

บทความวิจัย

การเปรียบเทียบรูปแบบการจัดกลุ่มและตัวแบบพยากรณ์ที่เหมาะสมที่สุดสำหรับการพยากรณ์ความต้องการอะไหล่สำรองในงานซ่อมบำรุง

ธวัชลักษณ์ จินะสี* และ ชูศักดิ์ พรสิงห์

ภาควิชาวิศวกรรมอุตสาหการและการจัดการงานวิศวกรรม คณะวิศวกรรมศาสตร์และเทคโนโลยีอุตสาหกรรม มหาวิทยาลัยศิลปากร

* ผู้นิพนธ์ประสานงาน โทรศัพท์ 09 1857 0299 อีเมล: jinasee_t@silpakorn.edu DOI: 10.14416/j.kmutnb.2026.09.006

รับเมื่อ 26 ธันวาคม 2568 แก้ไขเมื่อ 27 มีนาคม 2569 ตอบรับเมื่อ 15 พฤษภาคม 2569 เผยแพร่ออนไลน์ 14 กันยายน 2569

© 2026 King Mongkut's University of Technology North Bangkok. All Rights Reserved.

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อค้นหารูปแบบการจัดกลุ่มอะไหล่สำรองและวิธีการพยากรณ์ที่เหมาะสมสำหรับอะไหล่สำรองในงานซ่อมบำรุงที่มีลักษณะอุปสงค์แตกต่างกัน วิธีการศึกษาประกอบด้วยการจัดกลุ่มอะไหล่สำรองตามค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวน โดยเปรียบเทียบเทคนิคการจัดกลุ่ม 4 วิธี ได้แก่ การจัดกลุ่มตามรูปแบบอุปสงค์ การจัดกลุ่มแบบเค-มีนส์ การจัดกลุ่มแบบลำดับชั้น และการจัดกลุ่มเชิงพื้นที่ โดยอาศัยความหนาแน่นและจุดรวมก้อนข้อมูลที่ใช้ในการศึกษาเป็นข้อมูลความต้องการอะไหล่สำรองรายเดือนระหว่าง พ.ศ. 2563–2566 รวม 48 เดือน ครอบคลุมอะไหล่สำรอง 1,735 รายการ โดยใช้ข้อมูล พ.ศ. 2563–2565 (36 เดือน) เป็นชุดเรียนรู้เพื่อสร้างและปรับเทียบตัวแบบ และใช้ข้อมูล พ.ศ. 2566 (12 เดือน) เป็นชุดทดสอบเพื่อประเมินความแม่นยำของการพยากรณ์ วิธีการพยากรณ์ที่ใช้ในแต่ละกลุ่มอะไหล่สำรองประกอบด้วย ค่าเฉลี่ยเคลื่อนที่อย่างง่าย วิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์ วิธีของโครสตัน และตัวแบบถดถอยอัตโนมัติแบบอินทิเกรต และค่าเฉลี่ยเคลื่อนที่ ความแม่นยำของการพยากรณ์ประเมินด้วยค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร ผลการวิจัยพบว่าการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมก้อน เมื่อจับคู่กับการพยากรณ์ด้วยตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ให้ค่าความแม่นยำของการพยากรณ์ประเมินด้วยค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรต่ำที่สุดในทุกกลุ่ม โดยกลุ่มที่ 1 กลุ่มที่ 2 กลุ่มที่ 3 และกลุ่มที่ 4 มีค่าเท่ากับ 3.01% 27.03% 5.37% และ 4.17% ตามลำดับ แสดงให้เห็นถึงประสิทธิภาพของการผสมผสานเทคนิคการจัดกลุ่มและการพยากรณ์ดังกล่าวสำหรับการพยากรณ์ความต้องการอะไหล่สำรองในงานซ่อมบำรุง

คำสำคัญ: การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมก้อน การพยากรณ์ความต้องการอะไหล่สำรอง ค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร (sMAPE) งานซ่อมบำรุง ตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่

การอ้างอิงบทความ: ธวัชลักษณ์ จินะสี และ ชูศักดิ์ พรสิงห์, “การเปรียบเทียบรูปแบบการจัดกลุ่มและตัวแบบพยากรณ์ที่เหมาะสมที่สุดสำหรับการพยากรณ์ความต้องการอะไหล่สำรองในงานซ่อมบำรุง,” วารสารวิชาการพระจอมเกล้าพระนครเหนือ, ปีที่ 36, ฉบับที่ 4, หน้า 1–27, ต.ค.-ธ.ค. 2569. เลขที่บทความ 264-038347, doi: 10.14416/j.kmutnb.2026.09.006



Research Article

A Comparative Analysis of Clustering Schemes and Optimal Forecasting Models for Spare Parts Demand Forecast in Maintenance Operation

Thawanluck Jinasee* and Choosak Pornsing

Department of Industrial Engineering and Management, Faculty of Engineering and Industrial Technology, Silpakorn University, Nakhon Pathom, Thailand

* Corresponding Author, Tel.09 1857 0299, E-mail: jinasee_t@silpakorn.edu DOI: 10.14416/j.kmutnb.2026.09.006

Received 26 December 2025; Revised 27 March 2026; Accepted 15 May 2026; Published online: 14 September 2026

© 2026 King Mongkut's University of Technology North Bangkok. All Rights Reserved.

Abstract

This study aims to identify the most appropriate clustering scheme and forecasting method for spare parts used in maintenance operations with heterogeneous demand characteristics. The proposed methodology first classifies spare parts based on the Average Demand Interval (ADI) and the Squared Coefficient of Variation (CV^2). Four clustering techniques are compared: demand segmentation based on demand patterns, K-means clustering, hierarchical clustering, and Density-Based Spatial Clustering of Applications with Noise (DBSCAN). The empirical analysis is conducted using monthly demand data from 2020 to 2023 (48 months) covering 1,735 spare parts. Demand data from 2020 to 2022 (36 months) are used as the training set for model development and parameter calibration, while data from 2023 (12 months) are used as the test set to evaluate out-of-sample forecasting performance. For each cluster, demand is forecast using the Simple Moving Average (SMA), Holt's exponential smoothing method, Croston's method, and the Autoregressive Integrated Moving Average (ARIMA) model. Forecast accuracy is evaluated using the Symmetric Mean Absolute Percentage Error (sMAPE). The results indicate that DBSCAN combined with the ARIMA model provides the lowest sMAPE values across all clusters. Specifically, cluster 1, cluster 2, cluster 3, and cluster 4 achieve sMAPE values of 3.01%, 27.03%, 5.37%, and 4.17%, respectively. These findings demonstrate the effectiveness of integrating clustering and forecasting techniques for spare-parts demand forecasting in maintenance operations.

Keywords: Density-Based Spatial Clustering of Applications with Noise (DBSCAN), Spare-Parts Demand Forecasting, Symmetric Mean Absolute Percentage Error (sMAPE), Maintenance Operations, Autoregressive Integrated Moving Average (ARIMA)

Please cite this article as: T. Jinasee and C. Pornsing, "A comparative analysis of clustering schemes and optimal forecasting models for spare parts demand forecast in maintenance operation," *The Journal of KMUTNB*, vol. 36, no. 4, pp. 1-27, Oct.-Dec. 2026 (in Thai), Art. no. 264-038347, doi: 10.14416/j.kmutnb.2026.09.006

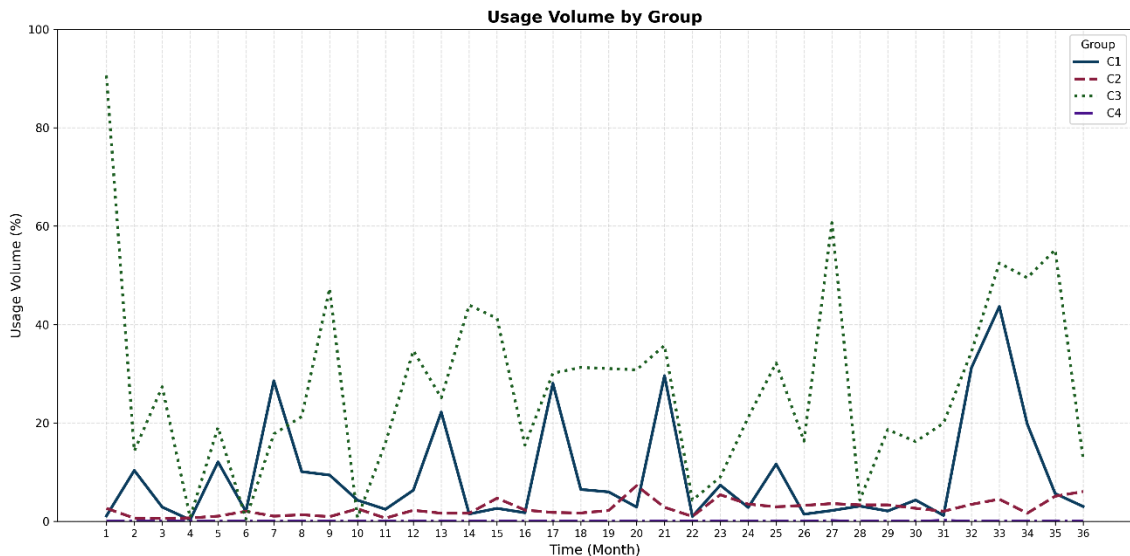
1. บทนำ

อุตสาหกรรมเครื่องตีของประเทศไทยมีบทบาทสำคัญต่อระบบเศรษฐกิจและมีพัฒนาการอย่างต่อเนื่องจากการผลิตเพื่อทดแทนการนำเข้าได้ขยายสู่การผลิตเชิงพาณิชย์ที่มีประสิทธิภาพและหลากหลายเพื่อตอบสนองความต้องการของผู้บริโภคทั้งในประเทศและต่างประเทศ การเติบโตของอุตสาหกรรมนี้ได้รับแรงสนับสนุนจากการฟื้นตัวของภาคการท่องเที่ยวและธุรกิจบริการที่เกี่ยวข้อง ส่งผลให้มูลค่าตลาดเครื่องตีของไทยใน พ.ศ. 2566 อยู่ในระดับสูงและมีแนวโน้มเติบโตต่อเนื่อง [1] ภายใต้บริบทดังกล่าว ความน่าเชื่อถือ (Reliability) และประสิทธิภาพของเครื่องจักรและสายการผลิตจึงเป็นปัจจัยสำคัญต่อความสามารถในการแข่งขันของผู้ประกอบการ โรงงานตัวอย่างในงานวิจัยนี้เป็นผู้ผลิตเครื่องตีรายใหญ่ของประเทศดำเนินการผลิตทั้งเพื่อการส่งออกและจำหน่ายภายในประเทศภายใต้ระบบการผลิตแบบต่อเนื่อง ทำให้การหยุดชะงักของสายการผลิตแม้เพียงระยะเวลาสั้นอาจก่อให้เกิดความสูญเสียทางเศรษฐกิจอย่างมีนัยสำคัญ ดังนั้น งานซ่อมบำรุงเชิงรุกและเชิงป้องกันจึงมีบทบาทสำคัญต่อการควบคุมต้นทุน การรักษาประสิทธิภาพของกระบวนการผลิต และการคงคุณภาพของผลิตภัณฑ์ให้เป็นไปตามมาตรฐาน

ประเด็นสำคัญของงานซ่อมบำรุงคือการบริหารจัดการอะไหล่สำรอง (Spare Parts Management) เนื่องจากอะไหล่แต่ละรายการมีลักษณะการใช้งานแตกต่างกัน ทั้งในด้านความสม่ำเสมอ ความถี่ และความผันผวนของอุปสงค์ ส่งผลให้การพยากรณ์ความต้องการมีความซับซ้อนกว่าสินค้าทั่วไป โดยความท้าทายสำคัญคือ การสร้างสมดุลระหว่างการมีสินค้าคงคลังมากเกินไป

(Overstocking) ซึ่งเพิ่มต้นทุนการถือครองกับการมีสินค้าคงคลังไม่เพียงพอ (Stock-out) ซึ่งกระทบต่อความพร้อมของเครื่องจักรและความต่อเนื่องของการผลิตโดยตรง

ในทางปฏิบัติโรงงานตัวอย่างยังใช้นโยบายกำหนดระดับการจัดเก็บอะไหล่สำรองตามระดับความสำคัญของอะไหล่ ได้แก่ ความสำคัญสูง (C1) ความสำคัญปานกลาง (C2) ความสำคัญน้อย (C3) และไม่มีความสำคัญ (C4) ซึ่งสะท้อนผลกระทบต่อการผลิต ต้นทุน และโอกาสในการยอมรับการขาดอะไหล่ที่แตกต่างกัน อย่างไรก็ตามการจำแนกตามระดับความสำคัญดังกล่าวยังไม่สามารถอธิบายความแตกต่างของรูปแบบอุปสงค์ของอะไหล่แต่ละรายการได้โดยตรง เนื่องจากอะไหล่ที่อยู่ในระดับความสำคัญเดียวกันอาจมีพฤติกรรมการใช้งานแตกต่างกันอย่างชัดเจน เมื่อพิจารณาปริมาณการใช้งานอะไหล่สำรองจำแนกตามกลุ่มความสำคัญ C1–C4 ดังแสดงในรูปที่ 1 ซึ่งแกน X แสดงลำดับเดือนที่ใช้ในการศึกษา 36 เดือน และแกน Y แสดงปริมาณการใช้งานอะไหล่สำรอง (Usage Volume) คิดเป็นร้อยละพบว่าอะไหล่สำรองแต่ละกลุ่มมีพฤติกรรมการใช้งานแตกต่างกันทั้งในด้านความถี่ ระดับการใช้ และความผันผวน โดยกลุ่ม C1 มีลักษณะการใช้งานไม่ต่อเนื่องและผันผวนสูง กลุ่ม C2 มีระดับการใช้งานต่ำและค่อนข้างสม่ำเสมอ ขณะที่กลุ่ม C3 มีปริมาณการใช้งานสูงและผันผวนเด่นชัดในหลายช่วงเวลา ส่วนกลุ่ม C4 มีการใช้งานต่ำมากเกือบตลอดช่วงการศึกษา ผลดังกล่าวสะท้อนว่า การจัดกลุ่มอะไหล่สำรองตามระดับความสำคัญเพียงอย่างเดียวอาจยังไม่เพียงพอต่อการพยากรณ์และการวางแผนอะไหล่สำรองอย่างมีประสิทธิภาพ



