



คู่มือการดำเนินงาน

กิจกรรมย่อยการประเมินเนื้อที่ยื่นต้นปาล์มน้ำมันด้วยเทคโนโลยีปัญญาประดิษฐ์
(ARTIFICIAL INTELLIGENCE: AI) ภายใต้โครงการประยุกต์ใช้เทคโนโลยี
ภูมิสารสนเทศเพื่อเพิ่มประสิทธิภาพการพยากรณ์ผลผลิตสินค้าเกษตร ปี พ.ศ. 2568

จัดทำโดย
ส่วนภูมิสารสนเทศการเกษตร
ศูนย์สารสนเทศการเกษตร
สำนักงานเศรษฐกิจการเกษตร

**กิจกรรมย่อยการประเมินเนื้อที่ยืนต้นปาล์มน้ำมันด้วยเทคโนโลยีปัญญาประดิษฐ์
(Artificial Intelligence: AI) ภายใต้โครงการประยุกต์ใช้เทคโนโลยีภูมิสารสนเทศเพื่อเพิ่มประสิทธิภาพ
การพยากรณ์ผลผลิตสินค้าเกษตร ปี พ.ศ. ๒๕๖๔**

๑. ที่มาและความสำคัญของโครงการ

สำนักงานเศรษฐกิจการเกษตร โดยศูนย์สารสนเทศการเกษตร ดำเนินงานโครงการเพิ่มประสิทธิภาพการจัดทำสารสนเทศการเกษตรและการบริหารจัดการข้อมูลขนาดใหญ่ ซึ่งมีความจำเป็นต้องมีการพัฒนาสารสนเทศเศรษฐกิจการเกษตรในด้านปริมาณการผลิต ผลผลิตต่อไร่ ต้นทุนการผลิต ราคาสินค้าเกษตร การพยากรณ์ปริมาณการผลิต การประยุกต์ใช้เทคโนโลยีภูมิสารสนเทศ รวมทั้งจัดทำฐานข้อมูลและฐานข้อมูลขนาดใหญ่ให้ครอบคลุม สามารถตอบสนองความต้องการของผู้ใช้ข้อมูลที่มีความหลากหลายมากขึ้น โดยเฉพาะสินค้าที่มีความสำคัญเชิงนโยบาย ซึ่งกิจกรรมหนึ่งที่สำคัญภายใต้โครงการฯ กำหนดให้มีการประยุกต์ใช้เทคโนโลยีภูมิสารสนเทศ (Geo-Informatics: GI) ที่เป็นเทคโนโลยีที่มีการบูรณาการความรู้และเทคโนโลยีจากศาสตร์ที่สำคัญ ๆ ประกอบด้วย เทคโนโลยีการสำรวจระยะไกล (Remote Sensing: RS) เทคโนโลยีระบบกำหนดตำแหน่งพิกัดบนพื้นโลก (Global Positioning System: GPS) หรือ Global Navigation Satellite System (GNSS) และระบบสารสนเทศภูมิศาสตร์ (Geographic Information System: GIS) ในการจัดทำข้อมูลเกษตร ในกิจกรรมย่อยการบูรณาการข้อมูลเชิงพื้นที่เนื้อที่เพาะปลูกพืชเศรษฐกิจของประเทศ กิจกรรมย่อยการประยุกต์ใช้เทคโนโลยีภูมิสารสนเทศในการพยากรณ์ผลผลิตสินค้าเกษตร และการบริหารจัดการข้อมูลเชิงพื้นที่ กิจกรรมย่อยการจัดทำแผนที่แสดงเนื้อที่เพาะปลูก/ยืนต้นพืชเศรษฐกิจเชิงเดี่ยว (Single Maps) และกิจกรรมย่อยการจัดทำข้อมูลในระดับจังหวัดของสำนักงานเศรษฐกิจการเกษตรที่ ๑ - ๑๒

ในกิจกรรมย่อยการประยุกต์ใช้เทคโนโลยีภูมิสารสนเทศในการพยากรณ์ผลผลิตสินค้าเกษตร และการบริหารจัดการข้อมูลเชิงพื้นที่ มีวัตถุประสงค์หลักในการนำเทคโนโลยีภูมิสารสนเทศ (Geo-Informatics) มาประยุกต์ใช้ในการแปลวิเคราะห์ และจำแนกเนื้อที่ยืนต้นปาล์มน้ำมัน ด้วยการนำข้อมูลภาพถ่ายดาวเทียมในระบบ Optical Sensor ที่มีรายละเอียดจุดภาพปานกลาง (Medium Satellite Imageries) ในช่วงเวลาที่สอดคล้องกับค่านิยามด้านเกษตร (Crop Definition) และปฏิทินการเพาะปลูกพืช (Crop Calendar) มาวิเคราะห์ร่วมกับแบบจำลองปัญญาประดิษฐ์ (Artificial Intelligence: AI) ที่หลากหลาย ประกอบด้วย วิธี Random Forest (RF), Support Vector Machine (SVM), RetinaNet และ Convolutional Neural Network (ConvNet/CNN) จากนั้น จึงนำผลการจำแนกเนื้อที่ยืนต้นปาล์มน้ำมันจากแต่ละวิธี มาเปรียบเทียบศักยภาพและความเหมาะสมในการนำมาใช้ในการจัดทำข้อมูลการเกษตร ด้วยการประเมินผลความถูกต้องทั้งหมด (Overall Accuracy: OA) และคุณลักษณะการจำแนกเชิงพื้นที่เมื่อเปรียบเทียบกับชุดข้อมูลเนื้อที่ยืนต้นปาล์มน้ำมัน ปี ๒๕๖๕ ที่ได้มีการจัดทำด้วยวิธีการแปลวิเคราะห์ข้อมูลด้วยสายตา (Visual Interpretation) อย่างไรก็ตาม การจัดทำข้อมูลข้างต้นยังมีอุปสรรคในด้านการแปลวิเคราะห์ภาพถ่ายจากดาวเทียมในระบบ Optical Sensor ที่ทำให้เจ้าหน้าที่ไม่สามารถแปลวิเคราะห์ภาพถ่ายจากดาวเทียมในบริเวณที่มีเมฆปกคลุมได้ ดังนั้น เพื่อเป็นการแก้ไขปัญหาดังกล่าว การดำเนินโครงการจึงได้นำภาพถ่ายจากดาวเทียมในระบบเรดาร์ช่องเปิดสังเคราะห์ (Synthetic Aperture Radar: SAR) คือ ดาวเทียม Sentinel-๑ ที่มีความสามารถในการถ่ายภาพทะลุเมฆมาจำแนกเนื้อที่ยืนต้นปาล์มน้ำมันร่วมกับเทคโนโลยี AI อันถือเป็นการเพิ่มประสิทธิภาพรวมถึงการขยายขีดความสามารถในการนำเซนเซอร์ (Sensor) อื่นของภาพถ่ายดาวเทียมมาใช้ในการแปล

