



การเปรียบเทียบประสิทธิภาพของโมเดลพยากรณ์ปริมาณฝุ่น PM2.5 ในกรุงเทพมหานคร ระหว่าง SARIMA และ XGBoost

Comparison of PM2.5 Forecasting Model Performance in Bangkok: SARIMA vs. XGBoost

ศรัณธร มั่งมี^{1*} และ พงศ์พัฒน์ ฉายศิริพันธ์²

Saranthon Maungmee^{1*} and Pongpat Chaisiripan²

¹ อาจารย์, สาขาวิชาธุรกิจดิจิทัลและเทคโนโลยี, คณะเทคโนโลยีสารสนเทศ, มหาวิทยาลัยสยาม

¹ Lecturer in Digital Business and Technology, Faculty of Information Technology, Siam University

² อาจารย์สาขาวิชาเทคโนโลยีสารสนเทศ, คณะเทคโนโลยีสารสนเทศ, มหาวิทยาลัยสยาม

² Lecturer in Information Technology, Faculty of Information Technology, Siam University

*Corresponding author, E-mail: saranthon.mau@siam.edu

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อเปรียบเทียบประสิทธิภาพของแบบจำลอง SARIMA และ XGBoost ในการพยากรณ์ค่าฝุ่นละอองขนาดเล็กไม่เกิน 2.5 ไมครอน (PM2.5) ในเขตกรุงเทพมหานคร โดยใช้ข้อมูลจาก สถานีตรวจวัดของกรมควบคุมมลพิษ จำนวน 8 สถานี ครอบคลุมช่วงปี พ.ศ. 2561 ถึง 2566 ข้อมูลได้รับการวิเคราะห์เบื้องต้นและผ่านกระบวนการเตรียมข้อมูลก่อนนำเข้าสู่การสร้างแบบจำลอง จากนั้นทำการ ประเมินผลการพยากรณ์ด้วยตัวชี้วัด ได้แก่ ค่าเฉลี่ยสัมบูรณ์ของความคลาดเคลื่อน (MAE) ค่ารากของ ความคลาดเคลื่อนกำลังสองเฉลี่ย (RMSE) และค่าสัมประสิทธิ์การตัดสินใจ (R^2) ผลการวิจัยพบว่า XGBoost ให้ผลการพยากรณ์ที่แม่นยำกว่า SARIMA ในทุกตัวชี้วัด โดยเฉพาะในช่วงฤดูหนาวที่ค่าฝุ่น มีความแปรปรวนสูง สะท้อนถึงศักยภาพของเทคนิค Machine Learning ในการวิเคราะห์ข้อมูลที่มีความ ซับซ้อนผลการศึกษานี้ สามารถนำไปใช้เป็นแนวทางในการเลือกแบบจำลอง เพื่อสนับสนุนการบริหาร จัดการคุณภาพอากาศเชิงคาดการณ์ในระดับพื้นที่

คำสำคัญ: ค่าฝุ่น PM2.5, การพยากรณ์, แบบจำลอง SARIMA, แบบจำลอง XGBoost,
การเรียนรู้ของเครื่อง, คุณภาพอากาศ

Abstract

This study aims to compare the forecasting performance of the SARIMA and XGBoost models for predicting fine particulate matter (PM2.5) levels in Bangkok, Thailand. The data used were obtained from eight air quality monitoring stations operated by the



Pollution Control Department, covering the period from 2018 to 2023. The data were preprocessed and analyzed through exploratory data analysis before being used to construct forecasting models. Model performance was evaluated using Mean Absolute Error (MAE), Root Mean Square Error (RMSE), and the coefficient of determination (R^2). The results indicate that the XGBoost model outperforms SARIMA in all evaluation metrics, particularly during the winter months when PM_{2.5} concentrations show greater variability. These findings highlight the potential of machine learning techniques in handling complex time series data and suggest their applicability in supporting air quality management and early warning systems.

Keywords: PM_{2.5}, Forecasting, SARIMA model, XGBoost model, Machine Learning, Air Quality

บทนำ

ปัญหาฝุ่นละอองขนาดเล็ก PM_{2.5} เป็นปัญหาสิ่งแวดล้อมที่ส่งผลกระทบต่อสุขภาพของประชาชนในหลายประเทศทั่วโลก รวมถึงประเทศไทย โดยเฉพาะในเขตเมืองใหญ่ เช่น กรุงเทพมหานคร ซึ่งมีแหล่งกำเนิดฝุ่นจากการจราจร อุตสาหกรรม และกิจกรรมของมนุษย์เป็นจำนวนมาก (กรมควบคุมมลพิษ, 2566) ฝุ่น PM_{2.5} เป็นอนุภาคขนาดเล็กกว่า 2.5 ไมครอนเมตร ซึ่งสามารถเข้าสู่ระบบทางเดินหายใจและกระแสเลือด ส่งผลกระทบต่อระบบหัวใจและหลอดเลือด รวมถึงเพิ่มความเสี่ยงต่อโรคระบบทางเดินหายใจและมะเร็งปอด (World Health Organization [WHO], 2021) การพยากรณ์ค่าฝุ่น PM_{2.5} มีความสำคัญอย่างยิ่งต่อการวางแผนและจัดการคุณภาพอากาศของหน่วยงานภาครัฐ ตลอดจนการเตรียมความพร้อมของประชาชนเพื่อลดความเสี่ยงด้านสุขภาพ การพยากรณ์ที่แม่นยำช่วยให้สามารถแจ้งเตือนล่วงหน้า และนำไปสู่การกำหนดมาตรการป้องกันที่มีประสิทธิภาพ เช่น การจำกัดการใช้น้ำมันในบางพื้นที่ หรือการแนะนำให้ประชาชนหลีกเลี่ยงกิจกรรมกลางแจ้ง

จากการทบทวนวรรณกรรม พบว่างานวิจัยเกี่ยวกับการพยากรณ์ค่าฝุ่น PM_{2.5} ได้มีการใช้เทคนิคหลากหลายรูปแบบ โดยสามารถแบ่งออกเป็น 2 กลุ่มหลัก ได้แก่ แบบจำลองเชิงสถิติและแบบจำลองเชิง Machine Learning (สมชาย วิริยะกุล, สุนิสา บุญประเสริฐ และศิริพร สุทธิเจริญ, 2562) ได้ศึกษาการใช้แบบจำลอง SARIMA ในการพยากรณ์ค่าฝุ่น PM_{2.5} ในเขตกรุงเทพมหานคร โดยใช้ข้อมูลรายวันย้อนหลัง 5 ปี (พ.ศ. 2557 - 2561) และทำการวิเคราะห์แนวโน้มและฤดูกาลของข้อมูล ผลการทดลองแสดงให้เห็นว่า SARIMA สามารถจับรูปแบบของข้อมูลตามฤดูกาลได้ดี และให้ค่าความคลาดเคลื่อนเฉลี่ยรากที่สอง (RMSE) อยู่ในระดับที่ยอมรับได้ อย่างไรก็ตาม แบบจำลองนี้มีข้อจำกัดในการรองรับปัจจัยแวดล้อมที่ซับซ้อน เช่น ความเร็วลม อุณหภูมิ และปริมาณการจราจร (อภิชัย นาคะเวช, 2563) ได้ศึกษาแบบจำลองการพยากรณ์ค่าฝุ่น PM_{2.5} โดยใช้ XGBoost ซึ่งเป็นเทคนิคการเรียนรู้ของเครื่องที่สามารถเรียนรู้รูปแบบ