รูปที่ 1 ปริมาณการใช้งานอะไหล่สำรองในแต่ละกลุ่ม

งานวิจัยที่ผ่านมาได้เสนอแนวทางที่หลากหลาย เพื่อเพิ่มความแม่นยำของการพยากรณ์และสนับสนุนการจัดการสินค้าคงคลัง โดยมีทั้งการประยุกต์ใช้การจัดกลุ่มร่วมกับการเรียนรู้แบบผสมผสานเพื่อเพิ่มความแม่นยำของการพยากรณ์ในบริบทค่าปลีก ซึ่งสะท้อนให้เห็นว่าการจัดกลุ่มข้อมูลก่อนการพยากรณ์สามารถช่วยเพิ่มประสิทธิภาพของแบบจำลองได้ อย่างไรก็ตามบริบทของข้อมูลดังกล่าวยังแตกต่างจากข้อมูลอะไหล่สำรองอย่างมีนัยสำคัญ [2] นอกจากนี้ยังมีการใช้ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ (Average Demand Interval; ADI) และค่ากำลังสองสัมประสิทธิ์ความแปรปรวน (Squared Coefficient of Variation; CV^2) เพื่อจำแนกรูปแบบอุปสงค์ และเชื่อมโยงผลการจำแนกไปสู่การกำหนดนโยบายสินค้าคงคลังที่เหมาะสม ซึ่งแสดงให้เห็นถึงความสำคัญของการวิเคราะห์ลักษณะอุปสงค์ต่อการตัดสินใจด้านคลังสินค้า แต่ยังไม่ได้เปรียบเทียบหลายวิธีการจัดกลุ่มร่วมกับหลายวิธีการพยากรณ์ภายใต้กรอบการทดลองเดียวกัน [3] ขณะเดียวกันงานวิจัยที่เปรียบเทียบ

วิธีพยากรณ์แบบดั้งเดิมกับวิธีการเรียนรู้ของเครื่องสำหรับยอดขายแบตเตอรี่พบว่าแบบจำลองผสมให้ผลดีที่สุด แต่ยังเป็นการศึกษาบนข้อมูลยอดขายสินค้าเชิงพาณิชย์และไม่ได้พิจารณาการจัดกลุ่มข้อมูลตามโครงสร้างอุปสงค์ก่อนเข้าสู่กระบวนการพยากรณ์ [4]

สำหรับงานวิจัยที่เกี่ยวข้องกับอะไหล่สำรองโดยตรง มีการเปรียบเทียบตัวแบบพยากรณ์หลายประเภทสำหรับอุปสงค์ไม่ต่อเนื่องในโซ่อุปทานอะไหล่สำรอง เช่น ค่าเฉลี่ยเคลื่อนที่ การทำให้เรียบแบบเอกซ์โพเนนเชียล วิธีของโครอสตัน ตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ (Autoregressive Integrated Moving Average; ARIMA) และโครงข่ายประสาทเทียม โดยชี้ให้เห็นว่าการเลือกตัวแบบพยากรณ์ควรพิจารณาให้สอดคล้องกับลักษณะของอุปสงค์แต่ละประเภท [5] ขณะเดียวกัน งานวิจัยด้านการพยากรณ์อะไหล่สำรองที่มีอุปสงค์เป็นก้อน ผันผวน และไม่ต่อเนื่องได้เสนอวิธีพยากรณ์แบบผสมที่ใช้วิธีสถิติสักร่วมกับบูตสตรอปปีง (Bootstrapping) เพื่อรองรับความ

ไม่แน่นอนของอุปสงค์และสนับสนุนการจัดการสินค้าคงคลัง [6] นอกจากนี้ยังมีการศึกษาการพยากรณ์อะไหล่สำรองในบริบทโดยสาร โดยเปรียบเทียบวิธีถดถอยวิธีฐานกฎ วิธีต้นไม้ และโครงข่ายประสาทเทียม พร้อมพิจารณาตัวแปรเชิงปฏิบัติการ เช่น จำนวนยานพาหนะ จำนวนการซ่อมบำรุง และระยะเวลาเฉลี่ยระหว่างความขัดข้อง [7]

จากงานวิจัยข้างต้นจะเห็นได้ว่าแม้งานวิจัยที่ผ่านมาได้ให้ความสำคัญกับการพัฒนาหรือเปรียบเทียบตัวแบบพยากรณ์ การจำแนกรูปแบบอุปสงค์ และการสนับสนุนการจัดการสินค้าคงคลัง แต่ยังขาดการศึกษาเชิงเปรียบเทียบที่พิจารณาความสัมพันธ์ระหว่างวิธีการจัดกลุ่มและตัวแบบพยากรณ์อย่างเป็นระบบบนข้อมูลอะไหล่สำรองจริง โดยเฉพาะการประเมินผลหลายวิธีการจัดกลุ่มร่วมกับหลายตัวแบบพยากรณ์ภายใต้ชุดข้อมูลและเกณฑ์การประเมินเดียวกันงานวิจัยนี้จึงมุ่งเติมเต็มช่องว่างดังกล่าว

งานวิจัยนี้ประยุกต์การจัดกลุ่มอะไหล่สำรองตามลักษณะอุปสงค์ โดยพิจารณาจากค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวน อันเป็นพื้นฐานของการแบ่งส่วนอุปสงค์ (Demand Segmentation) พร้อมเปรียบเทียบวิธีการจัดกลุ่ม 4 วิธี และตัวแบบพยากรณ์ 4 วิธีโดยใช้ค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร (Symmetric Mean Absolute Percentage Error; sMAPE) เป็นตัววัดความแม่นยำ ทั้งนี้ การวิเคราะห์อาศัยข้อมูลย้อนหลัง 36 เดือน ระหว่าง พ.ศ. 2563–2565 ครอบคลุมอะไหล่จำนวน 1,735 รายการ

ดังนั้นงานวิจัยนี้มีวัตถุประสงค์เพื่อเปรียบเทียบและคัดเลือกชุดค่าผสมที่เหมาะสมที่สุดระหว่างวิธีการจัดกลุ่มอะไหล่สำรองและตัวแบบพยากรณ์ เพื่อเพิ่ม

ความแม่นยำในการพยากรณ์ความต้องการอะไหล่สำรองและสนับสนุนการบริหารสินค้าคงคลังสำหรับงานซ่อมบำรุงอย่างมีประสิทธิภาพ

2. วิธีการทดลอง

การดำเนินการวิจัยครอบคลุมตั้งแต่การรวบรวมข้อมูล การจัดกลุ่มข้อมูล การพยากรณ์ และการประเมินผล โดยออกแบบกระบวนการทดลองให้สามารถเปรียบเทียบประสิทธิภาพของวิธีการจัดกลุ่มและตัวแบบพยากรณ์ได้อย่างเป็นระบบ ภายใต้ชุดข้อมูลและเกณฑ์การประเมินเดียวกัน เพื่อคัดเลือกแนวทางที่เหมาะสมที่สุดสำหรับการพยากรณ์ความต้องการอะไหล่สำรอง

2.1 การรวบรวมข้อมูล

ในงานวิจัยนี้ชุดข้อมูลที่น่ามาใช้เป็นข้อมูลที่รวบรวมย้อนหลัง 48 เดือน ตั้งแต่ เดือนมกราคม พ.ศ. 2563 ถึง เดือนธันวาคม พ.ศ. 2566 ซึ่งเป็นข้อมูลที่สามารถดึงออกมาจากระบบวางแผนทรัพยากรองค์กร (Enterprise Resource Planning; ERP) ได้แก่ ปริมาณการใช้อะไหล่สำรองย้อนหลัง ประเภทงานซ่อมเมื่อมีการเบิกอะไหล่ ประเภทความสำคัญของอะไหล่แต่ละรายการโดยใน พ.ศ. 2565 โดยมีสัดส่วนมูลค่าคงคลังเป็นดังตารางที่ 1 และ อะไหล่สำรองแต่ละประเภทความสำคัญมีจำนวนดังตารางที่ 2

ตารางที่ 1 มูลค่าอะไหล่สำรองตามประเภทความสำคัญ

ประเภทของอะไหล่	มูลค่า
อะไหล่สำรองประเภทที่มีความสำคัญสูง (C1)	67.58%
อะไหล่สำรองประเภทที่มีความสำคัญปานกลาง (C2)	26.83%
อะไหล่สำรองประเภทที่มีความสำคัญน้อย (C3)	5.22%
อะไหล่สำรองประเภทที่ไม่มีความสำคัญ (C4)	0.37%
รวม	100%

ตารางที่ 2 ะไหล่สำรองตามเกณฑ์ที่โรงงานกำหนด

ประเภทของอะไหล่	รายการ
อะไหล่สำรองประเภทที่มีความสำคัญสูง (C1)	336
อะไหล่สำรองประเภทที่มีความสำคัญปานกลาง (C2)	317
อะไหล่สำรองประเภทที่มีความสำคัญน้อย (C3)	1,074
อะไหล่สำรองประเภทที่ไม่มีความสำคัญ (C4)	8
รวม	1,735

ในการนำข้อมูลมาจัดกลุ่มอะไหล่สำรอง และพยากรณ์นั้น ต้องทำการแบ่งชุดข้อมูลออกเป็น 2 ชุด ได้แก่

2.1.1 ข้อมูลชุดเรียนรู้ (Training Dataset)

ข้อมูลชุดเรียนรู้ประกอบด้วยข้อมูลการใช้อะไหล่สำรองย้อนหลังระหว่าง พ.ศ. 2563 ถึง พ.ศ. 2565 รวม 36 เดือน ใช้สำหรับคำนวณตัวแปรเชิงลักษณะของอุปสงค์ ได้แก่ ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน และค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ ซึ่งคำนวณตามสมการ (1) และ (2) [3] ตามลำดับ เพื่อสะท้อนระดับความผันผวนและความถี่ของการเกิดอุปสงค์ค่าตัวแปรดังกล่าวถูกใช้เป็นคุณลักษณะหลักในการจัดกลุ่มอะไหล่สำรอง

2.1.2 ข้อมูลชุดทดสอบ (Testing Dataset)

ข้อมูลชุดทดสอบเป็นข้อมูลการใช้อะไหล่สำรองในปี พ.ศ. 2566 รวม 12 เดือน ใช้สำหรับประเมินประสิทธิภาพของตัวแบบพยากรณ์ จากนั้นทำการคำนวณค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวน

ก. ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน คือ ค่าที่แสดงให้เห็นถึงความผันแปรของข้อมูลเทียบกับค่าเฉลี่ยของข้อมูลทั้งหมด เพื่อพิจารณาอัตราการใช้งานวัสดุคงคลังในแต่ละช่วงเวลา

$$CV^2 = \left[\frac{\frac{\sum_{i=1}^n (\mu_{r_i} - \mu_a)^2}{N}}{\mu_a} \right]^2 \quad (1)$$

โดยที่ CV^2 คือ ค่ากำลังสองของสัมประสิทธิ์ความแปรปรวน

μ_{r_i} คือ ปริมาณความต้องการใช้วัสดุในแต่ละเดือน

μ_a คือ ปริมาณความต้องการใช้วัสดุคงคลังเฉลี่ย

N คือ จำนวนเดือนที่ใช้ของวัสดุคงคลังแต่ละรายการ

i คือ ดัชนีช่วงเวลา

T_i คือ ระยะเวลาที่ไม่มีความต้องการใช้วัสดุคงคลังในแต่ละเดือน

ข. ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์คือ ช่วงระยะเวลาเฉลี่ยของการเกิดอุปสงค์แต่ละครั้ง

$$ADI = \frac{\sum_{i=1}^n t_i}{N} \quad (2)$$

โดยที่ ADI คือ ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์

i คือ ดัชนีช่วงเวลา

t_i คือ ระยะเวลาที่ไม่มีความต้องการใช้วัสดุคงคลังในแต่ละเดือน

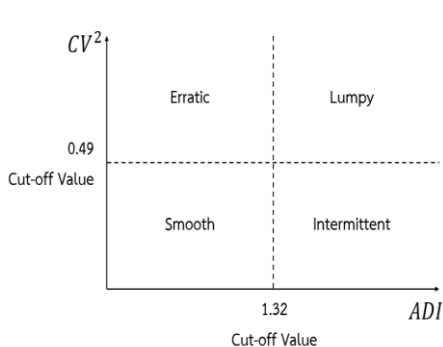
N คือ จำนวนเดือนที่ใช้ของวัสดุคงคลังแต่ละรายการ

2.2 การเตรียมข้อมูลก่อนการประมวลผล

นำค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนของทั้ง 1,735 รายการมาจัดกลุ่มอะไหล่สำรองโดยที่งานวิจัยนี้มีอยู่ทั้งหมด 4 วิธีดังนี้

2.2.1 การจัดกลุ่มตามรูปแบบอุปสงค์ (Demand Pattern)

จากงานวิจัย [8] อธิบายไว้ว่าการจัดกลุ่มอะไหล่สำรองด้วยรูปแบบอุปสงค์ควรใช้ค่าการตัดสินใจของค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์อยู่ที่ 1.32 และ ค่าการตัดสินใจของค่ากำลังสองสัมประสิทธิ์ความแปรปรวนอยู่ที่ 0.49 ดังแสดงในรูปที่ 2



รูปที่ 2 ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และ ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน

การจัดกลุ่มอะไหล่ด้วยรูปแบบอุปสงค์จะสามารถแบ่งกลุ่มได้เป็น 4 ประเภท [3] ดังนี้

1. อุปสงค์รูปแบบราบเรียบ (Smooth Demand) กลุ่มที่มีการเบิกใช้บ่อย ปริมาณคงที่ ($ADI < 1.32$ และ $CV^2 < 0.49$)
2. อุปสงค์รูปแบบไม่สม่ำเสมอ (Intermittent Demand) กลุ่มที่มีการเบิกใช้ไม่บ่อย ปริมาณคงที่ ($ADI \geq 1.32$ และ $CV^2 < 0.49$)
3. อุปสงค์รูปแบบผันผวน (Erratic Demand) กลุ่มที่มีการเบิกใช้บ่อย ปริมาณไม่คงที่ ($ADI < 1.32$ และ $CV^2 \geq 0.49$)
4. อุปสงค์รูปแบบเป็นก้อน (Lumpy Demand) กลุ่มที่มีการเบิกใช้ไม่บ่อย ปริมาณไม่คงที่ ($ADI \geq 1.32$ และ $CV^2 \geq 0.49$)

