธีการของการพยากรณ์ภาวะความล้มเหลว (Summary of Failure Prediction Methods) (ต่อ)

	นักวิจัย (Author)	Year	วิธีการ (Method)
Construction-specific Models	Abidali and Harris	1995	Multivariate Discriminant Analysis
	Russell and Jaselskis	1992	Logit Analysis
	Severson and Russell	1994	Logit Analysis
ตัวแบบผสมผสาน	Zaverger	1985	Logit Analysis and Expert System
Hybrid Models	Keasey and McGuinness	1992	Logit Analysis and Entropy

ที่มา: ดัดแปลงจาก (Bal et al., 2013)

นอกจากวิธีการพยากรณ์ภาวะความล้มเหลวดังกล่าวตามตารางที่ 1 ยังมีตัวแบบระบบผู้เชี่ยวชาญอัจฉริยะ (Artificially Intelligent Expert System Models: AIES) ซึ่งเป็นวิธีการพยากรณ์ที่ได้รับความนิยมมากขึ้นเรื่อยๆ รายละเอียดแสดงในตารางที่ 2

ตารางที่ 2 ตัวแบบระบบผู้เชี่ยวชาญอัจฉริยะชนิดต่างๆ (Different Types of AIES Models)

ตัวแบบ (Models)	นักวิจัย (Author)	ปี (Year)
Recursively Partitioned Decision Trees	Friedman	1977
(An Inductive Learning Model)	Pompe and Feelders	1993
Case-based reasoning (CBR) models	Kolodner	1993
Neural Networks (NN)	Salchenberger et al.	1992
	Coats and Fant	1993
	Yang et al.	1999
Genetic Algorithms (GA)	Shin and Lee	2002
	Varetto	1998
Rough Sets Model	Pawlak	1982
	Ziarko	1993
	Yimitras et al.	1999

ที่มา: (Aziz & Dar, 2006)

ในบทความนี้จะอธิบายถึงเฉพาะตัวแบบการพยากรณ์ภาวะความล้มเหลวที่นิยมทั่วไป ได้แก่ การวิเคราะห์จำแนกประเภทหลายตัวแปร (Multivariate Discriminant Analysis: MDA) การวิเคราะห์กำลังสอง (Quadratic Discriminant Analysis: QDA) แบบจำลองลอจิสติก (Logit Model) หรือการวิเคราะห์ความถดถอยโลจิสติก (Logistic Regression Analysis) แบบจำลองโพรบิต (Probit Model) ต้นไม้การตัดสินใจ (Decision Tree: DT) เครือข่ายประสาทเทียม (Artificial Neural Networks: ANN หรือ NN) และ ماشینเวกเตอร์แมชชีน (Support Vector Machine: SVM) รายละเอียดดังต่อไปนี้ (Okay, 2015)

1. การวิเคราะห์จำแนกประเภทหลายตัวแปร (Multivariate Discriminant Analysis: MDA)

MDA ได้ปรับปรุงมาจากการวิเคราะห์ตัวแปรเดียว (Univariate Analysis: UA) ของ Beaver (Beaver, 1966) โดย Altman ในปี 1968 เป็นที่รู้จักและยังคงใช้กันโดยทั่วไปในชื่อ Altman's Z-Score ข้อดีของ MDA ที่เหนือกว่า UA คือการนำตัวแปรเข้ามาพิจารณาและวิเคราะห์ในการพยากรณ์ Altman ได้ใช้ธุรกิจการผลิต (Manufacturing Corporations) ในประเทศสหรัฐอเมริกาเป็นตัวอย่างจำนวน 66 บริษัท แบ่งเป็น 2 กลุ่ม คือ กลุ่มบริษัทที่ล้มเหลวทางการเงินจำนวน 33 บริษัท และกลุ่มบริษัทที่ไม่ล้มเหลวจำนวน 33 บริษัท ระหว่างปี 1946-1965 ตัวแปรทางการเงิน (อัตราส่วนทางการเงิน) ที่ส่งผลต่อการพยากรณ์มีจำนวน 5 ตัวแปร คือ (1) อัตราส่วนเงินทุนหมุนเวียนต่อสินทรัพย์รวม (Net Working Capital/Total Assets) (2) อัตราส่วนกำไรสะสมต่อสินทรัพย์รวม (Retained Earning/Total Assets) (3) กำไรก่อนหักดอกเบี้ยและภาษีต่อสินทรัพย์รวม (Earnings before Interest and Taxes/Total Assets) (4) อัตราส่วนมูลค่าตลาดของหุ้นต่อราคาตามบัญชีของหนี้สินรวม (Market Value of Equity/Book Value of Liabilities) และ (5) อัตราส่วนยอดขายต่อสินทรัพย์รวม (Sales/Total Assets) ตัวแบบดังกล่าวสามารถพยากรณ์ได้ถูกต้องแม่นยำกว่า 1-3 ปีก่อนล้มเหลว ร้อยละ 95 ร้อยละ 72 ร้อยละ 48 ร้อยละ 29 และร้อยละ 36 ตามลำดับ (Altman, 1968)