ที่ซับซ้อนของข้อมูลได้ดี งานวิจัยนี้ใช้ข้อมูลค่าฝุ่นย้อนหลัง 3 ปี (พ.ศ. 2560 - 2562) และนำปัจจัยแวดล้อม เช่น ความชื้นสัมพัทธ์ อุณหภูมิ และปริมาณจากรวมเป็นตัวแปรนำเข้า ผลการทดลองพบว่า XGBoost ให้ค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (MAE) ต่ำกว่า SARIMA และสามารถคาดการณ์ค่าฝุ่นได้แม่นยำกว่า โดยเฉพาะในช่วงที่ค่าฝุ่นเปลี่ยนแปลงอย่างรวดเร็ว (ปณิธาน จันทศรี และ จารุวรรณ ทองก้อน, 2564) ได้พัฒนาแบบจำลองพยากรณ์ค่าฝุ่น PM2.5 โดยใช้การเปรียบเทียบเทคนิค Machine Learning 3 วิธี ได้แก่ 1) Random Forest Regression 2) Support Vector Machine for Regression (SVR) และ 3) Long Short-Term Memory (LSTM) โดยใช้ข้อมูลย้อนหลัง 7 ปี (พ.ศ. 2557 - 2563) และแบ่งข้อมูลออกเป็นชุดฝึกสอนและชุดทดสอบ ผลการทดลองพบว่า LSTM ซึ่งเป็นโครงข่ายประสาทเทียมเชิงลึก สามารถพยากรณ์ค่าฝุ่นได้แม่นยำที่สุด เมื่อพิจารณาจากค่า RMSE และ Mean Absolute Percentage Error (MAPE) แม้ว่า SARIMA และ XGBoost จะเป็นแบบจำลองที่มีประสิทธิภาพในการพยากรณ์ค่าฝุ่น PM2.5 แต่ยังไม่มีการศึกษาโดยตรงที่ทำการเปรียบเทียบโมเดลทั้งสอง โดยใช้ตัวชี้วัดทางสถิติที่หลากหลาย โดยเฉพาะ Mean Bias Error (MBE) ซึ่งใช้วัดแนวโน้มของโมเดลว่ามีการพยากรณ์ค่าต่ำกว่าหรือสูงกว่าค่าจริง และ Symmetric Mean Absolute Percentage Error (SMAPE) ซึ่งเป็นตัวชี้วัดข้อผิดพลาดเชิงเปอร์เซ็นต์ที่แก้ปัญหาการใช้ค่า MAPE ในกรณีที่ค่าจริงเป็นศูนย์ (Willmott & Matsuura, 2005) ดังนั้นงานวิจัยนี้มุ่งเน้นการเปรียบเทียบประสิทธิภาพของแบบจำลอง SARIMA และ XGBoost ในการพยากรณ์ค่าฝุ่น PM2.5 โดยพิจารณาตัวชี้วัดทางสถิติที่สำคัญพร้อมทั้งวิเคราะห์ข้อจำกัดของแต่ละโมเดล เพื่อนำไปสู่แนวทางในการปรับปรุงการพยากรณ์ในอนาคต

ในการศึกษานี้ ได้เลือกใช้แบบจำลองสองเทคนิคหลัก ได้แก่ Seasonal Autoregressive Integrated Moving Average (SARIMA) และ Extreme Gradient Boosting (XGBoost) เพื่อเปรียบเทียบประสิทธิภาพในการพยากรณ์ค่าฝุ่น PM2.5 ทั้งนี้เนื่องจาก SARIMA เป็นเทคนิคเชิงสถิติที่มีความเหมาะสมในการวิเคราะห์ข้อมูลอนุกรมเวลาที่มีลักษณะความต่อเนื่อง มีแนวโน้ม (Trend) และฤดูกาล (Seasonality) ซึ่งเป็นลักษณะเฉพาะของข้อมูล PM2.5 รายวัน ขณะที่ XGBoost เป็นเทคนิค Machine Learning ที่มีประสิทธิภาพสูงในการจัดการข้อมูลที่มีความซับซ้อน ไม่เป็นเชิงเส้น และสามารถลดความคลาดเคลื่อนจากความสัมพันธ์ที่ไม่ปกติระหว่างตัวแปรได้ การเลือกใช้ทั้งสองเทคนิคนี้ช่วยให้สามารถประเมินความสามารถของแบบจำลองที่มีพื้นฐานต่างกันในการพยากรณ์ค่าฝุ่นได้อย่างรอบด้านและเป็นกลาง การเลือกแบบจำลอง SARIMA และ XGBoost เพื่อเปรียบเทียบในการศึกษานี้มีที่มาจากลักษณะเฉพาะของข้อมูล PM2.5 และแนวโน้มของการพัฒนารูปแบบการพยากรณ์ในปัจจุบัน โดย SARIMA เป็นแบบจำลองเชิงสถิติที่ได้รับการยอมรับอย่างแพร่หลาย ในการวิเคราะห์ข้อมูลอนุกรมเวลาที่มีแนวโน้มและฤดูกาลที่ชัดเจน อีกทั้งยังมีความสามารถในการอธิบายโครงสร้างของข้อมูลตามช่วงเวลาได้ดี ส่วน XGBoost ซึ่งเป็นเทคนิค Machine Learning ประเภท Ensemble Learning ได้รับความนิยมเพิ่มขึ้นอย่างรวดเร็วในช่วงไม่กี่ปีที่ผ่านมา เนื่องจากมีความสามารถในการจัดการกับข้อมูลที่มีความซับซ้อนสูง และสามารถลดข้อผิดพลาดจากการพยากรณ์ได้อย่างมีประสิทธิภาพ การนำสองแนวทางที่มีพื้นฐานแตกต่างกันนี้มาเปรียบเทียบ จึงมีวัตถุประสงค์เพื่อศึกษาว่าแบบจำลองเชิงสถิติดั้งเดิมหรือเทคนิค Machine Learning สมัยใหม่ จะมีความเหมาะสมมากกว่ากันสำหรับการพยากรณ์ปริมาณฝุ่น PM2.5 ในบริบทของกรุงเทพมหานคร



วัตถุประสงค์ของการวิจัย

1. เพื่อสร้างแบบจำลองพยากรณ์ค่าฝุ่น PM2.5 ในเขตกรุงเทพมหานคร ด้วยโมเดล SARIMA และ XGBoost
2. เพื่อเปรียบเทียบประสิทธิภาพของโมเดล SARIMA และ XGBoost โดยใช้ตัวชี้วัด MAE, RMSE, R^2 , MBE และ SMAPE

สมมติฐาน

การพยากรณ์ค่าฝุ่น PM2.5 ในพื้นที่กรุงเทพมหานครด้วยเทคนิคเหมืองข้อมูล (Time Series Data Mining Techniques) โดยใช้แบบจำลอง SARIMA, XGBoost และ LSTM บนชุดข้อมูลเดียวกัน สามารถคาดการณ์แนวโน้มค่าฝุ่น PM2.5 ได้อย่างแม่นยำและมีประสิทธิภาพในระดับที่ใกล้เคียงกัน