วิเคราะห์เนื้อที่เกษตร โดยการดำเนินโครงการฯ กำหนดพื้นที่ศึกษา ปี พ.ศ. ๒๕๖๘ คือ จังหวัดชลบุรี ในฐานะแหล่งปลูกปาล์มน้ำมันสำคัญของภาคตะวันออก จากเดิมที่มีการดำเนินงานในพื้นที่จังหวัดสุราษฎร์ธานี ในปี พ.ศ. ๒๕๖๖ และจังหวัดกระบี่ ในปี พ.ศ. ๒๕๖๗

๒. ข้อมูลที่ใช้ในการวิเคราะห์ข้อมูล

๒.๑ ภาพถ่ายจากดาวเทียม Sentinel-๑

ในการดำเนินงานโครงการฯ กำหนดใช้ภาพถ่ายดาวเทียม Sentinel-๑A ในวันที่ ๑๔ ธันวาคม ๒๕๖๖ ซึ่งเป็นช่วงเวลาที่ใกล้เคียงกับการสำรวจภาคสนามมากที่สุด สามารถดาวน์โหลดได้จากผู้ให้บริการ European Space Agency (ESA) โดยกำหนดให้มีการดาวน์โหลดข้อมูลในรูปแบบ Interferometric Wide Swath Mode (IW) ที่มีความกว้างของภาพ (Swath Width) เป็น ๒๕๐ กิโลเมตร และมีขนาดจุดภาพเชิงพื้นที่ (Spatial Resolution) เป็น ๕ x ๒๐ เมตร กำหนดการรับ-ส่งสัญญาณ ทั้งแบบ VV และ VH Polarizations ที่มีค่า Incidence Angles อยู่ระหว่าง ๒๙.๑° - ๔๖° และใช้ผลิตภัณฑ์แบบ Level-๑ แบบ Single Look Complex (SLC) ซึ่งเป็นภาพที่ Slant Range ด้วย Azimuth Imaging Plane ในการได้มาซึ่งข้อมูลภาพถ่ายดาวเทียม

ตารางคุณลักษณะของภาพถ่ายดาวเทียม Sentinel-๑ ในโหมดที่เลือกปฏิบัติงาน

Feature	Sentinel-๑ (A/B)
ประเภทของเซนเซอร์	Active (Radar)
โหมดของดาวเทียม (Mode)	Interferometric Wide Swath
แบนด์ (Band)	C-band
Polarization	VH และ VV
Spatial Resolution (m)	๑๔
Swath Width (km)	๒๕๐

๒.๒ ข้อมูลพิกัดแปลงตัวอย่าง

การจัดเก็บข้อมูลค่าพิกัดเนื้อที่ยืนต้นปาล์มน้ำมันตัวอย่าง และพื้นที่อื่น ๆ ดำเนินการโดยการสุ่มจุดตัวอย่างแบบสุ่มตามชั้น (Random Stratified Sampling) ซึ่งเป็นเทคนิคการสุ่มตัวอย่างที่ช่วยเพิ่มความแม่นยำในการสำรวจ หรือวิจัย โดยการแบ่งประชากรเป้าหมายออกเป็นกลุ่มย่อย (หรือชั้น) ที่มีลักษณะเฉพาะร่วมกัน แล้วทำการสุ่มตัวอย่างจากแต่ละกลุ่มย่อยเหล่านั้น ด้วยโปรแกรมประยุกต์สำเร็จรูป ArcGIS Pro แบบอัตโนมัติ ซึ่งทุกจุดจะต้องอยู่ห่างกันอย่างน้อย ๕ กิโลเมตร แบ่งเป็นพื้นที่ปลูกปาล์มน้ำมัน จำนวน ๒๑๕ จุด และพื้นที่การใช้ประโยชน์ที่ดินประเภทอื่น จำนวน ๑๐๕ จุด จุดตัวอย่างเหล่านี้จะถูกนำมาแบ่งเป็นชุดข้อมูลเพื่อการฝึกหัด (Training Set) และการทดสอบแบบจำลอง (Testing Set) ในอัตราส่วนร้อยละ ๗๐ ต่อ ๓๐ ตามลำดับ รายละเอียดดังนี้