อย่างไรก็ตามการใช้ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนตามที่งานวิจัยแนะนำทำให้ 99.6% ของอะไหล่สำรองอยู่ในกลุ่มของอุปสงค์รูปแบบเป็นก้อน และอีก 0.4% อยู่ในกลุ่มของอุปสงค์รูปแบบผันผวน เนื่องจากข้อมูลของทางโรงงานตัวอย่างมีค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนค่อนข้างสูง จากงานวิจัย [9] มีผลลัพธ์ในลักษณะของการจับกลุ่มที่คล้ายกันและได้ทำการปรับเปลี่ยนค่าการตัดสินใจของค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวน ทางผู้วิจัยจึงได้ทำการวิเคราะห์คุณลักษณะพื้นฐานของข้อมูลเพิ่มเติมโดยพิจารณา ค่าเฉลี่ย และค่ามัธยฐานดังตารางที่ 3

ตารางที่ 3 คุณลักษณะพื้นฐานของข้อมูลค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวน

ตัวแปร	จำนวน	ค่าเฉลี่ย	ค่ามัธยฐาน
ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์	1,735	59.15	50
ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน	1,735	61.19	52.04

จากการวิเคราะห์คุณลักษณะของข้อมูลพบว่า คุณลักษณะพื้นฐานของข้อมูลค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนมีค่าเฉลี่ยเท่ากับ 59.15 และ 61.19 ตามลำดับ แสดงให้เห็นว่าอะไหล่ส่วนใหญ่มีการใช้งานที่ไม่สม่ำเสมอและมีความผันผวนของปริมาณการใช้สูง ซึ่งสะท้อนถึงลักษณะความต้องการแบบไม่ต่อเนื่อง ทั้งนี้ ค่ามัธยฐานของค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนที่อยู่ที่ 50 และ 52.04 ตามลำดับ และจะนำมาใช้เป็นค่าการตัดสินใจใหม่ในการ

จัดกลุ่มตามรูปแบบอุปสงค์แทนค่าการตัดสินใจเดิม เพื่อจัดกลุ่มอะไหล่ตามลักษณะการใช้งานได้อย่างเหมาะสม

2.2.2 การจัดกลุ่มแบบเค-มีนส์ (K-means)

เป็นการแบ่งชุดข้อมูลออกเป็น k คลัสเตอร์ ซึ่งสันนิษฐานว่ามีความแปรปรวนเท่ากัน โดยแต่ละคลัสเตอร์แทนด้วยจุดศูนย์กลาง (Centroid) ซึ่งคำนวณจากค่าเฉลี่ยของข้อมูลภายในคลัสเตอร์ กระบวนการจัดกลุ่มมีวัตถุประสงค์เพื่อลดความแปรปรวนภายในคลัสเตอร์ โดยทำให้ผลรวมของระยะทางระหว่างข้อมูลแต่ละจุดกับเซนทรอยด์ของคลัสเตอร์ที่สังกัดมีค่าน้อยที่สุด [2] ซึ่งนิยามเชิงคณิตศาสตร์ด้วยฟังก์ชันวัตถุประสงค์ดังแสดงในสมการ (3)

$$\sum_{i=0}^n \min_{\mu_j} \|x_i - \mu_j\| \quad (3)$$

โดยที่ x_i คือ ค่าข้อมูลตัวที่ i (Sample หรือ Record ใน dataset)

μ_j คือ ค่าเฉลี่ยของคลัสเตอร์ที่ j หรือ Centroid

$|x_i - \mu_j|$ คือ ระยะทางระหว่างข้อมูล x_i กับ Centroid

μ_j (โดยทั่วไปเป็น Euclidean Distance)

$\min_{\mu_j}$ คือ เลือกระยะทางที่น้อยที่สุดระหว่าง x_i กับ

centroid ทั้งหมด

$\sum_{i=0}^n$ คือ รวมผลของข้อมูลทุกตัวใน dataset

n คือ จำนวนข้อมูลทั้งหมดในชุดข้อมูล

k คือ จำนวนคลัสเตอร์ที่กำหนดไว้ล่วงหน้า

2.2.3 การจัดกลุ่มแบบลำดับขั้น (Hierarchical)

เป็นเทคนิคการทำเหมืองข้อมูล (Data Mining) แบบไม่กำกับ (Unsupervised Learning) ที่มีวัตถุประสงค์เพื่อจำแนกหรือจัดกลุ่มข้อมูลที่มีลักษณะคล้ายกันตามตัวแปรหลายมิติ โดยสร้างโครงสร้างลำดับขั้นของกลุ่ม

ข้อมูลจากระยะทางหรือความคล้ายคลึงระหว่างตัวอย่าง (Distance-based Grouping) ซึ่งช่วยให้มองเห็นความสัมพันธ์เชิงโครงสร้างของข้อมูลได้อย่างเป็นระบบ โดยเริ่มจากข้อมูลแต่ละหน่วยเป็นกลุ่มย่อยแล้วรวมกันทีละขั้นจนเกิดกลุ่มใหญ่ตามระดับความใกล้เคียงของคุณลักษณะ [10]

2.2.4 การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรบกวน (DBSCAN)

เป็นการจัดกลุ่มแบบไม่กำกับ (Unsupervised Clustering) ที่อาศัยแนวคิดของความหนาแน่นของข้อมูล (Data Density) ในการระบุและจัดกลุ่มจุดข้อมูล โดยไม่จำเป็นต้องกำหนดจำนวนคลัสเตอร์ล่วงหน้า อัลกอริทึมนี้จะจำแนกข้อมูลออกเป็นสามประเภท ได้แก่ จุดแกน (Core Points) ที่อยู่ในบริเวณที่มีความหนาแน่นสูง จุดขอบเขต (Border Points) ที่อยู่ใกล้จุดแกน และจุดรบกวนหรือข้อมูลผิดปกติ (Noise Points) ซึ่งอยู่นอกบริเวณความหนาแน่นนั้น [11]

ในการศึกษานี้ ผู้วิจัยใช้กราฟ k -distance plot เพื่อช่วยเลือกค่าพารามิเตอร์ ϵ (Epsilon) ของ DBSCAN และทำการปรับค่า ϵ เป็น 3 กรณี เพื่อให้จำนวนคลัสเตอร์ที่ได้ใกล้เคียง 2 3 และ 4 กลุ่ม ตามลำดับ ทั้งนี้ เพื่อให้สามารถเปรียบเทียบผลลัพธ์กับวิธี K-means และ Hierarchical ซึ่งต้องกำหนดจำนวนคลัสเตอร์ล่วงหน้าได้อย่างเป็นธรรมชาติ

จากนั้นทำการแบ่งกลุ่มด้วยวิธีที่กล่าวมาข้างต้น ทั้ง 4 วิธีโดยที่ K-means และ Hierarchical กำหนดจำนวนคลัสเตอร์เป็น 2 3 และ 4 กลุ่ม ส่วน DBSCAN ปรับค่า ϵ ให้ได้จำนวนคลัสเตอร์ใกล้เคียง 2 3 และ 4 กลุ่ม แล้วจึงนำผลลัพธ์ไปเปรียบเทียบกับวิธีการจัดกลุ่มด้วย Demand Pattern

2.3 การพยากรณ์ความต้องการอะไหล่สำรอง

นำข้อมูลที่ได้จากการจัดกลุ่มแล้วมาทำการพยากรณ์ปริมาณการใช้อะไหล่สำรองใน พ.ศ. 2566 โดยเลือกวิธีการพยากรณ์ 4 วิธีเพื่อให้ครอบคลุมลักษณะอุปสงค์ที่แตกต่างกันของข้อมูล กล่าวคือ วิธีค่าเฉลี่ยเคลื่อนที่อย่างง่าย (Simple Moving Average; SMA) ใช้เป็นวิธีฐาน วิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์ เหมาะสำหรับข้อมูลที่มีแนวโน้ม วิธีของครอสตัน เหมาะสำหรับอุปสงค์ไม่ต่อเนื่อง และตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ เหมาะสำหรับข้อมูลอนุกรมเวลาที่มีโครงสร้างซับซ้อน ทั้งนี้ เพื่อให้การเปรียบเทียบมีความเป็นระบบ ผู้วิจัยกำหนดช่วงพารามิเตอร์ของแต่ละวิธีล่วงหน้า และคัดเลือกค่าที่เหมาะสมจากข้อมูลชุดเรียนรู้ก่อนนำไปทดสอบกับข้อมูล พ.ศ. 2566 โดยมีรายละเอียดดังนี้

2.3.1 ค่าเฉลี่ยเคลื่อนที่อย่างง่าย (SMA)

เป็นวิธีพยากรณ์พื้นฐานที่คำนวณค่าเฉลี่ยของข้อมูลย้อนหลังจำนวน w ช่วงเวลา โดยให้น้ำหนักเท่ากันทุกช่วง ซึ่งช่วยให้ข้อมูลเรียบและลดความผันผวนระยะสั้นของอนุกรมเวลา [12] งานวิจัยนี้ใช้วิธีค่าเฉลี่ยเคลื่อนที่อย่างง่ายเป็นวิธีฐาน (Baseline) เพื่อเปรียบเทียบกับแบบจำลองที่มีความซับซ้อนสูงกว่า โดยกำหนดช่วงค่าพารามิเตอร์ w เท่ากับ 1-12 เดือน และเลือกค่าที่เหมาะสมจากข้อมูลชุดเรียนรู้ก่อนนำไปพยากรณ์ข้อมูลชุดทดสอบ ค่าพยากรณ์ในแต่ละช่วงเวลาคำนวณตามสมการ (4)

$$SMA_t = \frac{X_{t-1} + X_{t-2} + X_{t-3} + \dots + X_{t-n-1}}{n} \quad (4)$$

โดยที่ X_t คือ ค่าจริงของข้อมูล ณ เวลา t

n คือ จำนวนคาบในหน้าต่างค่าเฉลี่ยเคลื่อนที่

2.3.2 วิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์ (Holt's Exponential Smoothing Method)

เป็นแบบจำลองที่ขยายแนวคิดของการทำให้เรียบแบบเอกซ์โพเนนเชียลเพื่อรองรับข้อมูลที่มีแนวโน้ม โดยประมาณค่าระดับ (Level) และแนวโน้ม (Trend) โครงสร้างของวิธีการประกอบด้วยสมการหลักสามส่วน ได้แก่ สมการพยากรณ์ (Forecast Equation) สมการระดับ (Level Equation) และสมการแนวโน้ม (Trend Equation) ซึ่งแสดงในสมการ (5)-(7) ตามลำดับ [4] ในงานวิจัยนี้กำหนดขอบเขตการค้นหาค่าพารามิเตอร์อย่างเป็นระบบ ได้แก่ α สำหรับ level และ β สำหรับ trend ในช่วง 0-1 [13] โดยใช้วิธีค้นหาเชิงกริด (Grid Search) เพื่อคัดเลือกชุดพารามิเตอร์ที่เหมาะสมที่สุดในแต่ละกรณีสมการพยากรณ์

$$\hat{y}_{t+h|t} = L_t + hB_t \quad (5)$$

สมการระดับ

$$L_t = \alpha y_t + (1 - \alpha)(L_{t-1} + B_{t-1}) \quad (6)$$

สมการแนวโน้ม

$$B_t = \beta(L_t - L_{t-1}) + (1 - \beta)B_{t-1} \quad (7)$$

โดยที่ α คือ smoothing parameter ของระดับ โดยที่ $0 \leq \alpha \leq 1$

β คือ smoothing parameter ของแนวโน้ม โดยที่ $0 \leq \beta \leq 1$

L_t คือ ค่าระดับที่เวลา t

B_t คือ ค่าความชันที่เวลา t

h คือ จำนวนช่วงเวลาที่ต้องการพยากรณ์ล่วงหน้า

2.3.3 วิธีของครอสตัน (Croston's Method)

วิธีของครอสตันเป็นวิธีการพยากรณ์ที่พัฒนาขึ้นสำหรับข้อมูลอุปสงค์ไม่ต่อเนื่องโดยแยกการประมาณขนาดอุปสงค์ และช่วงเวลาระหว่างการเกิดอุปสงค์ออกจากกัน ในการวิจัยนี้ได้พิจารณาวิธีการพยากรณ์จำนวน 3 รูปแบบ ได้แก่ วิธีของครอสตัน วิธีการประมาณของซินเตโตส-บอยแลน (Syntetos-Boylan Approximation; SBA) และ วิธีของทอยน์เทอร์-ซินเตโตส-บาไบ (Teunter-Syntetos-Babai Method; TSB) เพื่อให้ครอบคลุมลักษณะของอุปสงค์ไม่ต่อเนื่องที่มีความหลากหลาย โดยกำหนดช่วงค่าพารามิเตอร์การทำให้เรียบ ไว้ที่ $0.05 \leq \alpha \leq 0.1$ และ $0.05 \leq \beta \leq 0.1$ สำหรับวิธี TSB [14] ทั้งนี้ ใช้วิธีการค้นหาแบบสุ่ม (Randomized Search) ในขั้นต้นเพื่อสำรวจบริเวณของค่าพารามิเตอร์ที่เหมาะสม ก่อนดำเนินการปรับละเอียดด้วยการค้นหาเชิงกริด โดยสมการของวิธีของครอสตัน วิธี SBA และวิธี TSB แสดงดังสมการ (8) (9) และ (10) ตามลำดับ

วิธีของครอสตัน

$$\hat{y}_t = \frac{z_t}{p_t} \quad (8)$$

วิธีการประมาณของซินเตโตส-บอยแลน

$$\hat{y}_t^{SBA} = \left(1 - \frac{\alpha}{2}\right) \cdot \frac{z_t}{p_t} \quad (9)$$