ขอบเขตของการวิจัย

1. ข้อมูลที่นำมาวิเคราะห์ เป็นข้อมูลค่าฝุ่น PM2.5 ในกรุงเทพมหานคร จากแหล่งข้อมูลของหน่วยงานราชการ เช่น กรมควบคุมมลพิษ (Pollution Control Department, PCD)
2. วิธีแบบจำลองพยากรณ์ค่าฝุ่น PM2.5 ในเขตกรุงเทพมหานคร ด้วยโมเดล SARIMA และ XGBoost เปรียบเทียบประสิทธิภาพของ SARIMA และ XGBoost ในการพยากรณ์ค่าฝุ่น PM2.5 ใช้ตัวชี้วัด เช่น Mean Absolute Error (MAE), Root Mean Squared Error (RMSE) และ R-squared (R^2)

ทั้งนี้ ในอนาคต ควรขยายขอบเขตของการวิจัยโดยการเปรียบเทียบและตรวจสอบความแม่นยำ (Validation) กับโมเดลอื่น ๆ เพิ่มเติม เช่น Random Forest, Support Vector Regression (SVR) หรือ Deep Learning Models อย่าง Long Short-Term Memory (LSTM) เพื่อประเมินประสิทธิภาพของเทคนิคที่หลากหลายยิ่งขึ้น นอกจากนี้ ควรพิจารณาใช้ฐานข้อมูลจากแหล่งอื่น ๆ นอกเหนือจากกรมควบคุมมลพิษ เช่น ข้อมูลจากเครือข่ายเซ็นเซอร์อากาศภาคประชาชน หรือข้อมูลสาธารณะจากหน่วยงานระหว่างประเทศ เพื่อเพิ่มความหลากหลายและความครอบคลุมของชุดข้อมูล รวมทั้งการรวบรวมข้อมูลจากพื้นที่อื่น ๆ นอกกรุงเทพมหานคร เช่น จังหวัดในภาคเหนือหรือภาคตะวันออกเฉียงเหนือที่ประสบปัญหาฝุ่น PM2.5 รุนแรง เพื่อให้ได้โมเดลที่มีความสามารถในการพยากรณ์ที่กว้างขวาง และนำไปประยุกต์ใช้ได้ในพื้นที่ที่แตกต่างกัน

แนวคิด ทฤษฎี กรอบแนวคิด

1. ข้อมูลอนุกรมเวลา (Time Series Data) คือชุดข้อมูลหรือค่าสังเกตเชิงปริมาณ (X) ที่มีการเก็บรวบรวม โดยเรียงลำดับของเวลาอย่างต่อเนื่องกัน การจัดเรียงข้อมูลจะมีลักษณะเป็นรายปี (Yearly) รายไตรมาส (Quarterly) รายเดือน (Monthly) รายอาทิตย์ (Weekly) รายวัน (Daily) หรือรายชั่วโมง (Hourly) ขึ้นอยู่กับความเหมาะสมในการนำไปใช้งาน (เฉลิมพล จตุพร, 2562) ส่วนประกอบของอนุกรมเวลาประกอบด้วย 4 ส่วน คือ



1. แนวโน้ม (Trend) คือ ส่วนที่ทำให้อนุกรมเวลามีค่าเพิ่มหรือลดลงเรื่อย ๆ เมื่อเวลาผ่านไป (ภูมิฐาน รังकुณวัฒน์, 2562)

2. วัฏจักร (Cycle) คือ ส่วนที่ทำให้อนุกรมเวลาที่เก็บข้อมูลมีค่าเพิ่มขึ้น และลดลงสลับกันไปรอบ ๆ ค่าแนวโน้ม การนับระยะเวลาของส่วนวัฏจักรจะนับจุดสูงสุดหนึ่ง (Peak) ไปยังอีกจุดสูงสุดหนึ่งหรือจากจุดต่ำสุดหนึ่ง (Trough) ไปยังอีกจุดต่ำสุดหนึ่ง เมื่อส่วนของวัฏจักรอยู่ในช่วงที่อนุกรมเวลามีค่าลดลงเรียกว่าช่วงถดถอย (Recession) ส่วนของวัฏจักรที่ทำให้อนุกรมเวลามีค่าเพิ่มขึ้นเรียกว่าช่วงฟื้นตัว (Recovery)

3. ความผันแปรจากฤดูกาล (Seasonal Variations) คือรูปแบบในช่วงเวลาหนึ่งของอนุกรมเวลาที่จะเป็นภายใน 1 ปีและเกิดขึ้นซ้ำกันทุกปี

4. ความผันผวนจากเหตุการณ์ไม่ปกติ (Irregular Fluctuations) คือ ส่วนที่ทำให้อนุกรมเวลามีค่าที่ผิดปกติไปจากรูปแบบที่เป็นมักเกิดจากเหตุการณ์ไม่คาดฝัน (Shock) คำนวณจากการนำค่าของส่วนแนวโน้ม ค่าของส่วนวัฏจักร และค่าของความผันแปรจากฤดูกาล ไปหักล้างค่าอนุกรมเวลานั้น

2. การวิเคราะห์การถดถอย (Regression analysis) เป็นเทคนิคทางสถิติที่ใช้ในการประมาณค่าความสัมพันธ์ระหว่างตัวแปรแบ่งออกเป็น 2 แบบคือ 1) การวิเคราะห์การถดถอยเชิงเส้น (Linear Regression Analysis) เป็นการวิเคราะห์การถดถอยที่ตัวแปรอิสระส่วนใหญ่เป็นตัวแปรเชิงปริมาณ ส่วนตัวแปรตามเป็นตัวแปรเชิงปริมาณเท่านั้น รูปแบบของความสัมพันธ์สามารถแทนได้ด้วยสมการทางคณิตศาสตร์ที่เป็นเชิงเส้น (Linear Model) 2) การวิเคราะห์การถดถอยแบบไม่เป็นเชิงเส้น (Non Linear Regression) เป็นการวิเคราะห์การถดถอยที่รูปแบบของความสัมพันธ์ระหว่างตัวแปรอิสระและตัวแปรตาม สามารถแทนได้ด้วยสมการทางคณิตศาสตร์ที่ไม่เป็นเชิงเส้น (Non – Linear Model) การวิเคราะห์การถดถอยแบ่งออกเป็น

1. แบบจำลองการถดถอยอย่างง่าย ประกอบด้วยตัวแปรตาม (Y) หนึ่งตัวแปร และตัวแปรอิสระ (X) เพียงหนึ่งตัวแปร ดังสมการที่ (1) (เฉลิมพล จตุพร, 2562)

$$Y_i = \alpha_0 + \beta_1 X_i + \epsilon_i \quad (1)$$

2. แบบจำลองการถดถอยพหุคูณ ประกอบด้วยตัวแปรตาม (Y) หนึ่งตัวแปร และตัวแปรอิสระ (X) มากกว่าหนึ่งตัวแปร ดังสมการที่ (2)

$$Y_t = \alpha_0 + \beta_1 X_{1t} + \alpha_0 + \beta_2 X_{2t} + \dots + \alpha_0 + \beta_n X_{nt} \epsilon_t$$

$$Y_t = \alpha_0 + \sum_{i=1}^n \beta_i X_{it} + \epsilon_t \quad (2)$$

3. SARIMA (Seasonal AutoRegressive Integrated Moving Average) เป็นการขยายแบบจำลอง ARIMA ให้รองรับข้อมูลที่มีองค์ประกอบฤดูกาล (Seasonality) โดย SARIMA(p,d,q)(P,D,Q)s (Box, Jenkins, & Reinsel, 2015) ดังสมการที่ (3)

$$\begin{aligned} & (1 - \varphi_1 B - \varphi_2 B^2 - \dots - \varphi_p B^p)(1 - \Phi_1 B^s \\ & \quad - \Phi_2 B^{2s} - \dots - \Phi_p B^{ps})(1 - B)^d(1 - B^s)^D Y_t \\ & = (1 - \theta_1 B - \theta - \theta_2 B^2 - \dots - \theta_q B^q)(1 - \Theta_1 B^s \\ & \quad - \Theta_2 B^{2s} - \dots - \Theta_Q B^{Qs})_{\varepsilon_t} \end{aligned} \quad (3)$$