MDA มีวัตถุประสงค์เพื่อหาการร่วมกันเชิงเส้นของตัวพยากรณ์ เพื่อเพิ่มความแปรปรวนระหว่างบริษัทที่ล้มเหลวกับบริษัทที่ไม่ล้มเหลว ในขณะที่ลดความแปรปรวนภายในบริษัทที่ล้มเหลวกับบริษัทที่ไม่ล้มเหลว โดยแต่เดิมได้ถูกใช้ในการจำแนกประเภทปัญหาของตัวแปรที่อยู่ในรูปแบบเชิงคุณภาพ รูปแบบพื้นฐานของตัวแบบปรากฏในสมการดังนี้

$$Z = \beta_0 + \beta_1 X_{1j} + \beta_2 X_{2j} + \dots + \beta_n X_{nj}$$

$\beta_i (i = 1, 2, \dots, n) = \text{Coefficient (Discriminant) Weights}$

$X_i (i = 1, 2, \dots, n) = \text{Independent Variables (Financial Ratios)}$

ความหมายของ Z-Scores แต่ผลของวิธี Conditional Logit Analysis สามารถตีความหมายโดยความน่าจะเป็นของภาวะล้มละลายภายใต้ข้อสมมติฐานจำกัดที่น้อยกว่า และได้ใช้จุดตัด (Cutoff Point) ที่คำนวณเพื่อให้ได้ค่าในการจำแนกผิดประเภทต่ำที่สุด (Minimize the Costs of Misclassification) ในการจำแนกประเภทของบริษัทที่ล้มละลายหรือบริษัทที่ไม่ล้มละลาย (Bankrupt or Non-Bankrupt) Ohlson ได้ใช้ตัวอย่างบริษัทอุตสาหกรรม (Industrial Companies) ในประเทศสหรัฐอเมริกาจำนวนมาก โดยแบ่งกลุ่มตัวอย่างเป็น 2 กลุ่ม ได้แก่ กลุ่มบริษัทที่ล้มละลายจำนวน 105 บริษัท และกลุ่มบริษัทที่ไม่ล้มละลายจำนวน 2,058 บริษัท ระหว่างปี 1970-1976 โดยไม่ได้จับคู่บริษัท และตัวแปรที่ส่งผลกระทบต่อผลการพยากรณ์ ได้แก่ (1) Log ของอัตราส่วนสินทรัพย์รวมต่อ GNP (การลอกลายตัวของราคา [Natural Log of (Total Assets/GNP Price-Level Index)]) (2) อัตราส่วนหนี้สินรวมต่อสินทรัพย์รวม (Total Liabilities/Total Assets) (3) อัตราส่วนเงินทุนหมุนเวียนต่อสินทรัพย์รวม (Working Capital/Total Assets) (4) อัตราส่วนหนี้สินหมุนเวียนต่อสินทรัพย์หมุนเวียน (Current Liabilities/Current Assets) (5) อัตราส่วนกำไรสุทธิต่อสินทรัพย์รวม (Net Income/Total Assets) (6) อัตราส่วนเงินทุนจากการดำเนินงานต่อหนี้สินรวม (Funds Provided by Operations/Total Liabilities) (7) กำหนดค่า 1 ถ้ากำไรสุทธิ 2 ปีล่าสุดติดลบ นอกนั้นให้กำหนดค่า 0 (One if net income was negative for the last two years, zero otherwise) (8) กำหนดค่า 1 ถ้าหนี้สินรวมมากกว่าสินทรัพย์รวม นอกนั้นให้กำหนดค่า 0 (One if total liabilities exceeds total assets, zero otherwise) และ (9) ค่าอัตราส่วนของกำไรสุทธิเวลาปัจจุบันหักกำไรสุทธีย้อนหลัง 1 ปีต่อค่าสัมบูรณ์ของกำไรสุทธิเวลาปัจจุบันบวกค่าสัมบูรณ์ของกำไรสุทธีย้อนหลัง 1 ปี $(NIt - NIt - 1) / (|NIt| + |NIt - 1|)$, where NIt is net income for the most recent period) ตัวแบบดังกล่าวสามารถพยากรณ์ได้ถูกต้องแม่นยำภายใน 1-2 ปี ร้อยละ 96.12 และร้อยละ 95.55 ตามลำดับ (Ohlson, 1980)

Logit มีวัตถุประสงค์เพื่อพยากรณ์ผลลัพธ์ที่ไม่ต่อเนื่อง (Discrete Outcome) ได้จากสมาชิกกลุ่มจากชุดของตัวแปร ซึ่งอาจเป็นตัวแปรต่อเนื่อง (Continuous) ตัวแปรไม่ต่อเนื่อง (Discrete) ตัวแปรทวิ (Dichotomous) หรือตัวแปรผสม (Mix) ตัวแบบปรากฏในสมการดังนี้

$$P = \frac{1}{1 + e^{\beta^T x}}$$

$\beta_i (i = 1, 2, \dots, n) = \text{Coefficient (Discriminant) Weights}$

$X_i (i = 1, 2, \dots, n) = \text{Independent Variables (Financial Ratios)}$

4. แบบจำลองโพรบิท (Probit Model)

ในปี 1983 Zmijewski ได้พัฒนาแบบจำลองการทำนายภาวะล้มละลายโดยใช้วิธี Probit Analysis ซึ่งได้เลือกอัตราส่วนทางการเงินจำนวน 3 อัตราส่วนโดยใช้อัตราส่วนอยู่ในกลุ่มต่างๆ กัน ได้แก่ Profitability Ratio, Financial Leverage Ratio และ Liquidity Ratio และเลือกตัวอย่างจำนวน 840 บริษัท แบ่งเป็นบริษัทที่ประสบปัญหาทางการเงินจำนวน 40 บริษัท และบริษัทที่ไม่ประสบปัญหาทางการเงินจำนวน 800 บริษัท ระหว่างปี 1972-1978 วิเคราะห์หาค่าทางสถิติที่เรียกว่า การวิเคราะห์ความสัมพันธ์ (Correlation) เพื่อหาความสัมพันธ์ระหว่างลักษณะของอัตราส่วนทางการเงินและหาค่าสัมประสิทธิ์ ต่อมาในปี 1984 ได้ทำการปรับปรุงแบบจำลองโดยการใช้ถ่วงน้ำหนัก (Weighted Original) อย่างไรก็ตามพบว่าทั้ง 2 แบบจำลองยังไม่แตกต่างกันมากนัก อัตราส่วนที่มีผลต่อการพยากรณ์คือ อัตราส่วนกำไรสุทธิต่อสินทรัพย์รวม (Net Income/Total Assets) อัตราส่วนหนี้สินรวมต่อสินทรัพย์รวม (Total Liabilities/Total Assets) และอัตราส่วนสินทรัพย์หมุนเวียนต่อหนี้สินหมุนเวียน (Current Assets/Current Liabilities) ซึ่งผลการวิจัยของ Zmijewski พบว่าแบบจำลองยังสามารถทำนายภาวะความล้มเหลวทางการเงินได้อย่างถูกต้องแม่นยำร้อยละ 92 (อภิสิทธิ์ อดทน, 2553; อ้างอิงจาก Zmijewski, 1983, 1984) (อภิสิทธิ์ อดทน, 2553)