ตารางจำนวนจุดตัวอย่างเพื่อการฝึกหัด (Training Set) และการทดสอบแบบจำลอง (Testing Set)

การใช้ประโยชน์ที่ดิน	จำนวนจุดตัวอย่างเพื่อการฝึกหัด (Training Set)	จำนวนจุดตัวอย่างเพื่อการทดสอบแบบจำลอง (Testing Set)
ปาล์มน้ำมัน	๑๕๐	๖๕
อื่น ๆ	๗๔	๓๑
รวมทั้งหมด	๒๒๔	๙๖

๓. แบบจำลอง (Models) ที่ใช้ในการดำเนินงาน

๓.๑ Random Forest (RF)

Random Forest (RF) เป็นเทคนิคการเรียนรู้ของเครื่อง (Machine Learning) ที่ใช้สำหรับงานการจำแนก (Classification) และการทำนาย (Regression) โดยพัฒนามาจากการรวมกันของหลาย ๆ ต้นไม้การตัดสินใจ (Decision Trees) เพื่อเพิ่มความแม่นยำและลดความเอนเอียงของการทำนาย หลักการทำงานของ Random Forest ประกอบด้วย ๓ ขั้นตอน ดังนี้

๑) การสร้างชุดข้อมูลย่อย (Bootstrap Sampling): จากชุดข้อมูลเดิม สุ่มตัวอย่างย่อย ๆ ด้วยการสุ่ม (Random Sampling with Replacement) ซึ่งหมายความว่าตัวอย่างเดียวกันอาจถูกเลือกหลายครั้ง

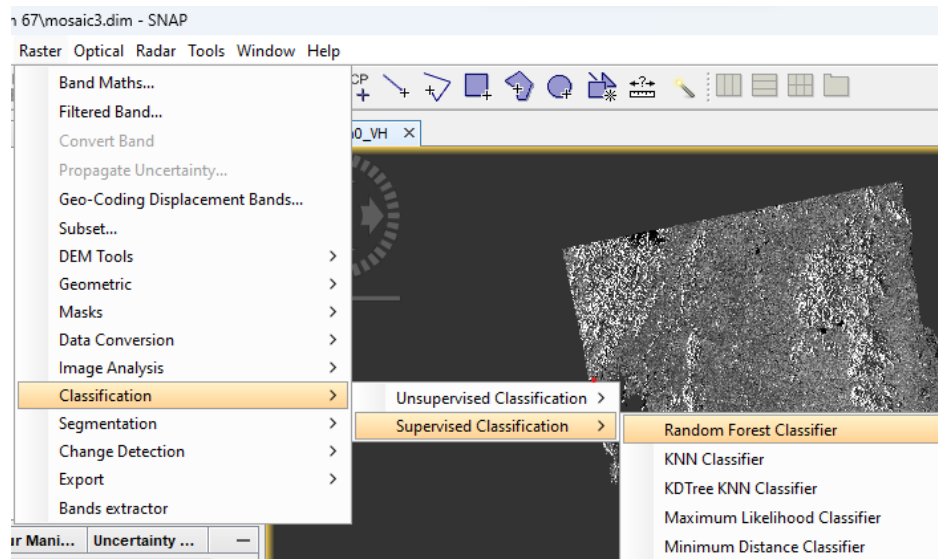
๒) การสร้างต้นไม้การตัดสินใจ (Decision Trees): จากชุดข้อมูลย่อยเหล่านี้ จะสร้างต้นไม้การตัดสินใจหลาย ๆ ต้น โดยการสร้างแต่ละต้นไม้จะใช้ส่วนของข้อมูลที่แตกต่างกัน

๓) การทำการตัดสินใจแบบรวบรวมผล (Aggregation): สำหรับการจำแนกประเภท ใช้วิธีการโหวตเสียงข้างมาก (Majority Voting) จากผลการทำนายของต้นไม้ทุกต้น ส่วนสำหรับการทำนายใช้ค่าเฉลี่ยจากผลการทำนายของต้นไม้ทุกต้น