4. XGBoost (eXtreme Gradient Boosting) เป็นอัลกอริทึมในการเรียนรู้ของเครื่องที่ใช้ Gradient Boosting โดยใช้แบบจำลอง Boosting ที่เป็น Ensemble Learning Algorithm ที่รวมแบบจำลองหลาย ๆ แบบจำลองเข้าด้วยกัน เพื่อสร้างแบบจำลองที่มีความแม่นยำสูงกว่าแบบจำลองเดี่ยว ซึ่ง XGBoost จะทำการเรียนรู้จากข้อมูลที่มีอยู่แล้ว และทำการ Optimize Hyperparameter เพื่อให้แบบจำลองทำนายได้อย่างแม่นยำ (ปัญหา ปะสีละเตสัง, 2564) โดยการเพิ่มต้นไม้ตัดสินใจ (Decision Trees) ทีละต้นเพื่อปรับปรุงผลลัพธ์ที่มีข้อผิดพลาดจากต้นไม้ก่อนหน้านี้ XGBoost ใช้การปรับแต่งหลายขั้นตอนและสูตรที่ช่วยเพิ่มประสิทธิภาพในการเรียนรู้ ซึ่งทำให้เป็นหนึ่งในโมเดลที่มีประสิทธิภาพสูงสำหรับการพยากรณ์ข้อมูล (Chen & Guestrin, 2016)

1. การรวมผลลัพธ์จากต้นไม้หลายต้น: สมมติว่า $f(x)$ คือฟังก์ชันของต้นไม้ตัดสินใจแต่ละต้นที่ใช้ในการพยากรณ์ ค่าพยากรณ์ของต้นไม้ที่ t -th ดังสมการที่ (4)

$$\hat{y}_t(x) = \sum_{i=1}^t f_i(x) \quad (4)$$

2. การคำนวณ Gradient Descent: ใน XGBoost, การคำนวณการอัปเดต (update) ของต้นไม้แต่ละต้นทำโดยใช้ Gradient Descent และ Second-order Approximation สูตรดัง ดังสมการที่ (5)

$$\mathcal{L}(t) = \sum_{i=1}^n \left[\text{loss}(y_i \hat{y}_i) + \frac{1}{2} \lambda \|\theta\|^2 \right] \quad (5)$$

3. การคำนวณ Gradient และ Hessian: การใช้ Gradient และ Hessian ช่วยในการหาค่าปรับปรุง (update) สำหรับการสร้างต้นไม้ใหม่:

- Gradient: เป็นอัตราการเปลี่ยนแปลงของ Loss Function:

$$g_i = \frac{\partial \mathcal{L}}{\partial \hat{y}_i} \quad (6)$$

- Hessian: คืออนุพันธ์อันดับสองของ Loss Function:

$$h_i = \frac{\partial^2 \mathcal{L}}{\partial \hat{y}_i^2} \quad (7)$$

4. การคำนวณผลลัพธ์สุดท้าย: การอัปเดตค่าในแต่ละต้นไม้จะใช้สูตรจาก Gradient และ Hessian ดังนี้:

$$\hat{y}_t(x) = \hat{y}_{i-1}(x) + \eta \sum_{i=1}^t g_i/h_i \quad (8)$$

5. การประเมินผลและเปรียบเทียบประสิทธิภาพ (Model Evaluation and Comparison)

1. ใช้ตัวชี้วัด MAE, RMSE และ R^2 เพื่อเปรียบเทียบประสิทธิภาพของโมเดล
2. วิเคราะห์ Residual Plots และ Forecast Error Distribution เพื่อประเมินความสอดคล้องของโมเดล

3. ทดสอบความสามารถในการพยากรณ์ระยะสั้นและระยะยาว (Short-term vs. Long-term Forecasting)

6. วิธีการวิเคราะห์ข้อมูล (Data Analysis Methodology) เพื่อเปรียบเทียบความแม่นยำของโมเดลพยากรณ์ จะใช้ตัวชี้วัดทางสถิติที่เป็นที่ยอมรับในงานวิจัยอนุกรมเวลา ได้แก่

1. Mean Absolute Error (MAE) ใช้สำหรับวัดค่าเฉลี่ยของความคลาดเคลื่อนระหว่างค่าจริงและค่าที่พยากรณ์

$$MAE = \frac{\sum_{i=1}^n |x_t - \hat{x}_t|}{n} \quad (9)$$

2. Root Mean Squared Error (RMSE) ใช้สำหรับวัดความคลาดเคลื่อนที่ให้ความสำคัญกับค่าผิดพลาดขนาดใหญ่

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (10)$$

3. Coefficient of Determination (R^2 Score) ใช้วัดว่าสัดส่วนของความแปรปรวนในข้อมูลที่สามารถอธิบายได้โดยโมเดล

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (11)$$

4. Mean Bias Error (MBE) เป็นตัวชี้วัดที่ใช้ในการวัดการเอียงของโมเดล (Bias) โดยการเปรียบเทียบระหว่างค่าพยากรณ์ ($\hat{y}$) กับค่าจริง (y) โดยคำนวณค่าเฉลี่ยของความแตกต่างระหว่างค่าจริงและค่าพยากรณ์

$$MBE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i) \quad (12)$$

5. Symmetric Mean Absolute Percentage Error (SMAPE) เป็นตัวชี้วัดที่ใช้ในการวัดข้อผิดพลาดในรูปแบบเปอร์เซ็นต์ โดยคำนวณจากความแตกต่างระหว่างค่าจริงและค่าพยากรณ์ โดยใช้ค่าเฉลี่ยของค่าจริงและค่าพยากรณ์เป็นตัวหาร ซึ่งช่วยในการจัดการกับค่าศูนย์ที่อาจเป็นปัญหาในการคำนวณ Mean Absolute Percentage Error (MAPE)

$$SMAPE = \frac{1}{n} \sum_{i=1}^n \frac{|Y_i - \hat{Y}_i|}{\frac{(Y_i + \hat{Y}_i)}{2}} \times 100\% \quad (13)$$

6. Residual Plots เป็นเครื่องมือที่ใช้ในการตรวจสอบความสัมพันธ์ระหว่างค่าความผิดพลาด (Residuals) กับค่าที่คาดการณ์หรือพยากรณ์ (Fitted Values) ของโมเดล โดยช่วยในการประเมินว่าโมเดลสามารถจับความสัมพันธ์ในข้อมูลได้ดีหรือไม่ และสามารถบ่งชี้ถึงปัญหาหรือข้อผิดพลาดที่อาจเกิดขึ้นในกระบวนการพยากรณ์ได้ ประเภทของ Residual Plots