ความแตกต่างระหว่างการเดิยวาระหว่างแบบจำลองโลจิสติกกับแบบจำลองโพรบิทคือ ฟังก์ชันเชื่อมต่อ (Link Function) แบบจำลองโพรบิทใช้ฟังก์ชันสะสมปกติ (Normal Cumulative Function) แทนฟังก์ชันโลจิสติก (Logistic Function) ส่วนการอธิบายประเด็นอื่นๆ เหมือนกับแบบจำลองโลจิสติก ตัวแบบปรากฏในสมการดังนี้

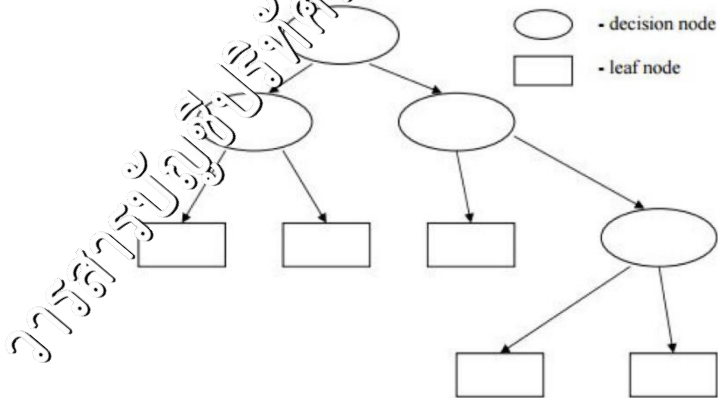
$$P = \int_{-\infty}^{\beta^i x} \frac{1}{\sqrt{2\pi}} e^{-t^2/2} dt = \Phi(\beta^i x)$$

5. ต้นไม้การตัดสินใจ (Decision Tree: DT)

DT ได้ใช้ Recursive Partitioning Technique ในการแยก The Observations เป็น Subsets กระบวนการจะดำเนินต่อเนื่องจนกระทั่ง Recursive Partitioning ไม่มีมูลค่าเพิ่มในการพยากรณ์ (No Longer Adds Value to Predictions) โมเดล DT มีประโยชน์หรือข้อได้เปรียบเหนือโมเดลทางสถิติดั้งเดิม (Classical Statistical Models) โดย DT ไม่มีข้อจำกัดทางสถิติ ได้แก่ ความไม่เป็นเชิงเส้น (Non-Linearity) และ การแจกแจงแบบไม่ปกติ (Non-Normality) และ DT มีรูปแบบและวิธีการที่เรียบง่ายต่อการนำไปใช้งาน

ในปี 1998 Joos และคณะ ได้เปรียบเทียบ DT และ Logit ในสภาพแวดล้อมของการจำแนกประเภทเครดิต (Credit Classification Environment) โดยใช้ An extensive Database ซึ่งเป็นหนึ่งในฐานข้อมูลที่ใหญ่ที่สุดของธนาคารเบลเยียม (Belgian Banks) และพบว่า Logit เหนือกว่า DT ในแง่ของประสิทธิภาพ ต้นทุนและคุณภาพของชุดข้อมูล (Cost Efficiency and Qualitative Data Set) และ DT เหนือกว่า Logit

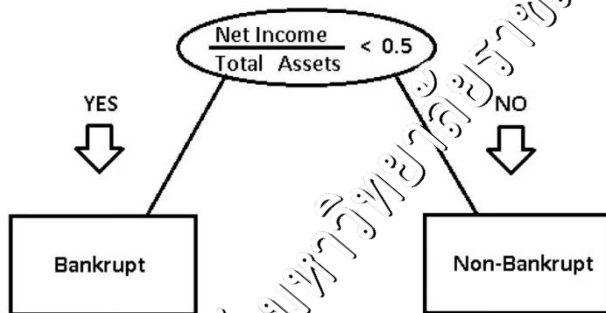
ส่วนใหญ่ DT จะใช้ในการแก้ปัญหาการจำแนกประเภทหรือการจัดหมวดหมู่ (Classification) ซึ่งประกอบด้วยโหนดการตัดสินใจ (Decision Nodes) และโหนดใบไม้ (Leaf Nodes) ทุกโหนดการตัดสินใจจะทำการวางทดสอบคุณลักษณะเฉพาะของข้อมูล รายละเอียดตามภาพที่ 1



ภาพที่ 1 ตัวอย่างต้นไม้การตัดสินใจ (Decision Tree Example)

ที่มา: (Okay, 2015)

ตัวอย่างที่เรียบง่ายของ DT เกี่ยวกับขั้นตอนวิธีการของการจำแนกประเภทของโมเดล ปรากฏในภาพที่ 2 โมเดลจะทำการตัดสินใจแยกเกี่ยวกับอัตราส่วนกำไรสุทธิต่อสินทรัพย์รวม (Net Income to Total Assets) ที่ 0.5 ดังนั้นถ้าบริษัทมีอัตราส่วนดังกล่าวต่ำกว่า 0.5 โมเดลจะแยกประเภทเป็นบริษัทที่ล้มละลาย (Bankrupt) แต่ถ้าอัตราส่วนดังกล่าวสูงกว่า 0.5 โมเดลจะแยกประเภทเป็นบริษัทที่ไม่ล้มละลาย (Non-Bankrupt) สามารถมีเงื่อนไขการตัดสินใจที่ซับซ้อน (Complex Decision Rules) ได้ คำถามว่าการสำคัญคือ ตัวแปรตาม (Explanatory Variable) ตัวไหนที่ควรใช้และจะใช้ตัวไหนที่จะใช้แยก (Which Value the Split Should Occur)