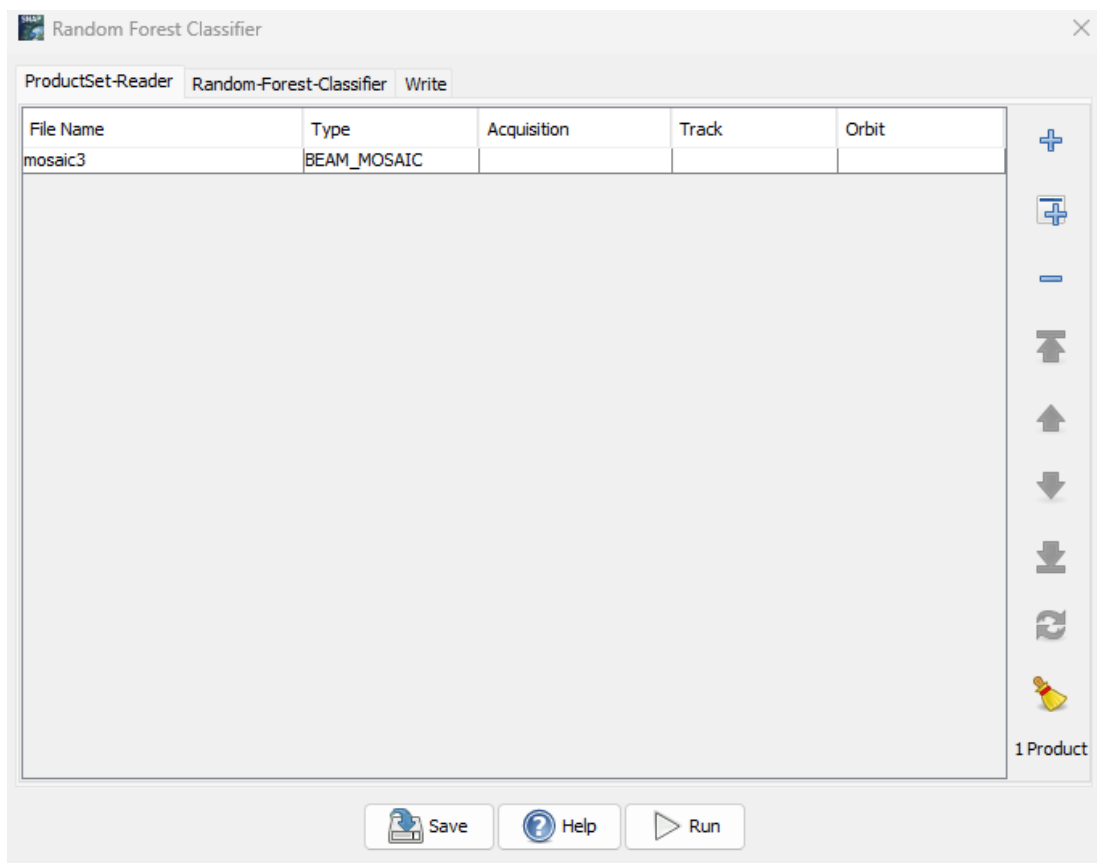
ข้อดีของการประยุกต์ใช้แบบจำลอง Random Forest เพื่อการจำแนกการใช้ประโยชน์ที่ดินบนภาพถ่ายดาวเทียม คือ การใช้ต้นไม้การตัดสินใจหลายต้นและการสุ่มชุดข้อมูลย่อย จะทำให้แบบจำลองมีความทนทานต่อปัญหา Overfitting ซึ่งมักเกิดจากการเรียนรู้จากข้อมูลจุดตัวอย่างมากเกินไปของแบบจำลอง ทั้งยังสามารถจัดการกับข้อมูลจุดตัวอย่างที่ไม่สมบูรณ์ได้เป็นอย่างดี เนื่องจากการตัดสินใจของต้นไม้แต่ละต้นเป็นอิสระต่อกัน

ในการดำเนินโครงการฯ จะใช้แบบจำลอง Random Forest ที่เป็นเครื่องมือหนึ่งของโปรแกรมสำเร็จรูป The Sentinel Application Platform (SNAP) ซึ่งสามารถดาวน์โหลดและติดตั้งได้ฟรีจากผู้ให้บริการ European Space Agency (ESA) มีความสามารถในการนำเข้า ประมวลผล วิเคราะห์ และส่งออกข้อมูลภาพถ่ายดาวเทียม ทั้งในระบบ SAR และ Optical Sensor โดยสามารถเรียกใช้แบบจำลอง Random Forest ได้จากแถบเครื่องมือ Raster → Classification → Supervised Classification → Random Forest Classifier จากนั้น จะปรากฏหน้าต่างเครื่องมือ เพื่อให้เลือกภาพถ่ายดาวเทียมที่จะทำการประมวลผล