- Residuals vs. Fitted Values Plot: การกระจายตัวของ Residuals เมื่อเทียบกับค่าพยากรณ์จะช่วยให้เห็นรูปแบบที่อาจเกิดขึ้นในข้อผิดพลาด เช่น การกระจายตัวที่ไม่เป็นแบบสุ่มอาจบ่งชี้ว่าโมเดลมีปัญหาหรือไม่สามารถอธิบายข้อมูลได้อย่างเหมาะสม

- Residuals vs. Time Plot: การตรวจสอบความสัมพันธ์ของ Residuals กับเวลา (หรือขั้นตอน) ช่วยในการระบุว่าโมเดลมีการจับแนวโน้มหรือฤดูกาลของข้อมูลได้ดีหรือไม่

- Histogram of Residuals: การตรวจสอบการกระจายตัวของ Residuals ว่ามีลักษณะเป็นการกระจายปกติ (Normal Distribution) หรือไม่ ซึ่งสามารถบ่งชี้ถึงข้อผิดพลาดในรูปแบบของ Bias หรือ Variance ของโมเดล

- Q-Q Plot (Quantile-Quantile Plot): ใช้เพื่อเปรียบเทียบการกระจายตัวของ Residuals กับการกระจายตัวปกติ (Normal Distribution) โดยหากจุดกระจายอยู่ใกล้เส้นทแยงมุม จะเป็นการยืนยันว่า Residuals มีการกระจายตัวแบบปกติ

วิธีดำเนินการวิจัย

การเปรียบเทียบประสิทธิภาพของโมเดลพยากรณ์ปริมาณฝุ่น PM2.5 ในกรุงเทพมหานคร ระหว่าง SARIMA และ XGBoost ในกระบวนการเก็บข้อมูลค่าฝุ่น PM2.5 ตั้งแต่ปี 2561 ถึง 2566 ผู้เขียนตั้งข้อสังเกตว่าช่วงต้นของชุดข้อมูล (ปี 2561-2562) และช่วงหลัง (ปี 2565-2566) อาจมีความแตกต่างในแง่แหล่งที่มาของฝุ่น PM2.5 ในกรุงเทพมหานคร ซึ่งส่งผลต่อรูปแบบข้อมูลและความแม่นยำของการพยากรณ์ ข้อมูลจากการทบทวนวรรณกรรมพบว่า ในช่วงปี 2561-2562 แหล่งที่มาหลักของ PM2.5 ได้แก่ การจราจรในเมือง การก่อสร้าง และการเผาในที่โล่งในพื้นที่รอบนอกเมือง (กรมควบคุมมลพิษ, 2562; World Bank, 2020) ขณะที่ในช่วงปีหลัง การเผาในที่โล่งในภาคเหนือและภาคตะวันออกเฉียงเหนือ มีบทบาทเพิ่มขึ้นอย่างมีนัยสำคัญในการส่งผลกระทบต่อคุณภาพอากาศในกรุงเทพมหานคร (กรมควบคุมมลพิษ, 2566) ดังนั้น ชุดข้อมูลที่ใช้ในงานวิจัยนี้จึงอาจมีลักษณะแปรเปลี่ยนของแหล่งกำเนิดฝุ่นตามช่วงเวลา ซึ่งเป็นปัจจัยที่ควรนำมาพิจารณาในการวิเคราะห์ความแม่นยำและความคลาดเคลื่อนของแบบจำลองพยากรณ์ เพื่อเพิ่มความถูกต้องของการตีความผลการวิเคราะห์ในภาพรวมดังนี้

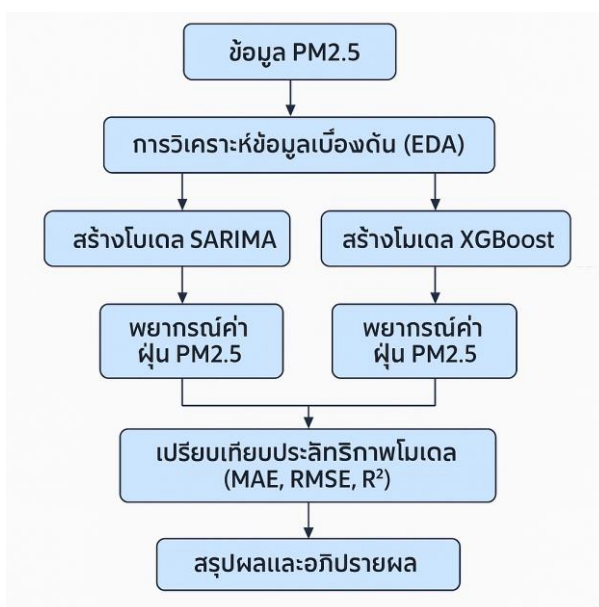
1. การเก็บข้อมูล ข้อมูลค่าฝุ่น PM2.5 ที่ใช้ในการวิเคราะห์ถูกรวบรวมจากสถานีตรวจวัดคุณภาพอากาศของกรมควบคุมมลพิษ จำนวนทั้งหมด 8 สถานี ได้แก่ สถานีดินแดง (03T) สถานีพระนคร (05T) สถานีบางเขน (13T) สถานีบางนา (50T) สถานีลาดกระบัง (39T) สถานีบางขุนเทียน (27T) สถานีสวนลุมพินี (17T) และสถานีริมถนนกาญจนาภิเษก (70T) ซึ่งครอบคลุมพื้นที่ในเขตกรุงเทพมหานครที่มีลักษณะการใช้ที่ดินที่แตกต่างกัน ได้แก่ เขตที่อยู่อาศัยหนาแน่น

2. เขตอุตสาหกรรม และเขตการค้า ข้อมูลที่นำมาใช้เป็นข้อมูลค่าเฉลี่ยรายวัน (Daily Average) ระหว่างปี 2561 ถึง 2566 และได้ผ่านกระบวนการทำความสะอาดข้อมูลเบื้องต้น เช่น การจัดการข้อมูลสูญหาย (Missing Data Handling) และการตรวจสอบค่าผิดปกติ (Outliers Detection) ก่อนนำไปใช้ในการสร้างและประเมินแบบจำลอง

3. สร้างแบบจำลอง SARIMA และ XGBoost โดยใช้ชุดข้อมูลฝึกสอน (Training Set) และชุดข้อมูลทดสอบ (Test Set)

4. คำนวณตัวชี้วัด MAE, RMSE, R^2 , MBE และ SMAPE เพื่อเปรียบเทียบประสิทธิภาพของโมเดล

5. วิเคราะห์ข้อจำกัดของแต่ละโมเดล และพิจารณาความเหมาะสมของการใช้งานในบริบทของข้อมูล



ภาพที่ 1 แสดงขั้นตอนการเปรียบเทียบประสิทธิภาพของแบบจำลองการพยากรณ์ค่าฝุ่น PM2.5 ด้วย SARIMA และ XGBoost