ภาพที่ 2 ตัวอย่างการจำแนกประเภทของต้นไม้ตัดสินใจ (Decision Tree Classification Example)

ที่มา: (Okay, 2015)

6. โครงข่ายประสาทเทียม (Artificial Neural Networks: ANN หรือ NN)

NN ได้รับแรงบันดาลใจจาก Biological Neural Networks ของระบบประสาทของมนุษย์ (The Human Nervous System) NN เรียนรู้จากสอนตัวอย่าง (Training Examples) โดยใช้อัลกอริทึมที่แตกต่างกัน (Different Algorithms) เหมือนการเรียนรู้ของมนุษย์ โครงสร้างของโมเดลมาจากข้อมูลจากการสอนและโดยทั่วไปจะใช้สมการที่ไม่เป็นเส้นตรง (Non-Linear Equations) ในการประมาณค่าส่งออก (Outputs) ซึ่งขึ้นอยู่กับตัวแปรนำเข้าหลายตัว (Many Inputs) NN ไม่มีข้อจำกัดเกี่ยวกับสมมติฐานทางสถิติ

(Statistical Assumptions) เหมือนโมเดลทางสถิติดั้งเดิม (Traditional Statistical Models)

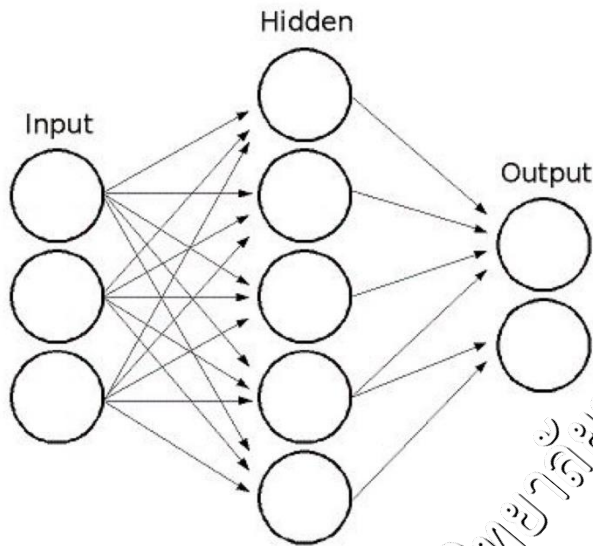
ในปี 1990 Odom และ Sharda เป็นผู้บุกเบิกในการใช้ NN เพื่อการพยากรณ์ภาวะล้มละลาย (Prediction of Bankruptcy) โดยใช้ตัวแปรทางการเงินเดียวกับ Altman's Z-Score ตัวอย่างจำนวน 129 บริษัท ซึ่งประกอบไปด้วยบริษัทที่ล้มละลายจำนวน 65 บริษัท และบริษัทที่ไม่ล้มละลายจำนวน 64 บริษัท ในระหว่างปี 1975-1986 โดยใช้ 74 บริษัทในการสอน (to train) และใช้ 55 บริษัทในการทดสอบ (to test) และใช้ MDA เป็นโมเดลเปรียบเทียบ (Benchmark Model) พบว่า NN มีความถูกต้องแม่นยำในการพยากรณ์ร้อยละ 82 สำหรับตัวอย่างทดสอบ (Holdout Sample) ในขณะที่ MDA มีความถูกต้องแม่นยำในการพยากรณ์ร้อยละ 74

ในปี 1993 Fletcher และ Goss ได้เปรียบเทียบ NN กับ Logit โดยใช้ตัวอย่างจำนวน 36 บริษัทที่ล้มละลายและบริษัทที่ไม่ล้มละลาย ซึ่งประกอบไปด้วย 3 อัตราส่วนทางการเงิน ได้แก่ Current Ratio, Quick Ratio และ Working Capital to Net Income พบว่า NN มีความถูกต้องแม่นยำในการพยากรณ์สูงกว่า Logit เกือบทุกจุดตัด (At Almost all Cutoff Points)

ในปี 1999 Zhang และคณะ ได้เปรียบเทียบประสิทธิภาพของ NN กับ Logit Model โดยใช้ตัวอย่างบริษัทการผลิต (Manufacturing Companies) พบว่า NN มีประสิทธิภาพเหนือกว่าอย่างมีนัยสำคัญเมื่อเทียบกับ The Logit Regression Model โดยมีความถูกต้องแม่นยำสำหรับการทดสอบขนาดเล็ก (Small Test Set) NN ร้อยละ 80 ในขณะที่ Logit ร้อยละ 74 และสำหรับกลุ่มทดสอบขนาดใหญ่ (Large Test Set) NN ร้อยละ 86 ในขณะที่ Logit ร้อยละ 78

NN เกี่ยวข้องกับชุดตัวแปรนำเข้า (Input Variables) จนถึงตัวแปรส่งออก (Output Variables) ซึ่งประกอบไปด้วยชั้นตัวแปรนำเข้า (Input Layers) หรือตัวแปรอิสระ (Independent Variables) ชั้นซ่อน (Hidden Layer) จะแปลงหรือคำนวณข้อมูลจากตัวแปรนำเข้า และชั้นส่งออก (Output Layers) คือผลลัพธ์จากการคำนวณของชั้นซ่อน ความแตกต่างที่สำคัญจากแบบจำลองอื่น ๆ คือแบบจำลองโครงข่ายประสาทเทียมมีชั้นซ่อน (Hidden Layers) ที่ซึ่งตัวแปรนำเข้าจะแปลงโดยชั้นที่พิเศษของอัลกอริทึม ด้วยการมีชั้นซ่อนนี้ก่อให้เกิดประสิทธิภาพในการประมวลแบบจำลองที่ไม่เป็นเชิงเส้น (Non-linear)