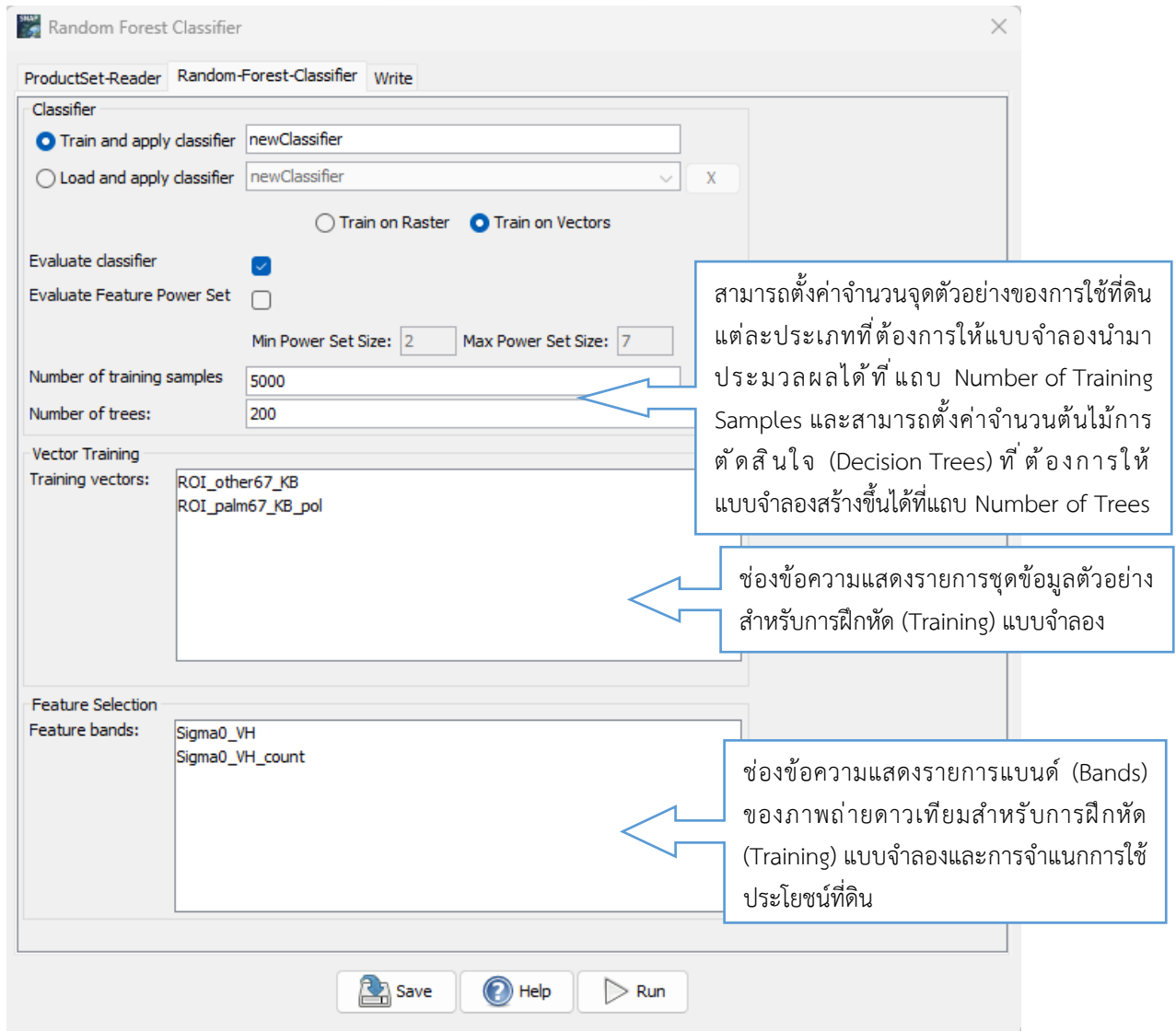




ภาพการเรียกใช้แบบจำลอง Random Forest จากแถบเมนูเครื่องมือ



ภาพแถบเครื่องมือ ProductSet-Reader ในหน้าต่าง Random Forest Classifier เพื่อการนำเข้าภาพถ่ายดาวเทียมที่ต้องการจำแนกการใช้ประโยชน์ที่ดิน



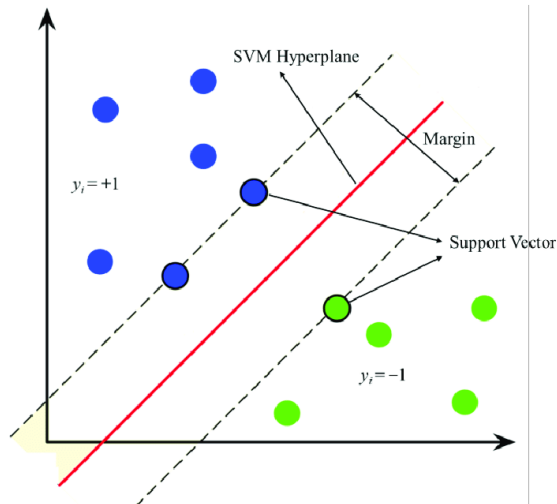
ภาพแถบเครื่องมือ Random-Forest Classifier ในหน้าต่าง Random Forest Classifier

การประมวลผลข้อมูลด้วยแบบจำลอง Random Forest จะกำหนดให้แบบจำลองใช้ต้นไม้การตัดสินใจ (Decision Trees) จำนวน ๕๐, ๑๐๐, ๒๐๐ และ ๕๐๐ ต้น เพื่อให้แบบจำลองประเมินศักยภาพของจำนวนต้นไม้ในเบื้องต้น โดยโปรแกรม SNAP จะคำนวณค่าความถูกต้องของการจำแนกเบื้องต้นของแต่ละวิธีเพื่อให้ผู้ใช้งานสามารถกำหนดจำนวนต้นไม้การตัดสินใจ (Decision Trees) ที่เหมาะสมต่องานวิจัยได้ ซึ่งในการดำเนินโครงการฯ นี้ ได้ผลลัพธ์ว่าควรใช้ต้นไม้การตัดสินใจ (Decision Trees) จำนวน ๑๐๐ ต้น จำแนกร่วมกับภาพถ่ายดาวเทียมใน VV polarization ซึ่งเป็นจำนวนที่ได้ค่าความถูกต้องเบื้องต้นสูงที่สุด

๓.๒ Support Vector Machine (SVM)

Support Vector Machine (SVM) เป็นอัลกอริทึมการเรียนรู้ของเครื่องที่ใช้ในการจัดหมวดหมู่ (Classification) และการวิเคราะห์การถดถอย (Regression) มีความสามารถโดดเด่นในการจัดการกับข้อมูลที่ไม่สามารถแบ่งแยกได้ด้วยเส้นตรง (Non-Linearly Separable) โดยใช้ฟังก์ชันเคอร์เนล (Kernel Function) ในการแปลงข้อมูลไปยังมิติที่สูงขึ้น ถึงแม้ว่า SVM จะมีลักษณะเป็นโมเดลกล่องดำ (Black-Box Model) แต่การวิเคราะห์เวกเตอร์สนับสนุนสามารถให้ข้อมูลเชิงลึกเกี่ยวกับลักษณะของข้อมูลที่สำคัญต่อการจำแนก