วิธีของทอยน์เทอร์-ซินเตโตส-บาไบ

$$\hat{y}_t^{TSB} = s_t \cdot p_t \quad (10)$$

โดยที่ y_t คือ ความต้องการจริงในงวด t

z_t คือ ขนาดความต้องการที่ปรับเรียบ (Smoothed Demand Size)

p_t คือ ช่วงระหว่างความต้องการที่ปรับเรียบ (Smoothed Demand Interval)

2.3.4 ตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ (ARIMA)

เป็นแบบจำลองสำหรับข้อมูลอนุกรมเวลา โดยผสมส่วนถดถอยอัตโนมัติ (Autoregressive; AR) การหาผลต่าง (Integrated; I) และค่าเฉลี่ยเคลื่อนที่ (Moving Average; MA) เพื่อรองรับข้อมูลที่มีแนวโน้ม ความไม่คงที่ หรือโครงสร้างเชิงเวลาที่ซับซ้อน [15] ในงานวิจัยนี้เลือกใช้ ARIMA เนื่องจากข้อมูลอุปสงค์จะไหลสำรอกในแต่ละกลุ่มอาจมีลักษณะอนุกรมเวลาที่แตกต่างกัน และอาจมีทั้งกรณีที่ใกล้เคียงข้อมูลคงที่ (Stationary) และไม่คงที่ (Non-stationary)

อย่างไรก็ตามงานวิจัยนี้ใช้ส่วน Differencing ของแบบจำลองเป็นกลไกในการรองรับความไม่คงที่ของข้อมูล และกำหนดพารามิเตอร์ของแบบจำลองอย่างเป็นระบบโดยพิจารณา $p \in [0,3]$, $d \in [0,2]$ และ $q \in [0,3]$ [4] โดยใช้วิธีการค้นหาเชิงกริดเพื่อคัดเลือกชุดพารามิเตอร์ที่เหมาะสมที่สุดในแต่ละกรณี ทั้งนี้ โครงสร้างทั่วไปของตัวแบบ ARIMA ที่ใช้ในการศึกษานี้แสดงดังสมการ (11)

$$\begin{aligned} \Delta_d y_t = & c + \phi_y \Delta_d y_{t-1} + \dots \\ & + \phi_p \Delta_d y_{t-p} + \varepsilon_t \\ & - \theta_1 \varepsilon_{t-1} - \dots \\ & - \theta_q \varepsilon_{t-q} \end{aligned} \quad (11)$$

โดยที่ y คือ ข้อมูลอนุกรมเวลาที่ใช้ในการพยากรณ์

t คือ ช่วงเวลาที่ต้องการพยากรณ์

c คือ ค่าคงที่

ϕ คือ ตัวแปรคงที่ (Parameter) สำหรับ Autoregressive

θ คือ ตัวแปรคงที่ (Parameter) สำหรับ Moving Average

p คือ จำนวนจุดข้อมูลก่อนหน้าที่ใช้ในการพยากรณ์

q คือ จำนวนค่าความคลาดเคลื่อนของจุดข้อมูลก่อนหน้าที่ใช้ในการพยากรณ์

d คือ จำนวนครั้งของการหาผลต่าง

ε คือ ค่าความคลาดเคลื่อน ณ ช่วงเวลาที่ต้องการพยากรณ์

ภายหลังจากกำหนดวิธีการจัดกลุ่มวิธีการพยากรณ์ และช่วงค่าพารามิเตอร์ของแต่ละแบบจำลองแล้ว งานวิจัยนี้ดำเนินการคัดเลือกแบบจำลองอย่างเป็นระบบ โดยมีเป้าหมายเพื่อคัดเลือกของวิธีการจัดกลุ่มร่วมกับวิธีการพยากรณ์ที่ให้ผลดีที่สุดในการพยากรณ์ ในขั้นแรกผู้วิจัยใช้กระบวนการค้นหาเชิงกริดเพื่อคัดเลือกค่าพารามิเตอร์ที่เหมาะสมสำหรับแต่ละวิธีการพยากรณ์ ภายใต้วงค่าที่กำหนดไว้ จากนั้นจึงนำแบบจำลองที่ได้จากแต่ละคู่ของวิธีการจัดกลุ่มและวิธีการพยากรณ์ไปประเมินและเปรียบเทียบในขั้นตอนถัดไป

2.4 การเลือกรูปแบบการจัดกลุ่มและวิธีพยากรณ์ที่เหมาะสมกับบ่อไหลสร้าง

ขั้นตอนสุดท้ายของการศึกษาเป็นการประเมินความแม่นยำของแบบจำลองการพยากรณ์ด้วยค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร ซึ่งเป็นตัวชี้วัดทางสถิติที่พัฒนาขึ้นจากค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ยร้อยละ (Mean Absolute Percentage Error; MAPE) เพื่อลดปัญหาความไม่สมมาตรของความคลาดเคลื่อนระหว่างค่าจริงและค่าพยากรณ์ ค่าดังกล่าวคำนวณตามสมการ (12) โดยพิจารณาส่วนต่างสัมบูรณ์

ระหว่างค่าจริงและค่าพยากรณ์เทียบกับค่าเฉลี่ยของทั้งสองค่าในแต่ละช่วงเวลา ส่งผลให้ค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรมีช่วงค่าอยู่ระหว่าง 0% ถึง 200% และเหมาะสมสำหรับการเปรียบเทียบประสิทธิภาพของแบบจำลองภายใต้ข้อมูลอุปสงค์ที่มีความไม่สม่ำเสมอ [16]

$$sMAPE = \frac{2}{h} \sum_{t=n+1}^{n+h} \frac{|Y_t - \hat{Y}_t|}{|Y_t| + |\hat{Y}_t|} \times 100(\%) \quad (12)$$

โดยที่ Y_t คือ ค่าจริงของตัวแปรในงวด t หรือค่าความต้องการจริงในช่วงเวลานั้น

$\hat{Y}_t$ คือ ค่าพยากรณ์ของ Y_t ที่ได้จากแบบจำลอง

h คือ จำนวนงวดที่ทำการพยากรณ์ล่วงหน้า (ช่วงคาบพยากรณ์)

n คือ จำนวนข้อมูลในชุดตัวอย่างที่ใช้สร้างโมเดล (ก่อนเริ่มพยากรณ์)

3. ผลการทดลองและอภิปรายผลการทดลอง

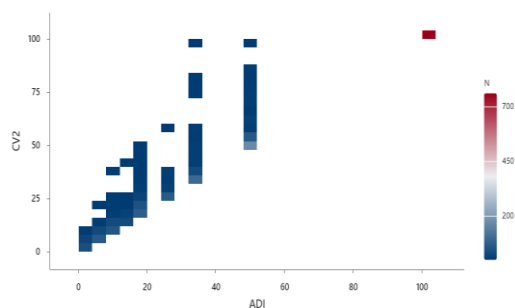
3.1 ผลการจัดกลุ่มบ่อไหลสร้าง

จากการจัดกลุ่มบ่อไหลสร้างจะมีผลลัพธ์ทั้งหมดดังนี้

3.1.1 ผลลัพธ์จากการจัดกลุ่มตามรูปแบบอุปสงค์

จากการวิเคราะห์กราฟ Hexbin Plot ดังแสดงในรูปที่ 3 ซึ่งแสดงความหนาแน่นของข้อมูลระหว่างค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนพบว่าข้อมูลส่วนใหญ่กระจุกตัวในช่วงที่ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์อยู่ในระดับต่ำถึงปานกลาง ขณะที่ค่ากำลังสองสัมประสิทธิ์ความแปรปรวนมีแนวโน้มเพิ่มขึ้นตามค่าดังกล่าว สะท้อนว่าบ่อไหลสร้างส่วนใหญ่มีลักษณะอุปสงค์ไม่สม่ำเสมอ และหลายรายการมีแนวโน้มใกล้เคียงกับ

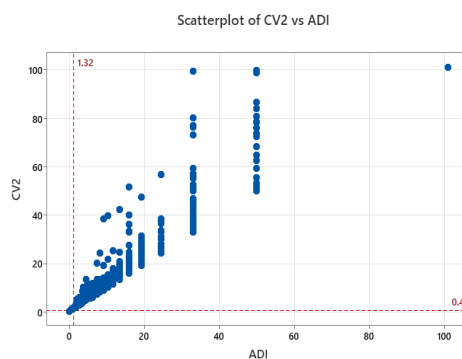
อุปสงค์รูปแบบเป็นก้อน กล่าวคือ มีการเบิกใช้ไม่บ่อย แต่มีความผันผวนของปริมาณการใช้งานสูง นอกจากนี้ ยังพบจุดข้อมูลบางส่วนที่มีค่าของทั้งสองตัวแปรสูงมาก ซึ่งสะท้อนความเบี่ยงเบนของข้อมูลในระบบอะไหล่สำรองได้อย่างชัดเจน ลักษณะดังกล่าวสอดคล้องกับที่ [3] พบในการศึกษาที่มีบริบทใกล้เคียงกัน ซึ่งชี้ว่าข้อมูลอะไหล่สำรองในโรงงานอุตสาหกรรมมักมีการกระจุกตัวในโซนอุปสงค์รูปแบบเป็นก้อนเป็นส่วนใหญ่



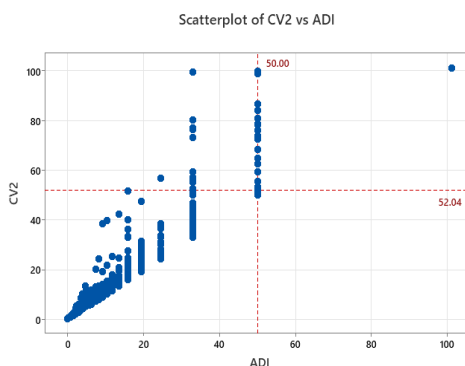
รูปที่ 3 แผนภาพ Hexbin Plot ของค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวน

รูปที่ 4 และรูปที่ 5 ซึ่งใช้ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนในการจำแนกรูปแบบอุปสงค์ของอะไหล่สำรองพบว่าเกณฑ์ดั้งเดิมที่ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์มีค่าเท่ากับ 1.32 และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนมีค่าเท่ากับ 0.49 ดังแสดงในรูปที่ 4 ทำให้อะไหล่สำรองร้อยละ 99.6 ถูกจัดอยู่ในกลุ่มอุปสงค์รูปแบบเป็นก้อน และร้อยละ 0.4 อยู่ในกลุ่มอุปสงค์รูปแบบผันผวน ดังตารางที่ 4 สะท้อนว่าเกณฑ์ดังกล่าวไม่สามารถแยกความแตกต่างของพฤติกรรมอุปสงค์ได้อย่างมีประสิทธิภาพ ซึ่งเป็นข้อจำกัดที่ [3] ได้ระบุไว้

เช่นกันว่าการใช้ค่าเกณฑ์ตัดสินใจดั้งเดิมอาจไม่เหมาะสมเมื่อนำไปประยุกต์กับข้อมูลจริงที่มีการกระจายตัวต่างไปจากสมมติฐานเดิม การใช้ค่ามัธยฐานเป็นเกณฑ์ใหม่ที่ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์มีค่าเท่ากับ 50 และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนมีค่าเท่ากับ 52.04 ดังแสดงในรูปที่ 5 ช่วยให้ข้อมูลกระจายตัวสมดุลขึ้นและจำแนกได้ครบทั้ง 4 กลุ่ม โดยรายละเอียดจำนวนอะไหล่สำรองในแต่ละกลุ่มแสดงดังตารางที่ 5 จึงเหมาะสมกว่าสำหรับการอธิบายโครงสร้างอุปสงค์ของอะไหล่สำรองในชุดข้อมูลนี้



รูปที่ 4 แผนภาพการกระจายของข้อมูลและค่าการตัดสินใจตามทฤษฎี



รูปที่ 5 แผนภาพการกระจายของข้อมูลและค่าการตัดสินใจตามค่ามัธยฐาน

ตารางที่ 4 จำนวนอะไหล่สำรองจากการจัดกลุ่มตามรูปแบบอุปสงค์

กลุ่ม	จำนวน	ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ (ADI)	ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน (CV^2)
อุปสงค์รูปแบบเป็นก้อน (Lumpy Demand)	1,728	50.00	52.99
อุปสงค์รูปแบบผันผวน (Erratic Demand)	7	0.92	1.47

ตารางที่ 5 จำนวนอะไหล่สำรองจากการจัดกลุ่มตามรูปแบบอุปสงค์โดยใช้ค่ามัธยฐาน

กลุ่ม	จำนวน	ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ (ADI)	ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน (CV^2)
อุปสงค์รูปแบบเป็นก้อน (Lumpy Demand)	860	101.00	101.00
อุปสงค์รูปแบบราบเรียบ (Smooth Demand)	694	16.00	19.40
อุปสงค์รูปแบบไม่สม่ำเสมอ (Intermittent Demand)	168	50.00	50.00
อุปสงค์รูปแบบผันผวน (Erratic Demand)	13	33.00	59.56