รายละเอียดตามภาพที่ 3



ภาพที่ 3 แบบจำลองโครงข่ายประสาทเทียม (Neural Networks Model)

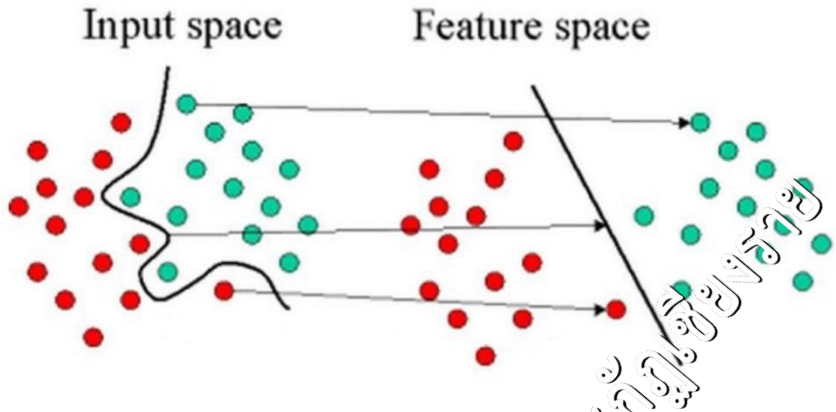
ที่มา: (Okay, 2015)

7. ซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM)

SVM เป็นอีกหนึ่งโมเดลการเรียนรู้ของเครื่องจักร (Machine Learning Model) วารณกรรมที่เกี่ยวข้องกับ SVM มีดังนี้ เมื่อเทียบกับ NN หรือ Traditional Statistical Models

ในปี 2005 Sung, Lee และ Kim ได้ใช้ SVM และเปรียบเทียบผลลัพธ์กับ NN พบว่า SVM มีประสิทธิภาพเหนือกว่า NN และในปีเดียวกัน Min และ Lee ได้เปรียบเทียบประสิทธิภาพของ SVM กับ MDA, Logit และ NN พบว่า SVM เหนือกว่าโมเดลอื่นๆ

SVM คือการผสมผสานของฟังก์ชันเชิงเส้นและอัลกอริทึมการเรียนรู้ ที่จะเลือกจำนวนจุดขอบเขตที่สำคัญ ซึ่งเรียกว่า Support Vectors โดยทั่วไป SVM ใช้ฟังก์ชันเชิงเส้น (Linear Function) ในการจำแนกข้อมูลโดยการหาแผนที่เวกเตอร์ของตัวแปรนำเข้าที่ไม่เป็นเชิงเส้น (Non-linear Input vectors) เข้าไปในพื้นที่หลายมิติ (Higher Dimensional Feature Space) รายละเอียดตามภาพที่ 4



ภาพที่ 4 แบบจำลอง SVM (SVM Model)

ที่มา: (Okay, 2015)

จากภาพที่ 4 ได้แสดงให้เห็นการทำงานของ SVM ซึ่งมีพื้นที่นำเข้าไป 2 มิติ จะเห็นว่าขอบเขตที่จำแนก 2 กลุ่มไม่ได้เป็นเชิงเส้น (Linear) โดยจะเริ่มจากการหาพื้นที่ที่ไม่มีตัวแยกเชิงเส้นที่มีลักษณะเด่นแยกกัน จึงจะเพิ่มขอบเขตระหว่างกลุ่ม SVM จะทำให้พื้นที่ที่หลากหลายมิติ (Higher Dimensional Feature Space) ที่กลุ่มจะกลายเป็นการแยกเชิงเส้น (Linearly Separable) โดยการใช้เทคนิค Kernel

ในปี 2015 Okay ได้ศึกษาความล้มเหลวของบริษัทที่ไม่ใช่การเงินในประเทศตุรกีระหว่างปี 2000-2015 และได้เปรียบเทียบความถูกต้องแม่นยำของโมเดลพยากรณ์แบบต่างๆ ได้แก่ MDA, Logistic, Probit, DT, NN และ SVM ตัวอย่างบริษัทจำนวน 64 บริษัท ได้แก่ บริษัทที่มีผลประกอบการประกอบด้วย 32 บริษัทตุรกีที่ถูกเพิกถอนจาก The Borsa Istanbul Market โดยจับคู่กับบริษัทที่ไม่ล้มเหลวที่มีขนาดสินทรัพย์รวมและปีใกล้เคียงกันจำนวน 32 บริษัท และใช้อัตราส่วนทางการเงิน (Financial Ratios) จำนวน 16 อัตราส่วนในการศึกษา

พบว่าตัวแปรทางการเงิน (Accounting Variables) ที่มีนัยสำคัญต่อการพยากรณ์ความล้มเหลวของบริษัทก่อน 1-2 ปี ได้แก่ เงินทุนหมุนเวียนต่อสินทรัพย์รวม (Working Capital to Total Assets) กำไรสุทธิต่อสินทรัพย์รวม (Net Income to Total Assets) และกำไรสุทธิต่อหนี้สินรวม (Net Income to Total Liabilities) รายละเอียดผลการศึกษาปรากฏดังนี้ (Okay, 2015)

ตารางที่ 3 เปรียบเทียบระดับความถูกต้องแม่นยำโดยใช้ตัวอย่างทั้งหมด (Comparison of Accuracy Levels Using Whole Sample)

Models/ Total Accuracy		Failed	Non- Failed	Type I Error	Type II Error
MDA	Failed	22	10	31.25%	
79.69%	Non-Failed	3	29		37%
QDA	Failed	19	13	49.62%	
76.55%	Non-Failed	2	30		6.25%
Logit	Failed	25	7	21.88%	
76.56%	Non-Failed	8	24		25%
Probit	Failed	26		18.75%	
79.69%	Non-Failed	7	25		21.88%
DT	Failed	28	4	12.5%	
89.06%	Non-Failed		29		9.37%