ประเภทได้ SVM มีเป้าหมายหลักคือการค้นหาไฮเปอร์เพลน (Hyperplane) ที่ดีที่สุดในการแบ่งข้อมูลออกเป็นกลุ่มต่าง ๆ ให้ได้อย่างชัดเจนที่สุด โดย SVM จะสร้างเส้นแบ่ง (สำหรับข้อมูล ๒ มิติ) หรือไฮเปอร์เพลน (สำหรับข้อมูลมากกว่า ๒ มิติ) ที่สามารถแบ่งข้อมูลออกเป็นกลุ่มได้อย่างชัดเจนที่สุด ด้วยการพิจารณาระยะห่างระหว่างไฮเปอร์เพลนกับจุดข้อมูลที่ใกล้ที่สุดของแต่ละกลุ่ม เรียกว่าระยะขอบ และ SVM จะพยายามหาไฮเปอร์เพลนที่ทำให้ระยะขอบนี้กว้างที่สุด ดังนั้น จุดข้อมูลที่อยู่ใกล้ไฮเปอร์เพลนที่สุดจะมีผลต่อการกำหนดตำแหน่งของไฮเปอร์เพลน



ภาพจำลองการทำงานของ SVM ด้วยการสร้างไฮเปอร์เพลน (Hyperplane) ระหว่างชุดข้อมูล ๒ กลุ่ม (Yang et al., ๒๐๒๒)

ฟังก์ชันเคอร์เนล (Kernel Function) เป็นเครื่องมือสำคัญที่ใช้ใน SVM เพื่อช่วยในการจัดการกับข้อมูลที่ไม่สามารถแบ่งแยกได้ด้วยเส้นตรง (Non-Linearly Separable Data) โดยหลักการคือ การแปลงข้อมูลจากปริภูมิที่มีมิติต่ำไปสู่ปริภูมิที่มีมิติสูงขึ้น ช่วยให้ SVM สามารถสร้างไฮเปอร์เพลนที่ซับซ้อนมากขึ้นเพื่อแบ่งข้อมูลได้อย่างแม่นยำ ทำให้ข้อมูลที่ซับซ้อนกลายเป็นข้อมูลเชิงเส้นที่สามารถแบ่งแยกได้ง่ายขึ้น ซึ่งในโลกแห่งความเป็นจริง ข้อมูลมักมีความสัมพันธ์ที่ซับซ้อนและไม่เป็นเชิงเส้น ฟังก์ชันเคอร์เนลจะช่วยให้ SVM สามารถจับความสัมพันธ์เหล่านี้ได้ การเลือกฟังก์ชันเคอร์เนลที่เหมาะสมขึ้นอยู่กับลักษณะของข้อมูล และปัญหาที่ต้องการแก้ไข โดยทั่วไปจะต้องทำการทดลองกับฟังก์ชันเคอร์เนลที่แตกต่างกันหลายๆ แบบ เพื่อหาฟังก์ชันที่ให้ผลลัพธ์ที่ดีที่สุด โดยประเภทของฟังก์ชันเคอร์เนลที่นิยมใช้ ได้แก่

- ๑) Linear kernel เหมาะสำหรับข้อมูลที่สามารถแบ่งแยกได้ด้วยเส้นตรง
- ๒) Polynomial kernel เหมาะสำหรับข้อมูลที่มีความสัมพันธ์แบบพหุนาม
- ๓) Radial Basis Function (RBF) kernel เป็นฟังก์ชันเคอร์เนลที่นิยมใช้มากที่สุด เนื่องจากมีความยืดหยุ่นสูง และสามารถจับความสัมพันธ์ที่ซับซ้อนได้ดี

ในการดำเนินโครงการฯ จะใช้โปรแกรมประยุกต์สำเร็จรูป ArcGIS Pro ซึ่งเป็นผลิตภัณฑ์จากบริษัท อีเอสอาร์ไอ (ประเทศไทย) จำกัด และจำเป็นต้องได้รับใบอนุญาต (License) ในการใช้งานโปรแกรม โดยเฉพาะชุดเครื่องมือ Machine Learning ที่เป็นเครื่องเฉพาะสำหรับผู้ใช้งานระดับสูง (Advance User) โดยทั่วไป ArcGIS Pro จะทำการเลือกเคอร์เนลที่เหมาะสมกับชุดข้อมูล ให้โดยอัตโนมัติ และมักเริ่มจากเคอร์เนลแบบ RBF เนื่องจากเป็นฟังก์ชันที่นิยมใช้และให้ผลลัพธ์ที่ดีในหลายกรณี นอกจากนี้

การปรับค่าพารามิเตอร์ของฟังก์ชันคอร์เนล เช่น ค่า C (Regularization Parameter) และ gamma ก็มีความสำคัญอย่างยิ่งต่อประสิทธิภาพของโมเดล