3.1.2 ผลลัพธ์จากการจัดกลุ่มแบบเค-มินส์

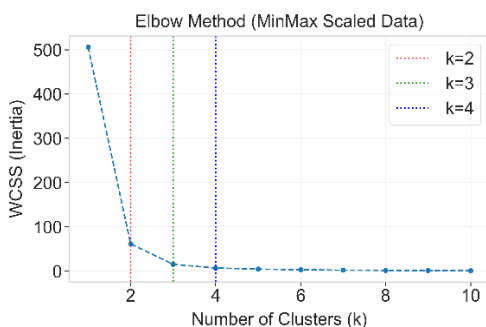
การกำหนดจำนวนคลัสเตอร์ของวิธีการจัดกลุ่มแบบเค-มินส์ ในงานวิจัยนี้พิจารณาจากกราฟ Elbow Method ร่วมกับค่า Silhouette Score และการตีความผลเชิงเปรียบเทียบระหว่างวิธีจัดกลุ่ม โดยรูปที่ 6 แสดงว่าค่า WCSS ลดลงชัดเจนในช่วงต้นและเริ่มชะลอตัวตั้งแต่วิธีเลือก $k = 3$ เป็นต้นไป สะท้อนว่าโครงสร้างข้อมูลไม่จำเป็นต้องเพิ่มจำนวนกลุ่มมากกว่านี้มากนัก ขณะเดียวกันรูปที่ 7-9 ซึ่งแสดงผลการจัดกลุ่มที่ $k = 2, 3$ และ 4 ช่วยให้เห็นว่าการเพิ่มจำนวนกลุ่มทำให้สามารถจำแนกความแตกต่างของรูปแบบอุปสงค์ได้ละเอียดขึ้น โดยค่า Silhouette Score ของ $k = 2, 3$ และ 4 เท่ากับ 0.8429, 0.8297 และ 0.8423 ตามลำดับ ซึ่งล้วนอยู่ในระดับสูงแสดงถึงคุณภาพการจัดกลุ่มที่ดี สอดคล้องกับแนวทางที่ [2] ใช้การจัดกลุ่มแบบเค-มินส์ เพื่อแบ่งข้อมูลก่อนการพยากรณ์ และพบว่าโครงสร้างกลุ่มที่ชัดเจนส่งผลให้แบบจำลองมีประสิทธิภาพสูงขึ้น อย่างไรก็ตามแม้ $k = 2$ จะให้ค่าสูงที่สุดเล็กน้อย แต่ยังเป็นการจำแนกในระดับหยาบ ขณะที่ $k = 4$ ให้ค่าที่ใกล้เคียงกัน แต่ให้รายละเอียดของโครงสร้างข้อมูลมากกว่า และสอดคล้องกับกรอบการเปรียบเทียบกับการจัดกลุ่มตามรูปแบบอุปสงค์ซึ่งแบ่งออกเป็น 4 กลุ่มหลัก ดังนั้นงานวิจัยนี้จึงเลือก $k = 4$ เป็นจำนวนคลัสเตอร์ที่เหมาะสมสำหรับการจัดกลุ่มอะไหล่สำรองด้วยวิธีการจัดกลุ่มแบบเค-มินส์

ผลการจัดกลุ่มในตารางที่ 6 สะท้อนความแตกต่างของโครงสร้างอุปสงค์อย่างชัดเจนในสองมิติหลัก ได้แก่ ความถี่การเบิกใช้และความผันผวนของอุปสงค์ โดยกลุ่ม 0 มีค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนมัธยฐานสูงสุด แสดงถึง

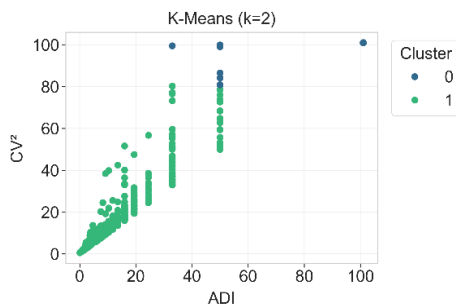
อะไหล่ที่มีการเบิกใช้ทางและผันผวนสูงมาก ขณะที่กลุ่ม 2 อยู่ในระดับปานกลาง ส่วนกลุ่มที่มีความผันผวนของคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนเท่ากับ 24.50 และ 33.00 สะท้อนการเบิกใช้ที่ถี่ขึ้นและมีความผันผวนปานกลางสำหรับกลุ่มที่มีความผันผวนของทั้งสองตัวแปรต่ำที่สุด บ่งชี้ถึงอะไหล่ที่มีการเบิกใช้ค่อนข้างถี่และมีความแปรปรวนน้อยกว่ากลุ่มอื่น การจำแนกดังกล่าวลดความหยابของการแบ่งกลุ่มเมื่อเทียบกับ $k = 2$ และแสดงความแตกต่างของข้อมูลได้ชัดเจนกว่ากรณี $k = 3$ ตามรูปที่ 7-9 จึงเอื้อต่อการเลือกแบบจำลองพยากรณ์และการเปรียบเทียบประสิทธิภาพระหว่างวิธีการจัดกลุ่มอย่างเป็นระบบ

ตารางที่ 6 จำนวนอะไหล่สำรองจากการจัดกลุ่มแบบเค-มินส์

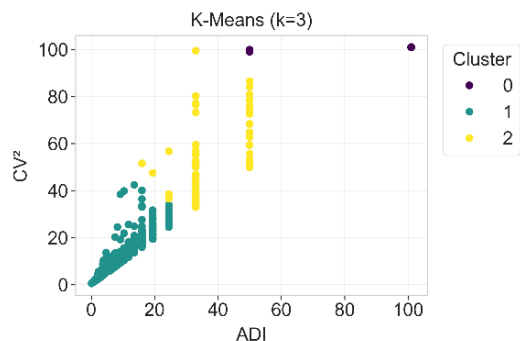
กลุ่ม	จำนวน	ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ (ADI)	ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน (CV^2)
K = 2	กลุ่ม 0	767	101.00
	กลุ่ม 1	968	24.50
K = 3	กลุ่ม 0	761	101.00
	กลุ่ม 1	537	13.57
	กลุ่ม 2	437	50.00
K = 4	กลุ่ม 0	759	101.00
	กลุ่ม 1	291	24.50
	กลุ่ม 2	289	50.00
	กลุ่ม 3	396	11.567



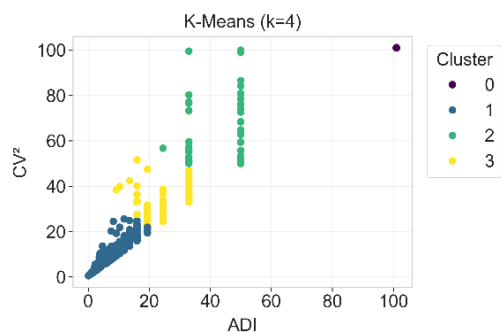
รูปที่ 6 กราฟ Elbow Method ของการจัดกลุ่มแบบเค-มินส์



รูปที่ 7 การจัดกลุ่มแบบเค-มินส์ โดยที่ $k = 2$



รูปที่ 8 การจัดกลุ่มแบบเค-มินส์ โดยที่ $k = 3$



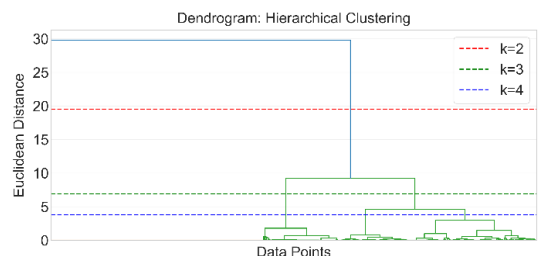
รูปที่ 9 การจัดกลุ่มแบบเค-มินส์ โดยที่ $k = 4$

3.1.3 ผลลัพธ์จากการจัดกลุ่มแบบลำดับขั้น

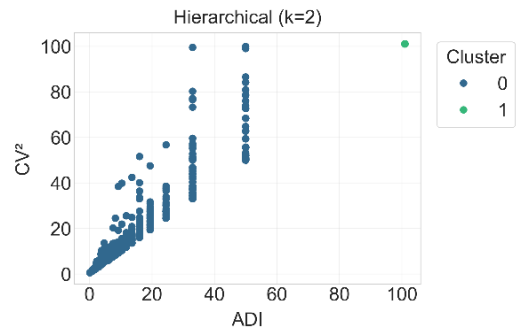
การจัดกลุ่มแบบลำดับขั้นในงานวิจัยนี้ใช้วิธีวอร์ด (Ward's Method) ร่วมกับระยะทางแบบยูคลิด (Euclidean Distance) โดยพิจารณาแผนภาพ Dendrogram ในรูปที่ 10 ร่วมกับค่า Silhouette Score เพื่อประเมินความเหมาะสมของจำนวนคลัสเตอร์ ขณะเดียวกัน รูปที่ 11 รูปที่ 12 และรูปที่ 13 ซึ่งแสดงผลการจัดกลุ่มเมื่อกำหนด $k = 2, 3$ และ 4 ตามลำดับ ช่วยให้เห็นว่าการเพิ่มจำนวนกลุ่มทำให้โครงสร้างข้อมูลถูกแยกละเอียดขึ้น โดยค่า Silhouette Score ของ $k = 2, 3$ และ 4 เท่ากับ 0.8227 0.7878 และ 0.7451 ตามลำดับ ซึ่งยังสะท้อนคุณภาพการจัดกลุ่มที่ดี ส่วนค่า Cut Height เท่ากับ 19.5040 6.9054 และ 3.7827 ตามลำดับ แสดงถึงการแยกโครงสร้างข้อมูลที่ละเอียดขึ้นเมื่อจำนวนกลุ่มเพิ่มขึ้น อย่างไรก็ตาม แม้ $k = 3$ จะให้ความเหมาะสมมากกว่าในเชิงโครงสร้างของวิธีการจัดกลุ่มแบบลำดับขั้น แต่งานวิจัยนี้เลือกใช้ $k = 4$ เพื่อให้สอดคล้องกับกรอบการเปรียบเทียบของวิธีอื่น โดยค่า Silhouette Score ของ $k = 4$ แม้ต่ำกว่า $k = 2$ และ $k = 3$ แต่ยังให้รายละเอียดของโครงสร้างอุปสงค์เพียงพอสำหรับการวิเคราะห์เปรียบเทียบอย่างเป็นระบบ

ผลการจัดกลุ่มในตารางที่ 7 ชี้ให้เห็นว่ากลุ่มที่มีค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนมัธยฐานสูงที่สุด สะท้อนอะไหล่ที่มีการเบิกใช้ห่างและมีความผันผวนสูงมาก ขณะที่กลุ่มที่มีค่ามัธยฐานเท่ากับ 50.00 และ 50.00 แสดงถึงอะไหล่ที่ยังมีอุปสงค์ไม่สม่ำเสมอ แต่มีความถี่และความผันผวนต่ำกว่ากลุ่มแรก ส่วนกลุ่มที่มีค่ามัธยฐานเท่ากับ 24.50 และ 26.06 สะท้อนอะไหล่ที่มีการเบิกใช้ถี่ขึ้นและมีความผันผวนปานกลาง ขณะที่กลุ่มที่มี

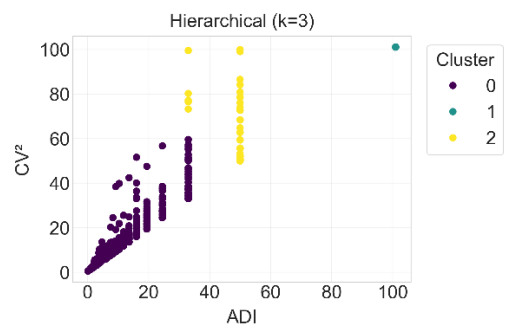
ค่ามัธยฐานต่ำที่สุดบ่งชี้ถึงอะไหล่ที่มีการเบิกใช้ค่อนข้างถี่และมีความแปรปรวนน้อยที่สุด การจำแนกในลักษณะนี้ช่วยลดความหยابของการแบ่งกลุ่มเมื่อเทียบกับ $k = 2$ และเพิ่มความละเอียดของโครงสร้างอุปสงค์เมื่อเทียบกับ $k = 3$ จึงเอื้อต่อการเปรียบเทียบประสิทธิภาพของวิธีการจัดกลุ่มและการเลือกแบบจำลองพยากรณ์ให้เหมาะสมกับลักษณะข้อมูลของแต่ละกลุ่ม



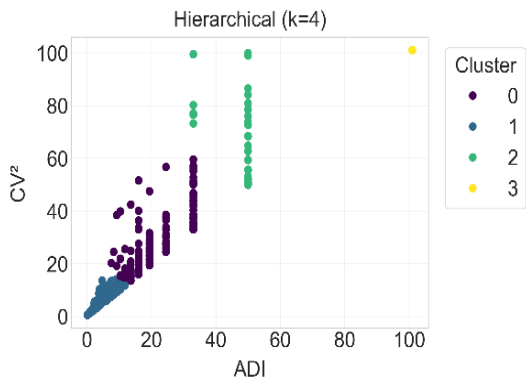
รูปที่ 10 Dendrogram ของการจัดกลุ่มแบบลำดับขั้น



รูปที่ 11 การจัดกลุ่มแบบลำดับขั้น โดยที่ $k = 2$



รูปที่ 12 การจัดกลุ่มแบบลำดับขั้น โดยที่ $k = 3$



รูปที่ 13 การจัดกลุ่มแบบลำดับขั้น โดยที่ $k = 4$

ตารางที่ 7 จำนวนอะไหล่สำรองจากการจัดกลุ่มแบบลำดับขั้น

กลุ่ม	จำนวน	ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์		ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน	
		อุปสงค์		แปรปรวน	
		(ADI)		(CV ²)	
K = 2	กลุ่ม 0	759	101.00	101.00	
	กลุ่ม 1	976	24.50	27.56	
K = 3	กลุ่ม 0	759	101.00	101.00	
	กลุ่ม 1	274	50.00	50.00	
	กลุ่ม 2	702	16.00	19.40	
K = 4	กลุ่ม 0	759	101.00	101.00	
	กลุ่ม 1	274	50.00	50.00	
	กลุ่ม 2	240	6.29	7.82	
	กลุ่ม 3	462	24.50	26.06	

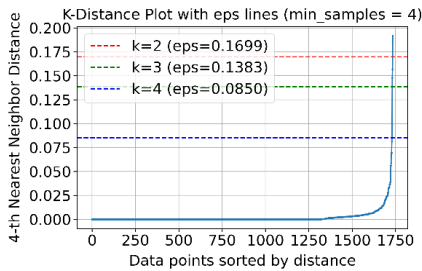
3.1.4 ผลลัพธ์จากการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมน

งานวิจัยนี้กำหนดพารามิเตอร์หลัก 2 ค่า ได้แก่ ค่า ε และ MinPts โดยกำหนด MinPts เท่ากับ 2 เพื่อให้สอดคล้องกับการสร้างกราฟ K-Distance Plot จากเพื่อนบ้านลำดับที่ 2 ซึ่งใช้ประกอบการเลือกค่า ε ที่