ที่มา: ดัดแปลงจาก (Okay, 2015)

จากตารางที่ 3 เป็นการเปรียบเทียบระดับความถูกต้องแม่นยำโดยใช้ตัวอย่างทั้งหมด จำนวน 64 บริษัท ในการเปรียบเทียบ 5 โมเดล ปรากฏว่าโมเดลต้นไม้การตัดสินใจ (Decision Tree Model: DT) มีประสิทธิภาพเหนือกว่าโมเดลอื่นๆ โดยมีความถูกต้องแม่นยำในการพยากรณ์รวมร้อยละ 89.06 (สูงที่สุด) และ Type I Error เท่ากับร้อยละ 12.5 (ต่ำสุด) ในขณะที่ MDA และ Probit มีความถูกต้องแม่นยำในการพยากรณ์รวมเท่ากันที่ร้อยละ 79.69 แต่อย่างไรก็ตาม Type I Error ของ MDA เท่ากับร้อยละ 31 ซึ่งสูงกว่า Type I Error ของ Probit เท่ากับร้อยละ 18.8 ในที่นี้ไม่ได้นำ NN และ SVM มาเปรียบเทียบ เนื่องจากเป็นโมเดลที่ต้องการตัวอย่างทดลอง (Training Samples) และตัวอย่างทดสอบ (Test Samples)

ตารางที่ 4 การเปรียบเทียบอัตราความถูกต้องแม่นยำของตัวอย่างทดสอบ (Comparison of Test Samples Accuracy Rates)

Models	% of Correct Classifications	Value of t
1.MDA	77.5	3.1**
2.QDA	79.4	3.3**
3.LOGIT	76.9	3.1**
4.PROBIT	76.9	2.1*
5.DT	68.1	2.1*
6.NN	81.3	3.5***
7.SVM	78.8	3.3**

นัยสำคัญทางสถิติ 0.025 (*) 0.0025 (**) และ 0.001 (***)

ที่มา: (Okay, 2015)

จากตารางที่ 4 เป็นการเปรียบเทียบอัตราความถูกต้องแม่นยำของตัวอย่างทดสอบ (Holdout Samples) พบว่า DT มีความถูกต้องแม่นยำต่ำที่สุดคือร้อยละ 68 ในขณะที่ DT มีความถูกต้องแม่นยำสูงที่สุดเมื่อใช้ตัวอย่างทั้งหมด (All Samples) สาเหตุอาจจะมาจากปัญหาการเข้ากันเกินไป (Overfitting Problem) คือ การที่โมเดลที่ได้จากการใช้ตัวอย่างทดลอง (Training Samples) มีความถูกต้องแม่นยำสูง แต่เมื่อนำไปใช้กับข้อมูลทดสอบ (Test Sample) กลับมีความถูกต้องต่ำ

โมเดลทางสถิติดั้งเดิม (The Traditional Statistical Models) ได้แก่ MDA, QDA, Logit และ Probit ได้ผลลัพธ์ไม่แตกต่างกัน ในขณะที่ NN มีความถูกต้องแม่นยำสูงที่สุดในบรรดาโมเดลทั้งหมด โมเดลที่มีความถูกต้องแม่นยำในการพยากรณ์ไม่แตกต่างกัน ในการใช้ตัวอย่างทั้งหมดหรือตัวอย่างทดสอบ ถือได้ว่าเป็นโมเดลที่ดี

สาเหตุที่ NN มีความถูกต้องแม่นยำในการพยากรณ์สูงที่สุดจากงานวิจัยของ Okay ในปี 2015 เนื่องจากแบบจำลองดังกล่าวเป็นแขนงย่อยของปัญญาประดิษฐ์ (Artificial Intelligence: AI) ซึ่งเป็นแบบจำลองที่มีรูปแบบและโครงสร้างในการทำงานของการ

ประมวลผลเหมือนกับสมองของมนุษย์ โดยสามารถเรียนรู้และปรับเปลี่ยนวิธีการประมวลผลให้ตอบสนองและสอดคล้องต่อข้อมูลนำเข้าได้ดี มีความยืดหยุ่นสูง และสามารถนำมาใช้กับตัวแปรที่ไม่เป็นเส้นตรง (จิรนนท์ เชมจันทร์ & สุรัชย์ จันทร์จรัส, 2556)

ตารางที่ 5 ความสามารถในการพยากรณ์จำแนกตามปีและวิธี (Predictive Ability by Decade and Method)

ช่วงปี	Lowest Accuracy	Highest Accuracy	Method(s) Used to Obtain Highest Accuracy
1960's	79%	92%	Univariate DA [Beatty, 1966]
1970's	56%	100%	Linear Probability; [Meyer and Pifer, 1970] MDA [Cmister, 1972]; [Santomero and Pardo, 1977])
1980's	20%	100%	MDA ([Marais, 1980]; [Betts and Behoul, 1982]; [El Hennawy and Morris, 1983]; [Izan, 1984]; [Takahashi et al.,1984]; [Frydman et al., 1985]) Recursive Partitioning Algorithm [Frydman et al., 1985] Neural Network [Messier and Hansen, 1988]

ตารางที่ 5 ความสามารถในการพยากรณ์จำแนกตามปีและวิธี (Predictive Ability by Decade and Method) (ต่อ)

ช่วงปี	Lowest Accuracy	Highest Accuracy	Method(s) Used to Obtain Highest Accuracy
1990's	27%	100%	Neural Network ([Guan, 1993]; [Tsukuda and Baba, 1994]; [El-Tahawy, 1995]) Judgmental [Koufteros and Puri, 1992] Cumulative Sums [Theodossiou, 1993]
2000's	27%	100%	ML [Atterson, 2001]