การวิจัยนี้เริ่มต้นจากการรวบรวมข้อมูลค่าฝุ่น PM2.5 รายวันจากสถานีตรวจวัดคุณภาพอากาศของกรมควบคุมมลพิษในช่วงปี 2561–2566 จากนั้นทำการวิเคราะห์ข้อมูลเบื้องต้น (Exploratory Data Analysis: EDA) เพื่อศึกษาลักษณะของข้อมูล เช่น แนวโน้ม ฤดูกาล และความแปรปรวน ก่อนนำข้อมูลผ่านการเตรียมพร้อมไปสร้างแบบจำลองการพยากรณ์ด้วยสองวิธี ได้แก่ SARIMA ซึ่งเป็นเทคนิคทางสถิติ



สำหรับข้อมูลอนุกรมเวลา และ XGBoost ซึ่งเป็นเทคนิค Machine Learning ที่สามารถรองรับข้อมูลที่มีความซับซ้อนและไม่เป็นเชิงเส้น หลังจากสร้างแบบจำลองแล้ว จะดำเนินการพยากรณ์ค่าฝุ่น PM2.5 และประเมินประสิทธิภาพของแต่ละแบบจำลองโดยใช้ตัวชี้วัด ได้แก่ MAE, RMSE และ R^2 เพื่อเปรียบเทียบประสิทธิภาพของโมเดล และสรุปผลว่าวิธีการใดเหมาะสมที่สุดในการพยากรณ์ค่าฝุ่น PM2.5 ในพื้นที่กรุงเทพมหานคร

หลักเกณฑ์ (Criteria) ที่ควรใช้ในการเลือกกรอบแนวคิด (Conceptual Framework)

1. ความสอดคล้องกับวัตถุประสงค์ของการวิจัย (Alignment with Research Objectives)

- กรอบแนวคิดที่เลือกต้องสามารถตอบโจทย์วัตถุประสงค์ของการวิจัยได้โดยตรง
- เช่น ในกรณีศึกษา: ต้องการพยากรณ์และเปรียบเทียบประสิทธิภาพของแบบจำลอง SARIMA กับ XGBoost → กรอบแนวคิดจึงต้องเชื่อมโยงโมเดลเหล่านี้กับตัวชี้วัดประสิทธิภาพ (เช่น MAE, RMSE, R^2)

2. ความมีฐานทางทฤษฎีและงานวิจัยรองรับ (Theoretical and Empirical Support)

- กรอบแนวคิดควรมีทฤษฎีหลักที่สนับสนุน เช่น ทฤษฎีอนุกรมเวลา (Time Series Theory) รองรับ SARIMA และทฤษฎี Machine Learning รองรับ XGBoost
- รวมถึงมีงานวิจัยที่ผ่านมาใช้โมเดลเหล่านี้กับข้อมูล PM2.5 หรือปัญหาสิ่งแวดล้อมคล้ายคลึงกัน

3. ความทันสมัยและสอดคล้องกับสถานการณ์ปัจจุบัน (Recency and Relevance)

- กรอบแนวคิดควรเลือกจากแนวโน้มวิจัยล่าสุด เช่น งานวิจัยปี 2020–2024 ที่แสดงให้เห็นว่าเทคนิค Machine Learning มีบทบาทมากขึ้นในพยากรณ์ PM2.5
- หากกรอบแนวคิดล้าสมัย (เช่น ใช้แต่ ARIMA โดยไม่เปรียบเทียบกับ ML) จะลดความน่าเชื่อถือของงานวิจัย

4. ความครอบคลุมและสมบูรณ์ (Comprehensiveness)

- กรอบแนวคิดควรรวมปัจจัยสำคัญทั้งหมด เช่น ตัวแปรต้น (โมเดลที่ใช้), ตัวแปรตาม (ผลลัพธ์การพยากรณ์), ตัวชี้วัดความแม่นยำ (เช่น RMSE, MAE)
- มีความเชื่อมโยงระหว่างตัวแปรครบถ้วน ไม่เว้นส่วนสำคัญที่มีผลต่อผลการศึกษา

5. ความเหมาะสมกับบริบทของข้อมูล (Context Appropriateness)

- เช่น ในกรณี: ข้อมูล PM2.5 ของกรุงเทพมหานคร มีลักษณะเป็นอนุกรมเวลา รายวัน และมีฤดูกาลชัดเจน → การเลือก SARIMA เหมาะสม และการเลือก XGBoost เพื่อรองรับความไม่เป็นเชิงเส้นของข้อมูลก็เหมาะสม

6. ความเป็นไปได้ในการนำไปใช้จริง (Practicality)

- กรอบแนวคิดต้องสามารถแปลงไปเป็นขั้นตอนการดำเนินงานจริงได้ เช่น การสร้างแบบจำลอง, การแบ่งข้อมูลฝึก-ทดสอบ, การคำนวณตัวชี้วัด
- ไม่ควรเลือกกรอบแนวคิดที่ดีทางทฤษฎี แต่ปฏิบัติจริงไม่ได้เพราะข้อจำกัดของข้อมูลหรือเครื่องมือ



7. ความสามารถในการขยายผลต่อ (Extendibility)

- กรอบแนวคิดที่ดี ควรนำไปปรับใช้ในบริบทอื่นหรือขยายขอบเขตได้ เช่น การทำนายฝุ่น PM10, NO₂ หรือพื้นที่เมืองอื่นนอกจากกรุงเทพฯ

ผลการวิจัย

การศึกษานี้ได้เปรียบเทียบประสิทธิภาพของโมเดล SARIMA และ XGBoost ในการพยากรณ์ปริมาณฝุ่น PM2.5 โดยใช้ตัวชี้วัดทางสถิติ ได้แก่ Mean Absolute Error (MAE), Root Mean Squared Error (RMSE) และ R²

ผลการวิเคราะห์พบว่า โมเดล SARIMA มีค่าความคลาดเคลื่อนต่ำกว่า XGBoost ในทุกตัวชี้วัด โดยมีค่า MAE เท่ากับ 7.64 และ RMSE เท่ากับ 10.65 ซึ่งต่ำกว่าค่าที่ได้จาก XGBoost ที่มีค่า MAE เท่ากับ 13.08 และ RMSE เท่ากับ 15.28 แสดงให้เห็นว่า SARIMA สามารถพยากรณ์ค่าได้ใกล้เคียงกับค่าจริงมากกว่า

นอกจากนี้ ค่า R² ของ SARIMA อยู่ที่ -0.26 ซึ่งสูงกว่าค่าของ XGBoost ที่มีค่า -1.60 แม้ว่าค่า R² ที่เป็นลบจะบ่งบอกว่าโมเดลยังไม่สามารถอธิบายความแปรปรวนของข้อมูลได้ดีนัก แต่ค่า R² ของ XGBoost ที่มีค่าต่ำกว่าบ่งชี้ว่า โมเดล XGBoost อาจไม่เหมาะสมกับชุดข้อมูลนี้ เนื่องจากอาจไม่สามารถจับแนวโน้มของข้อมูลได้ดีพอ

ตารางที่ 1 ผลการเปรียบเทียบประสิทธิภาพของแบบจำลองการพยากรณ์ PM2.5 ด้วย SARIMA และ XGBoost

Model	MAE	RMSE	R ²	MBE	SMAPE(%)
SARIMA	7.639376	10.653977	-0.263203	5.489365	38.07703
XGBoost	12.937441	15.045905	-1.519337	-9.822828	54.506402