การกำหนดค่า C ที่เหมาะสมจะช่วยควบคุมความสมดุล ระหว่างการลดข้อผิดพลาดในการฝึก (training error) และการลดความซับซ้อนของโมเดล (model complexity) หากกำหนดค่า C สูงเกินไป อาจทำให้แบบจำลองเกิดการเรียนรู้ข้อมูลตัวอย่างมากเกินไป หรือ Overfitting ทำให้แบบจำลองทำงานได้ดีกับเฉพาะข้อมูลฝึกหัด (Training Samples) แต่ไม่สามารถนำไปใช้กับชุดข้อมูลอื่นได้ ในทางกลับกัน หากกำหนดค่า C น้อยเกินไปจะทำให้แบบจำลองไม่สามารถสกัดคุณลักษณะของชุดข้อมูลตัวอย่างได้ดี ส่งผลต่อความผิดพลาดในการจำแนก

ในทำนองเดียวกัน การกำหนดค่า Gamma จะช่วยควบคุมอิทธิพลของแต่ละชุดข้อมูลต่อการสร้างไฮเปอร์เพลน (Hyperplane) หากกำหนดค่า Gamma สูง จะทำให้อิทธิพลของแต่ละชุดข้อมูลมีขอบเขตแคบ แบบจำลองจะสามารถตรวจจับรายละเอียดของชุดข้อมูลตัวอย่างได้ดี แต่มีความเสี่ยงต่อ overfitting ในขณะที่การกำหนดค่า Gamma ต่ำ จะทำให้อิทธิพลของแต่ละชุดข้อมูลมีขอบเขตกว้าง ทำให้โมเดลไม่สามารถจับลักษณะเฉพาะของชุดข้อมูลตัวอย่างได้ดีพอ

- ← กำหนดแบบจำลองสำหรับการจำแนกภาพถ่ายดาวเทียม
- ← เลือกชุดข้อมูลฝึกหัด (Training Samples) ที่เตรียมไว้
- ← กำหนดจำนวนตัวอย่างต่อประเภท/ชั้น ที่ต้องการจำแนก

ภาพการเรียกใช้เครื่องมือ Image Classification จาก ArcGIS Pro

๓.๓ RetinaNet

Retinanet เป็นหนึ่งในสถาปัตยกรรมที่ได้รับการพัฒนาขึ้นเพื่อการตรวจจับวัตถุ (Object Detection) อย่างมีประสิทธิภาพ โดยมีจุดเด่นในการสร้างสมดุลระหว่างความแม่นยำ (accuracy) และความเร็วในการประมวลผล โมเดลนี้ได้รับความนิยมอย่างแพร่หลายในงานด้านการมองเห็นด้วยคอมพิวเตอร์ (Computer Vision) รวมถึงการประยุกต์ใช้กับการจำแนกภาพถ่ายดาวเทียมที่ต้องการ

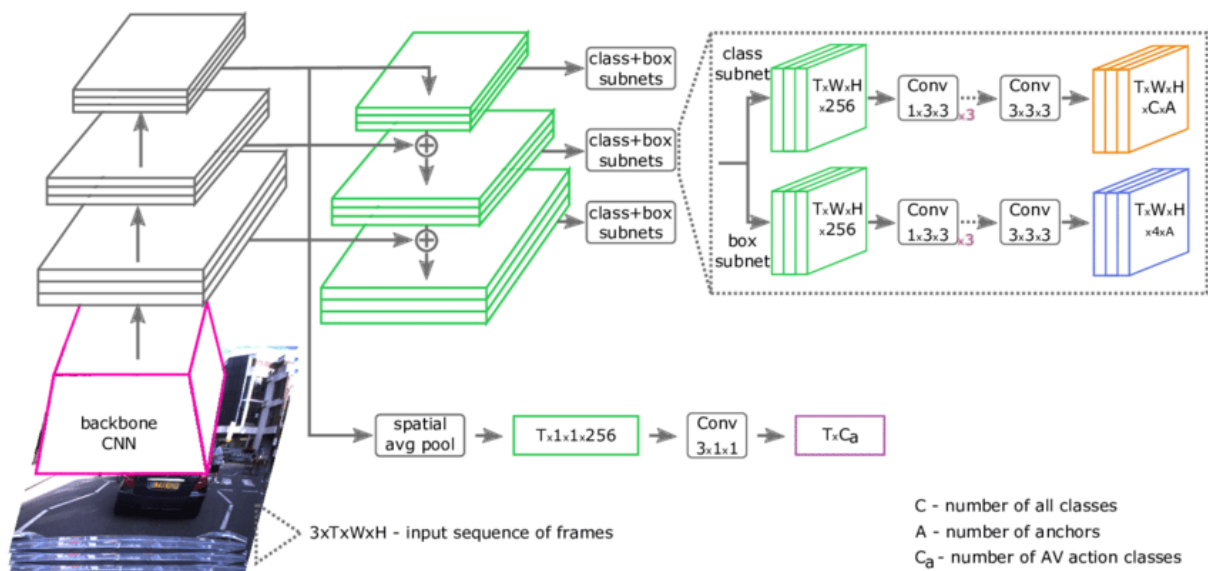
ความแม่นยำสูงเพื่อสนับสนุนการวิเคราะห์และการตัดสินใจในงานด้านสิ่งแวดล้อม การเกษตร และการจัดการทรัพยากร Retinanet ได้รับการพัฒนาขึ้นโดย Lin et al. (๒๐๑๗) โดยเป็นโมเดลที่นำแนวคิดของโครงข่ายประสาทเทียมแบบ Convolutional Neural Network (CNN) มาประยุกต์ใช้งานในลักษณะของโครงสร้างที่เรียกว่า Feature Pyramid Network (FPN) เพื่อรองรับการตรวจจับวัตถุที่มีขนาดแตกต่างกัน โมเดลนี้ถูกออกแบบมาให้มีประสิทธิภาพสูงทั้งในแง่ของความแม่นยำและความเร็ว และยังสามารถทำงานได้ดีกับข้อมูลที่ซับซ้อนหรือไม่สมดุล เช่น ภาพถ่ายดาวเทียมที่มีการกระจายตัวของวัตถุในระดับที่ต่างกัน Retinanet ใช้แนวคิดสำคัญสองประการ ได้แก่