ที่มา: (Bellocary, Giacomino, & Akers, 2007)

จากตารางที่ 5 เป็นการเปรียบเทียบความสามารถในการพยากรณ์จำแนกตามปีและวิธี จากงานวิจัยของ Bellocary, Giacomino, และ Akers ในปี 2007 จะเห็นได้ว่าโมเดลมีความถูกต้องแม่นยำในการพยากรณ์สูงสุด (Maximum Accuracy) ถึงร้อยละ 100 อย่างไรก็ตามความถูกต้องแม่นยำในการพยากรณ์ต่ำสุด (Minimum Accuracy) สูงถึงร้อยละ 79 ในช่วงปี 1990's แต่กลับมีค่าลดลงอย่างมากในช่วงปี 1980 ซึ่งเหลือเพียงร้อยละ 20 กล่าวคือโมเดลที่ใหม่กว่า (Newer Models) ไม่ได้นำมาซึ่งความถูกต้องแม่นยำกว่าโมเดลดั้งเดิม (Older Models)

ตารางที่ 6 ความสามารถในการพยากรณ์จำแนกตามโมเดล (Predictive Ability by Model)

Models	Lowest Accuracy	Highest Accuracy	Studies which Obtained Highest Accuracy
MDA	32%	100%	Edmister (1972); Santomero and Vinso (1977); Marais (1980); Potts and Belhoul (1982); El Hennawy and Morris (1983); Izan (1984); Takahashi et al. (1984); Frydman et al. (1985); Patterson (2001)
Logit Analysis	20%	98%	Damjanolena and Shulman (1988)
Probit Analysis	20%	84%	Silva et al. (1990)
Neural Networks	71%	100%	Messier and Hansen (1988); Guan (1993); Tsukuda and Baba (1994); El-Temtamy (1995)

ที่มา: (Bellovary et al., 2007)

จากตารางที่ 6 เป็นการเปรียบเทียบความสามารถในการพยากรณ์จำแนกตามโมเดล (Predictive Ability by Model) จากงานวิจัยจะเห็นว่า MDA และ Neural Network Models มีความถูกต้องแม่นยำในการพยากรณ์สูงที่สุดคือร้อยละ 100 รองลงมาคือ Logit Analysis มีความถูกต้องแม่นยำร้อยละ 98

อย่างไรก็ตามโมเดลที่มีช่วงของความถูกต้องแม่นยำดีที่สุด (The Best Accuracy Range) คือ Neural Networks ซึ่งมีค่าระหว่างร้อยละ 71-100 ในขณะที่ MDA มีค่าระหว่างร้อยละ 32-100 Logit มีค่าระหว่างร้อยละ 20-98 และ Probit มีค่าระหว่างร้อยละ 20-84 ตามลำดับ

ตารางที่ 7 เปรียบเทียบข้อดีและข้อจำกัดของแต่ละแบบจำลอง

แบบจำลอง	ข้อดี	ข้อจำกัด
MDA	1.เหมาะสำหรับการใช้ตัวแปรอิสระหลายตัวในการทำนายตัวแปรตาม 2.สามารถบอกได้ว่าตัวแปรใดจำแนกได้ดีมากน้อยกว่ากัน	1.มีข้อสมมติในเรื่องการแจกแจงปกติของตัวแปรและความแปรปรวนร่วมของเมตริกซ์ความแปรปรวนร่วมของตัวแปรอิสระในแต่ละกลุ่ม 2.เหมาะสำหรับการที่เป็นเส้นตรง
QDA	ใช้กับเมตริกซ์ความแปรปรวนร่วมของตัวแปรอิสระแต่ละกลุ่มไม่เท่ากัน	ใช้ได้กับเชิงแปรอิสระที่เป็นตัวแปรเชิงปริมาณชนิดต่อเนื่องเท่านั้น
Logit Analysis	1.ไม่ต้องมีข้อสมมติในเรื่องการแจกแจงปกติของตัวแปรและความเท่ากันของเมตริกซ์ความแปรปรวนร่วมของตัวแปรอิสระของแต่ละกลุ่ม 2.ยืดหยุ่นและง่ายต่อการใช้งาน	1.เหมาะสำหรับสมการที่เป็นเส้นตรง 2.อธิบายตัวแปรในรูปแบบของความน่าจะเป็น
Probit Analysis	1.ใช้การพยากรณ์ที่ดี ถ้าตัวแปรมีความสัมพันธ์เชิงเส้น 2.สะดวกและง่ายต่อการทำความเข้าใจ	1.เหมาะสำหรับสมการที่เป็นเส้นตรง 2.อธิบายตัวแปรในรูปแบบของความน่าจะเป็น
DT	ง่ายต่อการทำความเข้าใจและการแปลความหมาย	อาจเกิดปัญหา overfitting หรือ under fitting โดยเฉพาะชุดข้อมูลที่มีขนาดเล็ก

ตารางที่ 7 เปรียบเทียบข้อดีและข้อจำกัดของแต่ละแบบจำลอง (ต่อ)

แบบจำลอง	ข้อดี	ข้อจำกัด
Neural Networks	1.มีความยืดหยุ่นสูง 2.สามารถใช้ได้กับตัวแปรที่ไม่เป็นเส้นตรง 3.สามารถเรียนรู้ได้ดีและประยุกต์ใช้ได้หลากหลาย	1.อธิบายความสัมพันธ์ของตัวแปรได้ยาก 2.หลักการทำงานมีความซับซ้อน 3.ไม่มีหลักการตายตัวในการกำหนดโครงสร้างของโครงข่าย
SVM	1.ไม่ค่อยเกิดปัญหา overfitting มากเหมือนกับ Neural Network 2.มี kernel function ที่สามารถเปลี่ยนข้อมูลที่มีมิติ (dimension) ที่ต่ำกว่าให้มีมิติที่สูงขึ้นเพื่อใช้กับแบ่งข้อมูลแบบ linear model ได้	การทำงานของโครงข่ายของ SVM ค่อนข้างเข้าใจยากพอสมควร