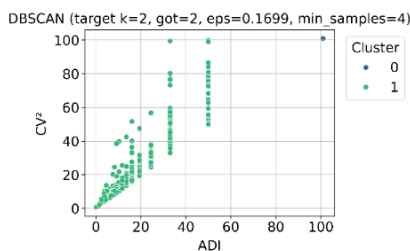
เหมาะสม ดังแสดงในรูปที่ 14 ผลการทดลองพบว่าค่า ε มีอิทธิพลโดยตรงต่อโครงสร้างการจัดกลุ่ม กล่าวคือเมื่อกำหนด $\varepsilon = 0.1699$ ตามรูปที่ 15 ข้อมูลถูกจัดรวมเป็นกลุ่มขนาดใหญ่เพียงไม่กี่กลุ่ม สะท้อนว่ารัศมีเพื่อนบ้านมีขนาดกว้างและลดความสามารถในการแยกความแตกต่างของข้อมูล ขณะที่ $\varepsilon = 0.1383$ ตามรูปที่ 16 ให้การจำแนกที่ชัดเจนขึ้น โดยสามารถแยกได้ 3 กลุ่ม และมีจุดรวมน 1 รายการ แสดงถึงความสามารถในการแยกโครงสร้างข้อมูลที่ดีขึ้น

เมื่อกำหนด $\varepsilon = 0.0850$ ตามรูปที่ 17 พบว่าให้โครงสร้างกลุ่มที่เหมาะสมที่สุดสำหรับข้อมูลชุดนี้ เนื่องจากสามารถสะท้อนความแตกต่างของข้อมูลได้อย่างชัดเจน โดยจำแนกออกเป็น 4 กลุ่ม และมีจุดรวมน 3 รายการ ซึ่งจุดรวมนดังกล่าวสะท้อนจุดเด่นของวิธีนี้ที่สามารถแยกข้อมูลที่มีพฤติกรรมแตกต่างออกจากกลุ่มหลักได้โดยตรง ส่งผลให้กลุ่มที่ได้มีความเป็นเนื้อเดียวกันมากขึ้น รายละเอียดของจำนวนอะไหล่สำรอง ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนของแต่ละกลุ่มแสดงในตารางที่ 8

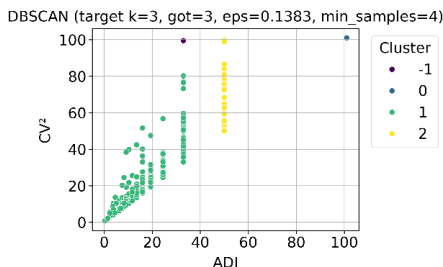
เมื่อพิจารณาคุณลักษณะของทั้ง 4 กลุ่ม พบว่ากลุ่ม 0 มีค่ามัธยฐานของคาบเวลาเฉลี่ยระหว่างอุปสงค์ และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนสูงที่สุด สะท้อนอุปสงค์ที่เกิดห่างและมีความผันผวนสูง ขณะที่กลุ่ม 1 มีค่าของทั้งสองตัวแปรต่ำที่สุด จึงเป็นกลุ่มที่มีอุปสงค์สม่ำเสมอมากที่สุด ส่วนกลุ่ม 2 มีลักษณะอยู่ในระดับปานกลางระหว่างสองกลุ่มดังกล่าว ขณะที่กลุ่ม 3 เป็นกลุ่มย่อยขนาดเล็กที่มีความผันผวนสูงเฉพาะตัว แม้ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์จะไม่สูงมากนัก ซึ่งสะท้อนลักษณะอุปสงค์ที่มีความไม่แน่นอนเฉพาะกลุ่ม



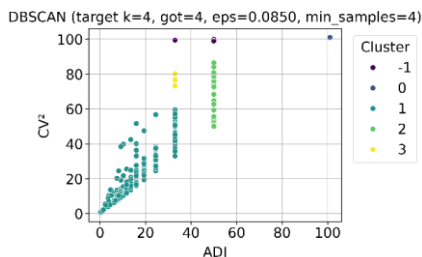
รูปที่ 14 K-Distance Plot สำหรับการเลือกค่าพารามิเตอร์ ϵ ของ DBSCAN



รูปที่ 15 ผลการจัดกลุ่มด้วยวิธี DBSCAN เมื่อปรับค่า ϵ ให้ได้จำนวนคลัสเตอร์ใกล้เคียง 2 กลุ่ม



รูปที่ 16 ผลการจัดกลุ่มด้วยวิธี DBSCAN เมื่อปรับค่า ϵ ให้ได้จำนวนคลัสเตอร์ใกล้เคียง 3 กลุ่ม



รูปที่ 17 ผลการจัดกลุ่มด้วยวิธี DBSCAN เมื่อปรับค่า ϵ ให้ได้จำนวนคลัสเตอร์ใกล้เคียง 4 กลุ่ม

โดยสรุปผลลัพธ์ดังกล่าวแสดงให้เห็นว่าการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรบกวนสามารถจำแนกโครงสร้างข้อมูลละเอียดได้อย่างมีประสิทธิภาพทั้งในมิติของความถี่และความผันผวนของอุปสงค์ อีกทั้งยังสามารถจัดการข้อมูลผิดปกติได้อย่างเหมาะสม

ตารางที่ 8 จำนวนอะไหล่สำรองจากการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรบกวน

กลุ่ม	จำนวน	ค่าค่าเวลาเฉลี่ยระหว่างอุปสงค์		ค่ากำลังสองสัมประสิทธิ์ความแปรปรวน
		(ADI)		(CV ²)
จำนวน	กลุ่ม 0	759	101.00	101.00
คลัสเตอร์	กลุ่ม 1	976	24.50	27.56
(k = 2)	กลุ่ม 0	759	101.00	101.00
คลัสเตอร์	กลุ่ม 1	706	16.00	19.40
(k = 3)	กลุ่ม 2	269	50.00	50.00
จำนวนจุดรบกวน 1 รายการ				
จำนวน	กลุ่ม 0	759	101.00	101.00
คลัสเตอร์	กลุ่ม 1	702	16.00	19.40
(k = 4)	กลุ่ม 2	267	50.00	50.00
	กลุ่ม 3	4	33.00	76.81
จำนวนจุดรบกวน 3 รายการ				

หมายเหตุ: สำหรับการจัดกลุ่มแบบเค-มีนส์ และการจัดกลุ่มแบบลำดับขั้นจำนวนคลัสเตอร์ k เป็นค่าพารามิเตอร์อินพุตของวิธีการ ในขณะที่การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรบกวนค่า k ในตารางนี้หมายถึงจำนวนคลัสเตอร์ที่ได้จากการปรับค่าพารามิเตอร์ ϵ ให้ผลลัพธ์ใกล้เคียง 2 3 และ 4 กลุ่มตามลำดับ

3.2 ผลการพยากรณ์และผลการวัดค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร

การประเมินประสิทธิภาพของการจัดกลุ่มและการพยากรณ์ในงานวิจัยนี้ไม่ได้พิจารณาเพียงด้านใดด้านหนึ่งเท่านั้น แต่พิจารณาร่วมกันระหว่างความแม่นยำของ



ตารางที่ 9 ผลการวัดค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรของวิธีพยากรณ์ด้วยวิธีค่าเฉลี่ยเคลื่อนที่
อย่างง่าย

วิธีการ	จำนวนกลุ่ม	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3	กลุ่มที่ 4
การจัดกลุ่มแบบเค-มีนส์	2	4.30	36.75	-	-
	3	4.31	52.55	16.87	-
	4	4.33	27.91	12.70	60.09
การจัดกลุ่มแบบลำดับชั้น	2	4.33	36.47	-	-
	3	4.33	11.21	46.32	-
	4	4.33	11.21	68.29	34.91
การจัดกลุ่มเชิงพื้นที่โดยอาศัย	2	4.33	36.47	-	-
ความหนาแน่นและจุดรวมวง	3	4.33	46.08	11.30	-
	4	4.33	46.32	11.38	4.17
การจัดกลุ่มตามรูปแบบอุปสงค์ (ค่ามัธยฐาน)		อุปสงค์รูปแบบ เป็นก้อน	อุปสงค์รูปแบบ ราบเรียบ	อุปสงค์รูปแบบ ผันผวน	อุปสงค์รูปแบบ ไม่สม่ำเสมอ
		4.83	46.83	3.85	12.90

การพยากรณ์ที่วัดด้วยค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร และความเหมาะสมของการจำแนกกลุ่ม สามารถอธิบายได้จากกราฟ Scatter Plot, Elbow Method, Dendrogram, K-Distance Plot และจากการพยากรณ์ด้วยวิธีค่าเฉลี่ยเคลื่อนที่อย่างง่าย วิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์ วิธีของโครสตัน และตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ได้ผลลัพธ์ดังตารางที่ 9 10 11 และ 12 ตามลำดับ

จากการวิเคราะห์ค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรของการพยากรณ์ด้วยวิธีค่าเฉลี่ยเคลื่อนที่อย่างง่ายตามตารางที่ 9 พบว่าการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมวงจำนวน 4 กลุ่มให้ผลแม่นยำที่สุด โดยมีค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรเฉลี่ยต่ำสุดที่ 16.55% สะท้อนว่าการจัดกลุ่มดังกล่าวช่วยลดความแตกต่างภายในกลุ่มและเสริมให้วิธีค่าเฉลี่ยเคลื่อนที่อย่างง่ายสร้าง

ค่าพยากรณ์ได้แม่นยำกว่าวิธีจัดกลุ่มอื่น อย่างไรก็ตามแม้วิธีค่าเฉลี่ยเคลื่อนที่อย่างง่ายจะใช้งานง่ายและเหมาะสมในฐานะตัวแบบฐานแต่ค่าความคลาดเคลื่อนรายกลุ่มยังสูงในหลายกรณี โดยเฉพาะกลุ่มที่มีอุปสงค์ผันผวนหรือเกิดไม่สม่ำเสมอ สะท้อนข้อจำกัดของวิธีนี้ซึ่งเหมาะกับข้อมูลที่มีรูปแบบค่อนข้างคงที่มากกว่าสอดคล้องกับ [5] ที่ระบุว่าการเลือกตัวแบบพยากรณ์ควรพิจารณาให้สอดคล้องกับลักษณะอุปสงค์แต่ละประเภท

ผลในตารางที่ 10 ซึ่งให้เห็นว่าภายใต้การเปรียบเทียบที่จำนวน 4 กลุ่ม การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมวงร่วมกับวิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์ให้ความแม่นยำดีที่สุด โดยมีค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรเฉลี่ยต่ำสุดเท่ากับ 97.64% สะท้อนว่าการจัดกลุ่มดังกล่าวสามารถลดความแตกต่างของข้อมูลภายในกลุ่มได้ดีกว่าวิธีอื่น อย่างไรก็ตามค่าความคลาดเคลื่อนโดยรวมยัง

อยู่ในระดับสูงมาก ซึ่งสะท้อนให้เห็นว่าวิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์ ซึ่งออกแบบมาสำหรับข้อมูลที่มีแนวโน้มต่อเนื่องไม่เหมาะสมกับลักษณะอุปสงค์ไม่ต่อเนื่อง

ของอะไหล่สำรอง ทั้งนี้ สอดคล้องกับที่ [4] พบว่าวิธีการทำให้เรียบแบบเอกซ์โพเนนเชียลให้ผลที่ดีน้อยกว่าแบบจำลองที่ซับซ้อนกว่าในบริบทที่ข้อมูลมีความผันผวนสูง

ตารางที่ 10 ผลการวัดค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรของวิธีพยากรณ์ด้วยวิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์

วิธีการ	จำนวนกลุ่ม	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3	กลุ่มที่ 4
การจัดกลุ่มแบบเค-มีนส์	2	59.32	125.06	-	-
	3	58.72	143.97	101.98	-
	4	58.87	128.45	94.90	144.12
การจัดกลุ่มแบบลำดับขั้น	2	58.87	124.88	-	-
	3	58.87	93.78	137.01	-
	4	58.87	93.78	142.39	134.22
การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวม	2	58.87	124.88	-	-
	3	58.87	136.80	93.98	-
	4	58.87	137.01	94.68	100.00
การจัดกลุ่มตามรูปแบบอุปสงค์ (คำมัธยฐาน)		อุปสงค์รูปแบบเป็นก้อน	อุปสงค์รูปแบบราบเรียบ	อุปสงค์รูปแบบผันผวน	อุปสงค์รูปแบบไม่สม่ำเสมอ
		61.61	137.27	102.56	101.08

ตารางที่ 11 ผลการวัดค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรของวิธีพยากรณ์ด้วยวิธีของโครสตัน

วิธีการ	จำนวนกลุ่ม	รูปแบบ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3	กลุ่มที่ 4
การจัดกลุ่มแบบเค-มีนส์	2	Croston	61.93	147.18	-	-
		SBA	61.93	147.20	-	-
		TSB	24.64	146.77	-	-
	3	Croston	61.34	174.60	113.31	-
		SBA	61.34	174.62	113.34	-
		TSB	24.81	173.88	58.14	-
	4	Croston	61.24	147.55	104.38	177.67
		SBA	61.24	147.67	104.38	177.69
		TSB	24.88	82.63	50.17	176.70
การจัดกลุ่มแบบลำดับขั้น	2	Croston	61.24	147.02	-	-
		SBA	61.24	147.03	-	-
		TSB	24.88	146.61	-	-

ธวัชลักษณ์ จินะสี และ ชุติศักดิ์ พรสิงห์, “การเปรียบเทียบรูปแบบการจัดกลุ่มและตัวแบบพยากรณ์ที่เหมาะสมที่สุดสำหรับการพยากรณ์ความต้องการอะไหล่สำรองในงานซ่อมบำรุง.”