จากตารางที่ 1 พบว่า SARIMA มีค่าความคลาดเคลื่อนต่ำกว่า XGBoost ในทุกตัวชี้วัด ซึ่งแสดงให้เห็นว่า SARIMA สามารถพยากรณ์ค่า PM2.5 ได้ดีกว่าในบริบทของข้อมูลชุดนี้

ตารางที่ 2: ค่าความคลาดเคลื่อนของการพยากรณ์ PM2.5 โดย SARIMA และ XGBoost

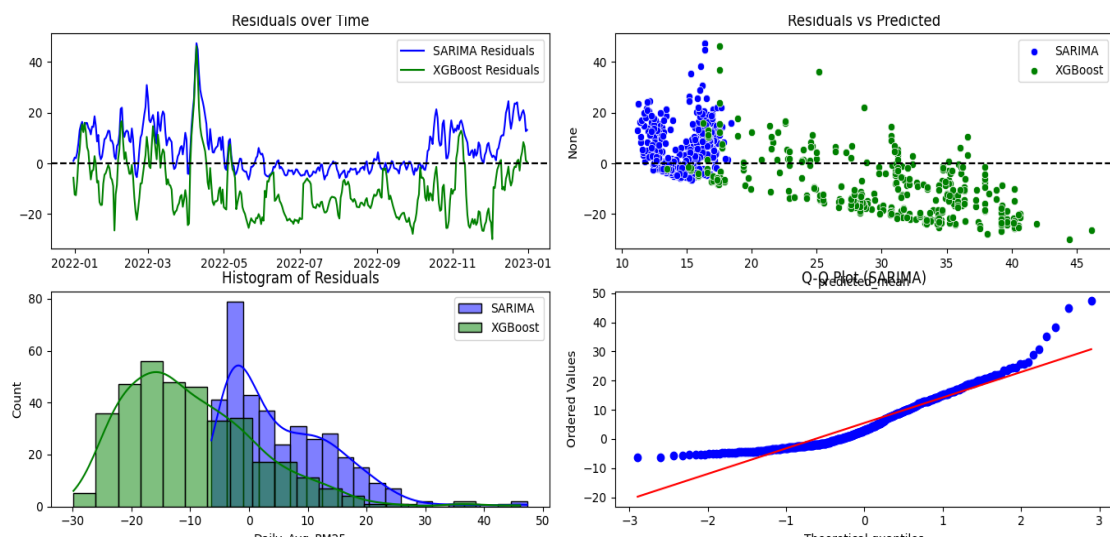
วันที่	Actual	SARIMA Forecast	XGBoost Forecast	SARIMA Residuals	XGBoost Residuals	SARIMA Abs Error	XGBoost Abs Error
2021-12-31	18.1	17.2	23.8	0.9	-5.7	0.9	5.7
2022-01-01	19.3	17.0	31.6	2.3	-12.3	2.3	12.3
2022-01-02	19.0	17.0	31.6	2.0	-12.5	2.0	12.5



ตารางที่ 2: (ต่อ)

วันที่	Actual	SARIMA Forecast	XGBoost Forecast	SARIMA Residuals	XGBoost Residuals	SARIMA Abs Error	XGBoost Abs Error
2022-01-03	21.4	17.3	28.4	4.1	-7.0	4.1	7.0
2022-01-04	25.3	17.4	23.9	7.9	1.3	7.9	1.3
2022-01-05	29.4	16.4	23.9	13.0	5.4	13.0	5.4
2022-01-06	32.3	17.3	20.0	15.0	12.4	15.0	12.4
2022-01-07	32.9	16.6	17.5	16.3	15.4	16.3	15.4
2022-01-08	29.9	16.6	17.5	13.3	12.3	13.3	12.3
2022-01-09	32.3	17.4	16.3	14.9	16.0	14.9	16.0

จากตารางที่ 2 พบว่า SARIMA มีแนวโน้มทำนายค่าสูงกว่าค่าจริงเล็กน้อย ขณะที่ XGBoost มีความแม่นยำมากขึ้นในช่วง



ภาพที่ 2 การเปรียบเทียบผลการพยากรณ์ PM2.5 ระหว่างโมเดล SARIMA และ XGBoost

สรุปและอภิปรายผล

จากการศึกษาเปรียบเทียบประสิทธิภาพของแบบจำลอง SARIMA และ XGBoost ในการพยากรณ์ค่าฝุ่น PM2.5 พบว่าโมเดล SARIMA มีประสิทธิภาพดีกว่า XGBoost ในทุกตัวชี้วัด โดยมีค่าความคลาดเคลื่อนต่ำกว่าและสามารถพยากรณ์ค่าฝุ่น PM2.5 ได้ใกล้เคียงกับค่าจริงมากกว่า อย่างไรก็ตาม ค่า R^2 ที่เป็นค่าติดลบยังสะท้อนให้เห็นว่าโมเดลทั้งสองยังมีข้อจำกัดในการพยากรณ์ ผลการวิจัยพบว่าโมเดล SARIMA มีประสิทธิภาพในการพยากรณ์ค่าฝุ่น PM2.5 ดีกว่า XGBoost เมื่อพิจารณาจากค่า MAE, RMSE และ SMAPE ซึ่งสอดคล้องกับงานวิจัยของ สมชาย วิริยะกุล และคณะ (2562) ที่พบว่า SARIMA



สามารถจับแนวโน้มและรูปแบบตามฤดูกาลได้อย่างมีประสิทธิภาพในข้อมูล PM2.5 ของกรุงเทพมหานคร อย่างไรก็ตาม เมื่อเปรียบเทียบกับงานของ อภิชัย นาคะเวช และคณะ (2563) ซึ่งใช้ XGBoost ร่วมกับตัวแปรแวดล้อม เช่น อุณหภูมิและความชื้น พบว่า XGBoost ให้ผลการพยากรณ์ที่แม่นยำกว่าการใช้ข้อมูล PM2.5 เพียงอย่างเดียว ดังนั้นข้อจำกัดของการศึกษานี้ อาจเกิดจากการใช้เฉพาะข้อมูล PM2.5 โดยไม่มีตัวแปรเสริม อีกทั้ง เมื่อพิจารณาจากงานของปณิธาน จันทร์ศรี และจารุวรรณ ทองก้อน (2564) ที่ใช้โมเดล LSTM พบว่าเทคนิค Deep Learning สามารถจับความสัมพันธ์ที่ซับซ้อนได้ดีกว่าโมเดลเชิงสถิติและ Machine Learning ทั่วไป ซึ่งชี้ให้เห็นว่าในอนาคตการนำ LSTM หรือโมเดลเชิงลึก มาใช้ อาจช่วยเพิ่มประสิทธิภาพการพยากรณ์ PM2.5 ได้ดียิ่งขึ้น