ที่มา: (จิรนนท์ เชมพันธ์ & สุรัชย์ จักรินทร์, 2556; เอกสิทธิ์ พัทธวงศ์ศักดิ์, 2557)

สรุป

แบบจำลองการพยากรณ์ภาวะความล้มเหลวมีความสำคัญและจำเป็นต่อบริษัทและผู้มีส่วนได้เสียที่เกี่ยวข้อง โดยแบบจำลองการพยากรณ์ภาวะความล้มเหลวมีหลายชนิด ได้แก่ แบบจำลองทางสถิติดั้งเดิม คือ MDA, QDA, Logit และ Probit และแบบจำลองการเรียนรู้ของเครื่องจักร คือ DT, ANN และ SVM ซึ่งแบบจำลองแต่ละชนิดจะมีข้อได้เปรียบและข้อจำกัดที่แตกต่างกันไป ดังนั้นการเลือกใช้แบบจำลองแต่ละชนิดจึงจำเป็นต้องมีการศึกษาข้อมูลอย่างละเอียดถี่ถ้วน รวมถึงข้อดีหรือข้อได้เปรียบและข้อจำกัดของแต่ละแบบจำลอง เพื่อให้ได้แบบจำลองที่เหมาะสมและสอดคล้องกับวัตถุประสงค์ในการพยากรณ์และลักษณะของข้อมูลที่ใช้ในการพยากรณ์ ซึ่งหากเลือกใช้แบบจำลองที่เหมาะสมก็จะส่งผลต่อความถูกต้องแม่นยำในการพยากรณ์มากยิ่งขึ้น โดยสามารถส่งสัญญาณเตือนภัยล่วงหน้า

ทางการเงินได้ทันเวลา และมีประโยชน์อย่างมากในทางปฏิบัติต่อผู้มีส่วนได้เสีย ผู้สอบบัญชี ผู้จัดการ เจ้าหนี้และนักวิเคราะห์ (จิรนนท์ เชมขันธ & สุรัชย์ จันทรจรัส, 2556)

รายการอ้างอิง

- จิรนนท์ เชมขันธ, & สุรัชย์ จันทรจรัส. (2556). เครื่องมือพยากรณ์ความล้มเหลวทางการเงิน. **วารสารนักบริหาร**, 33(4), 34-41.
- ประเสริฐ ลิฬหาวาสน์, & มนวิภา ผดุงสิทธิ์. (2552). การพยากรณ์ภาวะล้มเหลวทางธุรกิจจากข้อมูลทางบัญชี. **วารสารวิชาชีพบัญชี**, 5(13), 65-69.
- ปานรดา พิลาศรี. (2553). แบบจำลองภาวะความล้มเหลวทางการเงินของบริษัทจดทะเบียนในตลาดหลักทรัพย์แห่งประเทศไทย. (วิทยานิพนธ์ปริญญาโท), มหาวิทยาลัยธรรมศาสตร์กรุงเทพฯ.
- สิทธิพล สมชม, & ณัฐวุฒิ คุ้มฒันเจริญชัย. (2557). ความสัมพันธ์ของอันดับความน่าเชื่อถือกับโอกาสประสบภาวะตกต่ำทางการเงินของบริษัทที่จดทะเบียนในตลาดหลักทรัพย์แห่งประเทศไทย. Paper presented at the SEC Working Papers Forum ครั้งที่ 3/2557 กรุงเทพฯ.
- อภิญญา ออดทน. (2553). การเปรียบเทียบความสามารถของแบบจำลองเพื่อทำนายภาวะล้มเหลวทางการเงิน. (วิทยานิพนธ์ปริญญาโท), มหาวิทยาลัยธรรมศาสตร์.
- เอกสิทธิ์ พิชรวงศ์. (2557). เทคนิค Support Vector Machine (SVM) และ Soft Margin HR-SVM. Retrieved 18 มีนาคม 2560, from <http://datamined.com/2014/support-vector-machine-svm/>
- Altman, E. I. (1968). Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. **The Journal of Finance**, 23(4), 598-609.
- Aziz, M. A., & Dar, H. A. (2006). Predicting corporate bankruptcy: where we stand? **CORPORATE GOVERNANCE**, 6(1), 18-33.

- Bal, J., Cheung, Y., & Wu, H.-C. (2013). Entropy for business failure prediction: an improved prediction model for the construction industry. **Advances in Decision Sciences**, 2013, 1-14.
- Beaver, W. H. (1966). Financial ratios as predictors of failure. **Journal of Accounting Research**, 4, 71-111.
- Beaver, W. H. (1968). Market prices, financial ratios, and the prediction of failure. **Journal of Accounting Research**, 6(2), 179-192.
- Bellovary, J., Giacominio, D., & Akers, M. (2007). A review of bankruptcy prediction studies: 1930-present. **Journal of Financial Education**, 33, 1-42.
- Kpodoh, B. (2009). **Bankruptcy and financial distress prediction in the mobile telecom industry**. (Master's Degree in Business Administration), Blekinge Institute of Technology.
- National Institute of Open Schooling. (2016). **Objectives of business**. Retrieved 12 พฤศจิกายน 2559, from <http://old.nios.ac.in/Secbus/cour/cc03.pdf>
- Ohlson, J. A. (1980). Financial ratios and the probabilistic prediction of bankruptcy. **Journal of Accounting Research**, 18(1), 109-131.
- Okay, K. (2015). **Predicting business failures in non-financial Turkish companies**. (Master of Science), İhsan Doğramacı Bilkent University.
- RapidMiner. (2016). **Quadratic Discriminant Analysis**. Retrieved 22 ธันวาคม 2559, from http://docs.rapidminer.com/studio/operators/modeling/predictive/discriminant_analysis/quadratic_discriminant_analysis.