ตารางที่ 11 ผลการวัดค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรของวิธีพยากรณ์ด้วยวิธีของครอสตัน (ต่อ)

วิธีการ	จำนวนกลุ่ม	รูปแบบ	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3	กลุ่มที่ 4
การจัดกลุ่มแบบลำดับขั้น (ต่อ)	3	Croston	61.24	103.41	164.04	-
		SBA	61.24	103.41	164.05	-
		TSB	24.88	48.69	163.47	-
	4	Croston	61.24	103.71	173.09	158.91
		SBA	61.24	103.81	173.13	159.34
		TSB	24.88	95.74	173.10	95.71
	2	Croston	61.24	147.02	-	-
		SBA	61.24	147.03	-	-
		TSB	24.88	146.61	-	-
การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดบริเวณ	3	Croston	61.24	163.68	103.77	-
		SBA	61.24	163.69	103.78	-
		TSB	24.88	49.57	163.13	-
	4	Croston	61.24	164.00	103.79	99.98
		SBA	61.24	164.05	103.81	99.98
		TSB	24.888	49.55	163.47	4.17
			อุปสงค์รูปแบบ	อุปสงค์รูปแบบ	อุปสงค์รูปแบบ	อุปสงค์รูปแบบ
			เป็นก้อน	ราบเรียบ	ผันผวน	ไม่สม่ำเสมอ
	การจัดกลุ่มตามรูปแบบอุปสงค์ (ค่ามัธยฐาน)	Croston	65.31	164.47	110.25	108.53
		SBA	65.31	164.48	110.25	108.54
		TSB	27.24	163.89	50.00	51.71

ผลในตารางที่ 11 พบว่าภายใต้การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดบริเวณและวิธีพยากรณ์ด้วย TSB ให้ค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรเฉลี่ยต่ำที่สุดเท่ากับ 60.52% เมื่อเปรียบเทียบกับวิธีการพยากรณ์อื่นภายใต้โครงสร้างการจัดกลุ่มเดียวกัน อย่างไรก็ตามค่าความคลาดเคลื่อนดังกล่าวยังอยู่ในระดับค่อนข้างสูงซึ่งสะท้อนให้เห็นว่าแม้วิธีของครอสตันจะพัฒนามาสำหรับอุปสงค์ไม่ต่อเนื่องโดยเฉพาะ แต่ความซับซ้อนของข้อมูลอะไหล่สำรองที่มีหลายช่วงเวลาที่ไม่เกิดอุปสงค์สลับกับช่วงที่เกิดความ

ต้องการอย่างไม่สม่ำเสมอ และบางรายการยังมีความผันผวนของขนาดอุปสงค์สูง ทำให้การประมาณทั้งความถี่และปริมาณความต้องการทำได้ยาก ซึ่งสอดคล้องกับที่ [6] ชี้ว่าอุปสงค์ที่มีลักษณะเป็นก้อนผันผวนและไม่ต่อเนื่องต้องการแนวทางการพยากรณ์ที่รองรับความไม่แน่นอนได้อย่างเฉพาะเจาะจงมากกว่าวิธีมาตรฐาน

ตารางที่ 12 พบว่าการพยากรณ์ด้วยตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ภายหลังการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดบริเวณให้ผลดีที่สุด โดยมีค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์

ตารางที่ 12 ผลการวัดค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรของวิธีพยากรณ์ด้วยตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่

วิธีการ	จำนวนกลุ่ม	กลุ่มที่ 1	กลุ่มที่ 2	กลุ่มที่ 3	กลุ่มที่ 4
การจัดกลุ่มแบบเค-มีนส์	2	3.00	21.30	-	-
	3	3.00	32.99	6.67	-
	4	3.01	11.86	5.36	39.48
การจัดกลุ่มแบบลำดับชั้น	2	3.01	21.14	-	-
	3	3.01	5.35	27.30	-
	4	3.01	5.35	52.64	14.14
การจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวม	2	3.01	21.14	-	-
	3	3.01	27.17	5.33	-
	4	3.01	27.30	5.37	4.17
การจัดกลุ่มตามรูปแบบอุปสงค์ (ค่ามัธยฐาน)		อุปสงค์รูปแบบเป็นก้อน	อุปสงค์รูปแบบราบเรียบ	อุปสงค์รูปแบบผันผวน	อุปสงค์รูปแบบไม่สม่ำเสมอ
		3.24	27.59	3.85	5.56

เฉลี่ยแบบสมมาตรต่ำที่สุดเท่ากับ 9.38% สะท้อนให้เห็นว่าวิธีการจัดกลุ่มดังกล่าวสามารถจำแนกข้อมูลอะไหล่สำรองที่มีลักษณะอุปสงค์ใกล้เคียงกันได้อย่างเหมาะสม และคัดแยกข้อมูลที่มีพฤติกรรมแตกต่างหรือเป็นจุดรวมออกจากกลุ่มหลักได้อย่างมีประสิทธิภาพส่งผลให้ตัวแบบสามารถเรียนรู้โครงสร้างเชิงเวลาของข้อมูลได้อย่างแม่นยำยิ่งขึ้น ผลนี้สอดคล้องกับที่ [4] และ [7] พบว่าตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ให้ผลการพยากรณ์ที่ดีเมื่อข้อมูลมีโครงสร้างอนุกรมเวลาที่ชัดเจน สำหรับความแตกต่างของค่าความคลาดเคลื่อนระหว่างกลุ่มที่ปรากฏอยู่นั้นเป็นผลมาจากลักษณะอุปสงค์ในบางกลุ่มที่มีความผันผวนสูงเกิดขึ้นอย่างไม่สม่ำเสมอ และมีค่าความต้องการจริงในบางช่วงเวลาดำมากหรือเข้าใกล้ศูนย์ ซึ่งส่งผลให้ค่าความคลาดเคลื่อนขยายขึ้นอย่างมีนัยสำคัญ

3.3 อภิปรายผลการทดลอง

การศึกษานี้ประเมินประสิทธิภาพของการจัดกลุ่มและการพยากรณ์ความต้องการอะไหล่สำรองผ่าน 3 ขั้นตอน ได้แก่ 1) การจัดกลุ่มด้วยเทคนิคเชิงคำนวณหลายรูปแบบ 2) การพยากรณ์ความต้องการด้วยวิธีพยากรณ์ 4 วิธี และ 3) การประเมินความแม่นยำด้วยค่า sMAPE เพื่อคัดเลือกชุดค่าผสมที่เหมาะสมที่สุดโดยมีผลลัพธ์ดังตารางที่ 13

3.3.1 การจัดกลุ่มอะไหล่สำรอง

ผลการจัดกลุ่มพบว่าการจัดกลุ่มตามรูปแบบอุปสงค์แบบเดิมไม่เหมาะสมกับข้อมูลชุดนี้ เนื่องจากเกณฑ์ตัดสินใจที่ค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนเท่ากับ 1.32 และ 0.49 ตามลำดับ ทำให้อะไหล่สำรองกว่าร้อยละ 99.6 ถูกจัดอยู่ในกลุ่มอุปสงค์รูปแบบเป็นก้อนเพียงกลุ่มเดียว จึงไม่



ตารางที่ 13 สรุปผลการจับคู่ชุดค่าผสมของรูปแบบการจัดกลุ่มอะไหล่สำรอง และตัวแบบพยากรณ์ที่ดีที่สุด

วิธีการพยากรณ์	วิธีการจับกลุ่ม	ค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรเฉลี่ย (%)
วิธีค่าเฉลี่ยเคลื่อนที่อย่างง่าย	DBSCAN (K = 4)	16.55
วิธีทำให้เรียบแบบเอกซ์โพเนนเชียลของโฮลท์	DBSCAN (K = 4)	97.64
วิธีของครอสตัน (TSB)	DBSCAN (K = 4)	60.52
ตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่	DBSCAN (K = 4)	9.38

สามารถจำแนกความแตกต่างเชิงพฤติกรรมอุปสงค์ได้อย่างมีประสิทธิภาพ ซึ่งสอดคล้องกับข้อค้นพบของ [3] ที่ระบุว่าเกณฑ์ดั้งเดิมถูกพัฒนาขึ้นบนสมมติฐานของข้อมูลที่มีการกระจายตัวในช่วงต่ำ จึงจำเป็นต้องปรับให้สอดคล้องกับลักษณะเฉพาะของข้อมูลในแต่ละบริบทอุตสาหกรรม แม้การจัดกลุ่มแบบเค-มีนส์ และแบบลำดับชั้นจะช่วยให้โครงสร้างกลุ่มชัดเจนขึ้น แต่ยังไม่ปรากฏการซ้อนทับของข้อมูลระหว่างกลุ่ม ส่งผลให้ความเป็นเนื้อเดียวกันภายในกลุ่มยังไม่เพียงพอสอดคล้องกับที่ [2] ชี้ว่าคุณภาพของการจัดกลุ่มมีผลโดยตรงต่อความแม่นยำของการพยากรณ์ เนื่องจากความหลากหลายภายในกลุ่มที่สูงจะขัดขวางความสามารถของแบบจำลองในการเรียนรู้รูปแบบอุปสงค์ที่แท้จริง ในทางกลับกันการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมกันให้ผลการจำแนกที่ดีที่สุด เนื่องจากสามารถแบ่งกลุ่มตามความ

หนาแน่นของข้อมูลและคัดแยกจุดรบกวนออกจากข้อมูลหลักได้อย่างชัดเจนทำให้แต่ละกลุ่มสะท้อนโครงสร้างอุปสงค์ได้ถูกต้องและเหมาะสมกว่าแนวทางอื่น ดังที่ [5] และ [7] ได้ยืนยันในบริบทของอะไหล่สำรองว่าการจำแนกลักษณะอุปสงค์ที่ถูกต้องแม่นยำเป็นเงื่อนไขพื้นฐานของการเลือกตัวแบบพยากรณ์ที่เหมาะสม และความเป็นเนื้อเดียวกันภายในกลุ่มที่สูงขึ้นจะช่วยลดการปะปนของข้อมูลต่างลักษณะ ทำให้แบบจำลองเรียนรู้รูปแบบอุปสงค์ได้ชัดเจนและนำไปสู่การค้นหาค่าผสมที่เหมาะสมที่สุดได้อย่างเป็นระบบ

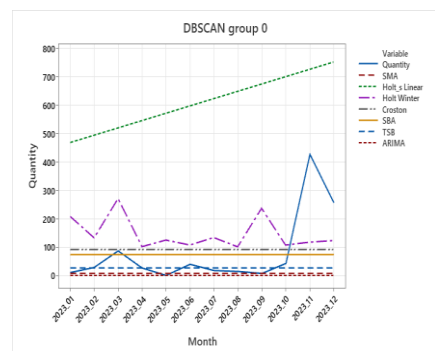
3.3.2 การพยากรณ์ความต้องการ

เมื่อพิจารณาผลการพยากรณ์หลังจัดกลุ่มพบความสัมพันธ์เชิงบวกที่ชัดเจนระหว่างคุณภาพของการจำแนกกลุ่มกับความแม่นยำของการพยากรณ์ โดยกลุ่มที่มีความเป็นเนื้อเดียวกันสูงจะลดความผันผวนภายในคลัสเตอร์และส่งผลกระทบต่อแบบจำลองพยากรณ์มีประสิทธิภาพสูงขึ้น ทั้งนี้ พบว่าในทุกวิธีพยากรณ์การจับคู่กับการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมแบบ 4 กลุ่ม ($k = 4$) ให้ผลดีกว่าการจัดกลุ่มรูปแบบอื่นอย่างสม่ำเสมอ ผลดังกล่าวสอดคล้องกับ [2] ที่ระบุว่าการจัดกลุ่มข้อมูลก่อนการพยากรณ์สามารถเพิ่มประสิทธิภาพของแบบจำลองได้อย่างมีนัยสำคัญ โดยงานวิจัยนี้ขยายข้อค้นพบดังกล่าวไปสู่บริบทของอะไหล่สำรองซึ่งมีความท้าทายสูงกว่า เนื่องจากลักษณะอุปสงค์ไม่ต่อเนื่องและความหนาแน่นของจุดรวมที่สูง ดังที่ [5] และ [6] ได้ชี้ให้เห็นว่าการพยากรณ์ในบริบทโซ่อุปทานอะไหล่สำรองต้องการแนวทางเฉพาะทางที่แตกต่างจากการพยากรณ์ข้อมูลอุปสงค์ทั่วไปอย่างมีนัยสำคัญ สำหรับตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ที่ให้ผลดีที่สุดในการศึกษานี้ สอดคล้องกับ [4]

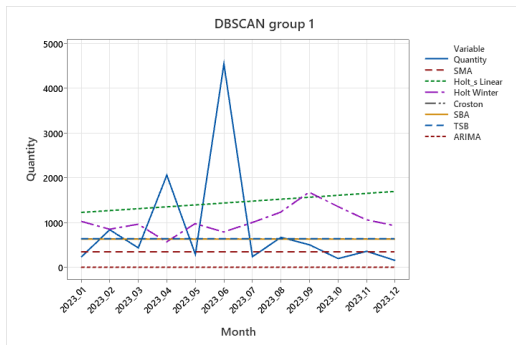
ที่พบว่าตัวแบบดังกล่าวมีประสิทธิภาพสูงในการจับโครงสร้างเชิงเวลาของข้อมูลที่มีความซับซ้อน โดยเฉพาะอย่างยิ่งเมื่อข้อมูลภายในกลุ่มมีความเป็นเนื้อเดียวกันในระดับที่เพียงพอ ขณะที่วิธีของโครสตัน ซึ่งพัฒนาขึ้นสำหรับอุปสงค์ไม่ต่อเนื่องโดยเฉพาะกลับให้ค่าความคลาดเคลื่อนสูงกว่าในหลายกลุ่ม เนื่องจากหลังการจัดกลุ่มด้วยวิธีการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรบกวน ข้อมูลในแต่ละกลุ่มมีโครงสร้างอนุกรมเวลาที่ชัดเจนเพียงพอให้ตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่จับรูปแบบได้อย่างมีประสิทธิภาพ นอกจากนี้ผลการศึกษายังขยายข้อค้นพบของ [7] โดยแสดงให้เห็นว่าการบูรณาการกระบวนการจัดกลุ่มตามโครงสร้างอุปสงค์เข้ากับตัวแบบพยากรณ์สามารถเพิ่มความแม่นยำได้อย่างมีนัยสำคัญเมื่อเทียบกับการใช้ตัวแบบพยากรณ์โดยปราศจากการแบ่งกลุ่มข้อมูลล่วงหน้า

3.3.3 การประเมินความแม่นยำด้วยค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร ผลการประเมินด้วยค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรยืนยันว่าการจัดกลุ่มที่ลดความแปรปรวนภายในคลัสเตอร์ได้อย่างมีประสิทธิภาพส่งผลโดยตรงต่อความแม่นยำของการพยากรณ์ โดยเฉพาะกับตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ซึ่งมีความไวสูงต่อจุดรบกวนในข้อมูล การคัดแยกจุดรบกวนออกจากคลัสเตอร์หลัก โดยวิธีการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรบกวนจึงทำให้ค่าความคลาดเคลื่อนต่ำกว่าวิธีอื่นอย่างมีนัยสำคัญทางสถิติ สอดคล้องกับ [2] ที่พบว่าการจัดกลุ่มข้อมูลก่อนการพยากรณ์ช่วยลดความแปรปรวนภายในกลุ่มและเพิ่มความแม่นยำของแบบจำลอง ขณะที่