ผลการวิจัยแสดงให้เห็นว่าแบบจำลอง SARIMA มีประสิทธิภาพในการพยากรณ์ค่าฝุ่น PM2.5 สูงกว่าแบบจำลอง XGBoost โดยพิจารณาจากค่าตัวชี้วัดทางสถิติ ได้แก่ MAE, RMSE และ SMAPE ที่มีค่าต่ำกว่า ซึ่งสอดคล้องกับงานวิจัยของสมชาย วิริยะกุล และคณะ (2562) ที่พบว่า SARIMA มีความสามารถในการจับแนวโน้มและรูปแบบเชิงฤดูกาลของข้อมูล PM2.5 ได้ดีในพื้นที่กรุงเทพมหานคร อย่างไรก็ตาม ค่า R^2 ของ SARIMA ในการศึกษานี้ยังคงเป็นค่าติดลบ (-0.26) สะท้อนให้เห็นว่ามีข้อจำกัดในการอธิบายความแปรปรวนของข้อมูล ซึ่งแตกต่างจากงานของ สมชาย วิริยะกุล และคณะ ที่ได้ค่า R^2 เป็นบวก ในระดับที่น่าพอใจ สาเหตุอาจเกิดจากความแปรปรวนของข้อมูล PM2.5 ที่เพิ่มขึ้นในช่วงปีที่ใช้ศึกษา หรือปัจจัยแวดล้อมอื่น ๆ ที่ไม่ได้นำมาพิจารณา เมื่อเปรียบเทียบกับงานของอภิชัย นาคะเวช และคณะ (2563) ซึ่งใช้ XGBoost ร่วมกับตัวแปรนำเข้าหลายตัว เช่น อุณหภูมิ ความชื้น และปริมาณจากรถ พบว่าการเพิ่มปัจจัยแวดล้อมสามารถช่วยเพิ่มประสิทธิภาพของโมเดล XGBoost ได้อย่างมีนัยสำคัญ ในขณะที่การศึกษานี้ใช้เฉพาะข้อมูล PM2.5 เป็นตัวแปรอิสระ ส่งผลให้ XGBoost ไม่สามารถเรียนรู้ความสัมพันธ์เชิงซับซ้อนได้อย่างเต็มที่ และมีค่าความคลาดเคลื่อนสูงกว่า SARIMA นอกจากนี้ ผลการศึกษาของ ปณิธาน จันทร์ศรี และ จารุวรรณ ทองก้อน (2564) ที่ใช้เทคนิค LSTM พบว่าแบบจำลองเชิงลึกสามารถพยากรณ์ข้อมูล PM2.5 ได้แม่นยำกว่าเทคนิค Machine Learning แบบดั้งเดิมและแบบจำลองเชิงสถิติ แสดงให้เห็นว่าข้อมูล PM2.5 ซึ่งมีความแปรปรวนและความซับซ้อนสูงนั้น อาจเหมาะกับการใช้โมเดล Deep Learning มากกว่าการใช้ SARIMA หรือ XGBoost เพียงอย่างเดียว โดยสรุปการศึกษานี้ยืนยันว่า SARIMA เหมาะสมสำหรับการพยากรณ์ข้อมูล PM2.5 ที่มีแนวโน้มและฤดูกาลชัดเจน แต่หากต้องการเพิ่มความแม่นยำในการพยากรณ์ ควรพิจารณาการนำตัวแปรแวดล้อมเพิ่มเติมมาใช้ในการสร้างแบบจำลอง รวมถึงการทดลองใช้โมเดลเชิงลึก เช่น LSTM หรือ Hybrid Models เพื่อปรับปรุงความสามารถในการจับรูปแบบข้อมูลที่ซับซ้อนในอนาคต

ข้อเสนอแนะ

1. ควรเพิ่มตัวแปรแวดล้อม เช่น อุณหภูมิ ความชื้น และความเร็วลม เพื่อช่วยเพิ่มความแม่นยำของโมเดล



2. ควรทดสอบโมเดลเชิงลึก เช่น LSTM หรือ Hybrid Models
3. ควรใช้ชุดข้อมูลที่ครอบคลุมช่วงเวลาที่ยาวขึ้นเพื่อช่วยให้โมเดลสามารถเรียนรู้แนวโน้มของข้อมูลได้ดีขึ้น

เอกสารอ้างอิง

- กรมควบคุมมลพิษ. (2566). รายงานสถานการณ์คุณภาพอากาศในประเทศไทย. สืบค้นเมื่อ 10 พฤษภาคม 2568 จาก <https://www.pcd.go.th/publication/22073>
- บัญชา ปะสีละเตสัง. (2564). บทเรียน XGBoost และการประยุกต์ใช้ในงานพยากรณ์. สืบค้นเมื่อ 10 พฤษภาคม 2568 จาก <https://www.datasciencethailand.com/xgboost-tutorial>
- ปณิธาน จันทศรี และ จารุวรรณ ทองก้อน. (2564). การเปรียบเทียบเทคนิค Machine Learning ในการพยากรณ์ค่าฝุ่น PM2.5. วารสารการคำนวณเชิงปัญญาและระบบอัจฉริยะ, 18(1), 75–89.
- ภูมิฐาน รังคกุลวัฒน์. (2556). การวิเคราะห์อนุกรมเวลาสำหรับเศรษฐศาสตร์และธุรกิจ (พิมพ์ครั้งที่ 1). กรุงเทพฯ : จุฬาลงกรณ์มหาวิทยาลัย.
- สมชาย วิริยะกุล, สุนิสา บุญประเสริฐ และ ศิริพร สุทธิเจริญ. (2562). การพยากรณ์ค่าฝุ่น PM2.5 ในเขตกรุงเทพมหานครโดยใช้แบบจำลองเชิงสถิติ SARIMA. วารสารวิทยาศาสตร์และเทคโนโลยี, 27(3), 45–60.
- อภิชัย นาคะเวช, อนุสรณ์ พิมพ์ทอง และ วิศิษฐ์ บุญชม. (2563). การพยากรณ์ค่าฝุ่น PM2.5 โดยใช้ Extreme Gradient Boosting (XGBoost). ใน อนุชิต คำผา (ประธาน), การประชุมวิชาการวิศวกรรมศาสตร์แห่งชาติ ครั้งที่ 14 (หน้า 225–232). สถาบันวิจัยวิศวกรรมแห่งชาติ, มหาวิทยาลัยเทคโนโลยีสุรนารี, นครราชสีมา.
- เฉลิมพล จตุพร. (2562). การวิเคราะห์อนุกรมเวลาเบื้องต้น (Basic of Time Series Analysis). สืบค้นเมื่อ 16 กุมภาพันธ์ 2567 จาก <https://cj007blog.files.wordpress.com/2020/04/03-basic-of-time-series-analysis.pdf>
- Box, G. E. P., Jenkins, G. M., & Reinsel, G. C. (2015). Time series analysis: Forecasting and control (5th ed.). John Wiley & Sons.
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (pp. 785–794). <https://doi.org/10.1145/2939672.2939785>
- Willmott, C. J., & Matsuura, K. (2005). Advantages of the mean absolute error (MAE) over the root mean square error (RMSE) in assessing average model performance. Climate Research, 30(1), 79–82. <https://doi.org/10.3354/cr030079>



มหาวิทยาลัยหาดใหญ่
HATYAI UNIVERSITY

World Bank. (2020). Thailand economic monitor: Thailand in the time of COVID-19.

Retrieved from <https://documents.worldbank.org/en/publication/documents-reports/documentdetail/882361593210407472>

World Health Organization. (2021). WHO global air quality guidelines: Particulate matter (PM_{2.5} and PM₁₀), ozone, nitrogen dioxide, sulfur dioxide and carbon monoxide.

<https://www.who.int/publications/i/item/9789240034228>