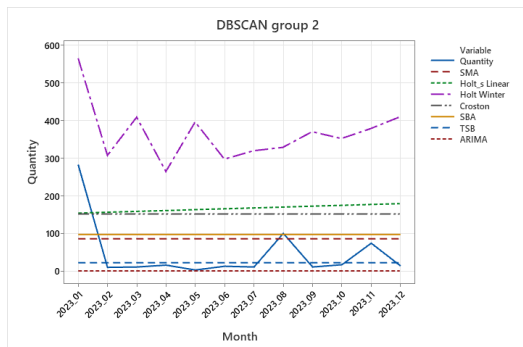
[5] และ [6] ชี้ให้เห็นว่าอุปสงค์ไม่ต่อเนื่องที่มีความผันผวนสูงยังคงเป็นความท้าทายหลัก ในการพยากรณ์อะไหล่สำรองซึ่งสะท้อนชัดเจนจากค่าความคลาดเคลื่อนของกลุ่มที่ 2 ที่สูงกว่ากลุ่มอื่น (27.03%) นอกจากนี้ความสม่ำเสมอของผลที่ดีขึ้นในทุกตัวแบบพยากรณ์ เมื่อใช้คลัสเตอร์จากการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรบกวนยังทำหน้าที่เป็นตัวตรวจสอบความเสถียรของผลการวิจัย และยืนยันว่าการเลือกตัวแบบพยากรณ์ควรพิจารณาควบคู่กับวิธีการจัดกลุ่มที่เหมาะสมกับโครงสร้างข้อมูล ดังที่ [4] และ [7] ได้เสนอในบริบทที่แตกต่างกัน จากรูปที่ 18-21 ซึ่งแสดงผลการพยากรณ์เปรียบเทียบกับความต้องการจริงใน พ.ศ. 2566 พบว่ากลุ่มที่ 1 กลุ่มที่ 2 กลุ่มที่ 3 และกลุ่มที่ 4 มีค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตรเท่ากับ 3.01% 27.03% 5.37% และ 4.17% ตามลำดับ ความแตกต่างดังกล่าวสะท้อนให้เห็นว่าลักษณะเฉพาะของอุปสงค์แต่ละกลุ่มมีผลต่อความสามารถในการพยากรณ์ ซึ่งยืนยันความสำคัญของการจัดกลุ่มข้อมูลตามโครงสร้างอุปสงค์ก่อนการพยากรณ์ตามแนวทางที่ [2] ได้นำเสนอไว้



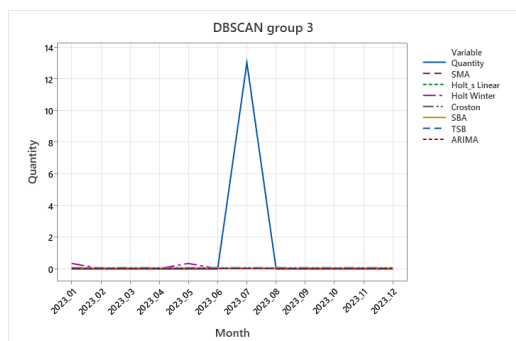
รูปที่ 18 ผลการพยากรณ์ด้วยตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ในกลุ่มที่ 0



รูปที่ 19 ผลการพยากรณ์ด้วยตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ในกลุ่มที่ 1



รูปที่ 20 ผลการพยากรณ์ด้วยตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ในกลุ่มที่ 2



รูปที่ 21 ผลการพยากรณ์ด้วยตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ในกลุ่มที่ 3

4. สรุปผลการทดลอง

ผลการศึกษายืนยันว่าการจัดกลุ่มอะไหล่สำรองโดยอาศัยค่าคาบเวลาเฉลี่ยระหว่างอุปสงค์และค่ากำลังสองสัมประสิทธิ์ความแปรปรวนเป็นปัจจัยกำหนดประสิทธิภาพของการพยากรณ์ความต้องการ เนื่องจากคุณภาพของการจำแนกโครงสร้างอุปสงค์ในขั้นก่อนการพยากรณ์มีผลโดยตรงต่อความสามารถของแบบจำลองในการเรียนรู้พฤติกรรมของข้อมูล ผลการทดลองแสดงให้เห็นว่าการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมควบรวมกับตัวแบบถดถอยอัตโนมัติแบบอินทิเกรตและค่าเฉลี่ยเคลื่อนที่ให้ผลการพยากรณ์ที่ดีที่สุดเมื่อประเมินด้วยค่าความคลาดเคลื่อนเปอร์เซ็นต์สัมบูรณ์เฉลี่ยแบบสมมาตร สะท้อนให้เห็นว่าการลดความไม่แน่นอนเดียวกันของข้อมูลภายในกลุ่มก่อนเข้าสู่กระบวนการพยากรณ์ช่วยเพิ่มขีดความสามารถของแบบจำลองในการจับโครงสร้างอุปสงค์ได้อย่างแม่นยำยิ่งขึ้นข้อค้นพบดังกล่าวชี้ให้เห็นว่าความแม่นยำของการพยากรณ์มีได้ขึ้นอยู่กับสมรรถนะของตัวแบบพยากรณ์แต่เพียงลำพัง หากแต่ยังขึ้นอยู่กับความเหมาะสมของกระบวนการเตรียมข้อมูลและการจัดกลุ่มในขั้นต้นด้วย กล่าวคือวิธีการจัดกลุ่มทำหน้าที่เป็นกลไกสำคัญในการลดการปะปนของข้อมูลที่มีพฤติกรรมอุปสงค์แตกต่างกัน ซึ่งเป็นต้นเหตุประการหนึ่งของความคลาดเคลื่อนในการพยากรณ์ โดยเฉพาะในบริบทของอะไหล่สำรองที่มีลักษณะอุปสงค์ไม่ต่อเนื่อง ผันผวนสูง และมีข้อมูลรบกวนในสัดส่วนมาก

เมื่อเปรียบเทียบกับงานวิจัยที่ผ่านมา งานวิจัยนี้มีความโดดเด่นในเชิงกรอบการวิเคราะห์ กล่าวคือ งานวิจัยก่อนหน้านี้ส่วนหนึ่งมุ่งประยุกต์การจัดกลุ่มในบริบทที่ไม่ใช่อะไหล่สำรอง ส่วนหนึ่งใช้ตัวชี้วัดลักษณะอุปสงค์

เพื่อสนับสนุนการกำหนดนโยบายสินค้าคงคลัง และอีกส่วนหนึ่งมุ่งเปรียบเทียบหรือพัฒนาตัวแบบพยากรณ์เป็นหลัก แม้งานที่เกี่ยวข้องกับอะไหล่สำรองโดยตรงจะชี้ให้เห็นความซับซ้อนของอุปสงค์ประเภทไม่ต่อเนื่อง ผันผวน และเป็นก้อน แต่ส่วนใหญ่ยังพิจารณาการจัดกลุ่มและการพยากรณ์แยกออกจากกัน งานวิจัยนี้จึงขยายขอบเขตดังกล่าวโดยออกแบบกรอบการทดลองที่ประเมินผลร่วมกันระหว่างหลายวิธีการจัดกลุ่มและหลายตัวแบบพยากรณ์ภายใต้ข้อมูลจริง ชุดข้อมูลเดียวกัน และเกณฑ์ประเมินเดียวกัน ทำให้สามารถวิเคราะห์ความสัมพันธ์ระหว่างวิธีการจัดกลุ่มอุปสงค์กับความแม่นยำของการพยากรณ์ได้อย่างเป็นระบบและมีความเชื่อถือได้สูง

ในเชิงคุณูปการทางวิชาการงานวิจัยนี้ขยายองค์ความรู้ด้านการบูรณาการระหว่างการจัดกลุ่มข้อมูลและการพยากรณ์สำหรับระบบอะไหล่สำรองที่มีลักษณะอุปสงค์ไม่สม่ำเสมอ โดยแสดงให้เห็นว่าการจัดกลุ่มเชิงพื้นที่โดยอาศัยความหนาแน่นและจุดรวมสามารถลดผลกระทบของข้อมูลรบกวนและจำแนกโครงสร้างอุปสงค์เชิงพฤติกรรมได้เหนือกว่าวิธีการจัดกลุ่มทางเลือกอื่น ส่งผลให้ตัวแบบพยากรณ์ในขั้นถัดไปมีประสิทธิภาพสูงขึ้น ทั้งนี้ นัยสำคัญของผลการศึกษาได้จำกัดเพียงการระบุวิธีที่ให้ผลดีที่สุด หากแต่ยังเน้นย้ำว่าการเลือกตัวแบบพยากรณ์ต้องพิจารณาควบคู่กับโครงสร้างข้อมูลและวิธีการจัดกลุ่มที่เหมาะสม งานวิจัยนี้จึงมีคุณค่าในฐานะการศึกษาเชิงเปรียบเทียบที่เชื่อมโยงทั้งสองกระบวนการเข้าด้วยกันอย่างบูรณาการ ช่วยเติมเต็มช่องว่างทางวิชาการเกี่ยวกับความสัมพันธ์ระหว่างการจำแนกโครงสร้างอุปสงค์และความแม่นยำของการพยากรณ์ในบริบทของอะไหล่สำรอง และสามารถนำไปใช้เป็นฐาน

ความรู้สำหรับการพัฒนาแนวทางการบริหารอะไหล่สำรองและการตัดสินใจด้านสินค้าคงคลังได้อย่างมีประสิทธิภาพในเชิงปฏิบัติ

เอกสารอ้างอิง

- [1] Krungsri Research. (2021). Industry Outlook 2021–2023,” Krungsri.com. [Online] Available: <https://www.krungsri.com/th/research/industry/summary-outlook/industry-summary-outlook-2021-2023> (in Thai).
- [2] M. Seyedan, F. Mafakheri, and C. Wang, “Cluster-based demand forecasting using Bayesian model averaging: An ensemble learning approach,” *Decision Analytics Journal*, vol. 3, 2022, Art. no. 100033, doi: 10.1016/j.dajour.2022.100033.
- [3] T. Chunpong, “Inventory management for integrated circuit packaging materials under variable demand,” Master’s thesis, Department of Industrial Engineering Faculty of Engineering Chulalongkorn University, Bangkok, Thailand, 2021 (in Thai).
- [4] S. Wongkamphu, “Battery sales forecasting using traditional forecasting and machine learning techniques,” Master’s thesis, Faculty of Engineering Chulalongkorn University, Bangkok, Thailand, 2023 (in Thai).
- [5] N. Paranthaman, H. N. Perera, N. T. Thalagala, and D. Kosgoda, “Evaluating intermittent demand forecasting techniques



- for spare part supply chains,” *IFAC-PapersOnLine*, vol. 58, no. 14, pp. 136–141, 2024, doi: 10.1016/j.ifacol.2024.09.025.
- [6] T. B. Affonso, S. V. Conceição, L. R. Muniz, J. F. D. F. Almeida, and J. C. Lima, “A new hybrid forecasting method for spare part inventory management using heuristics and bootstrapping,” *Decision Analytics Journal*, vol. 10, 2024, Art. no. 100415, doi: 10.1016/j.dajour.2024.100415.
- [7] M. İfraz, A. Aktepe, S. Ersöz, and T. Çetinyokuş, “Demand forecasting of spare parts with regression and machine learning methods: Application in a bus fleet,” *Journal of Engineering Research*, vol. 11, no. 2, 2023, Art. no. 100057, doi: 10.1016/j.jer.2023.100057.
- [8] A. A. Syntetos, J. E. Boylan, and J. D. Croston, “On the categorization of demand patterns,” *Journal of the Operational Research Society*, vol. 56, no. 5, pp. 495–503, 2005, doi: 10.1057/palgrave.jors.2601841.
- [9] B. A. Kannan, G. Kodi, O. Padilla, D. Gray, and B. C. Smith, “Forecasting spare parts sporadic demand using traditional methods and machine learning – a comparative study,” *SMU Data Science Review*, vol. 3, no. 2, 2020, Art. no. 9.
- [10] D. Visbal-Cadavid, E. Delahoz-Domínguez, and A. Mendoza-Mendoza, “A multiple factor analysis and hierarchical clustering of global logistics governance and development,” *Decision Analytics Journal*, vol. 15, 2025, Art. no. 100579, doi: 10.1016/j.dajour.2025.100579.
- [11] R. Ma, J. Sha, S. Zhang, D. Zhu, W. Kang, and J. Liu, “Fast grouping fusion method of dual carbon monitoring data based on DBSCAN clustering algorithm,” *Results in Engineering*, vol. 26, 2025, Art. no. 105057, doi: 10.1016/j.rineng.2025.105057.
- [12] A. Oyemaja, “Moving averages in time series analysis: Understanding trends & forecasting,” *ResearchGate Preprint*, Mar. 2024, doi: 10.13140/RG.2.2.14811.17441.
- [13] W. Dechadumrongchai, “Improvement of demand forecasting and replenishment policy for agricultural machine spare parts,” Master’s thesis, Faculty of Engineering Chulalongkorn University, Bangkok, Thailand, 2021 (in Thai).
- [14] M. Doszyn, “Accuracy of intermittent demand forecasting systems in the enterprise,” *European Research Studies Journal*, vol. 23, no. 4, pp. 912–930, 2020, doi: 10.35808/ersj/1723.
- [15] T. Wongchitsukkasem, “Thailand steel scrap price forecasting using decomposition-based hybrid model,” Master’s thesis, Department of Industrial Engineering,



- Faculty of Engineering, Chulalongkorn University, Bangkok, Thailand, 2023 (in Thai).
- [16] S. Makridakis, E. Spiliotis, and V. Assimakopoulos, "The M4 competition: 100,000 time series and 61 forecasting methods," *International Journal of Forecasting*, vol. 36, no. 1, pp. 54–74, 2020, doi: 10.1016/j.ijforecast.2019.04