๑) Feature Pyramid Network (FPN): ใช้โครงสร้างแบบพีระมิดในการสกัดคุณลักษณะจากภาพในหลายระดับความละเอียด ทำให้สามารถตรวจจับวัตถุขนาดเล็กและใหญ่ได้ในภาพเดียวกัน

๒) Focal Loss: Loss function ที่ถูกออกแบบมาเพื่อลดผลกระทบจากความไม่สมดุลของข้อมูล โดยการลดน้ำหนักของตัวอย่างที่ง่ายต่อการจำแนก และเพิ่มน้ำหนักให้กับตัวอย่างที่ยากต่อการจำแนก

โครงสร้างของ Retinanet แบ่งออกเป็นสามส่วนหลัก ดังนี้

- ๑) Backbone Network: เช่น ResNet-๕๐ หรือ ResNet-๑๐๑ สำหรับการสกัดคุณลักษณะเริ่มต้น
- ๒) FPN: ช่วยในการสร้างพีระมิดคุณลักษณะจากข้อมูลหลายระดับ
- ๓) Subnetwork: สำหรับการทำนาย Bounding Box และการจัดประเภทวัตถุ



ภาพแสดงโครงสร้างของแบบจำลอง RetinaNet (ที่มา Singh et al., ๒๐๒๑)

ในการดำเนินโครงการฯ แบบจำลอง RetinaNet สามารถพัฒนาโมเดลด้วยภาษา Python และใช้ไลบรารีที่เกี่ยวข้องสำหรับการสร้าง ติดตั้ง และประมวลผลโมเดลได้อย่างเหมาะสม ดังนี้

- ๑) TensorFlow หรือ PyTorch: เป็นเฟรมเวิร์กหลักที่ใช้สร้างและฝึกอบรมโมเดล
- ๒) Keras (หากใช้ TensorFlow): ช่วยให้การพัฒนาโมเดลง่ายและมีประสิทธิภาพ
- ๓) OpenCV: สำหรับการประมวลผลภาพเบื้องต้น เช่น การโหลดและแปลงภาพถ่ายดาวเทียม
- ๔) Matplotlib หรือ Seaborn: สำหรับการแสดงผลภาพและผลลัพธ์

๕) Pre-trained Weights หรือ Model Zoo: เช่น Weights จาก COCO Dataset เพื่อเพิ่มประสิทธิภาพการฝึกอบรม

ขั้นตอนการพัฒนา RetinaNet สำหรับการจำแนกภาพถ่ายดาวเทียม

๑) โหลดและเตรียมข้อมูล โดยภาพถ่ายดาวเทียมจะต้องถูกจัดเตรียมในรูปแบบของชุดข้อมูล เช่น COCO หรือ Pascal VOC เป็นต้น และข้อมูลจะต้องมีการสร้าง Annotation File ซึ่งกำหนด Bounding Boxes และ Labels สำหรับวัตถุในภาพ

๒) โมเดล RetinaNet สามารถสร้างได้ด้วยไลบรารี TensorFlow หรือ PyTorch โดยใช้สถาปัตยกรรม Backbone เช่น ResNet-๕๐ หรือ ResNet-๑๐๑

๓) ตั้งค่าพารามิเตอร์ที่สำคัญต่อการเรียนรู้ของแบบจำลอง ประกอบด้วย Anchor Ratios (ค่าอัตราส่วนของ Bounding Box ที่กำหนดรูปร่างของวัตถุ) Learning Rate (กำหนดความเร็วในการอัปเดตน้ำหนัก) Batch Size (จำนวนข้อมูลที่ใช้ในแต่ละรอบการฝึกอบรม) Focal Loss (ฟังก์ชัน Loss ที่ใช้เพื่อลดปัญหาความไม่สมดุลของข้อมูล) ตลอดจนการปรับแต่งค่าพารามิเตอร์เพื่อให้ได้ผลลัพธ์ที่ดีที่สุด เช่น Backbone Fine-tuning (การปรับแต่ง Backbone เช่น ResNet-๑๐๑ สำหรับข้อมูลเฉพาะ เช่น ข้อมูลดาวเทียม) และ Training Data Augmentation (เทคนิคเพิ่มข้อมูล เช่น การหมุนภาพหรือปรับแสง) เป็นต้น

๔) การฝึกอบรม (Training) ด้วยการใช้ชุดข้อมูลที่มีการ Annotate อย่างครบถ้วนสำหรับการฝึกอบรมโมเดล

๕) การทำนาย (Prediction) ด้วยการใช้โมเดลที่ผ่านการฝึกอบรมเพื่อจำแนกวัตถุในภาพถ่ายดาวเทียม

๖) การแสดงผล ด้วยการใช้ไลบรารีเช่น Matplotlib เพื่อแสดงผลภาพพร้อม Bounding Boxes

๓.๔ Convolutional Neural Network (CNN)

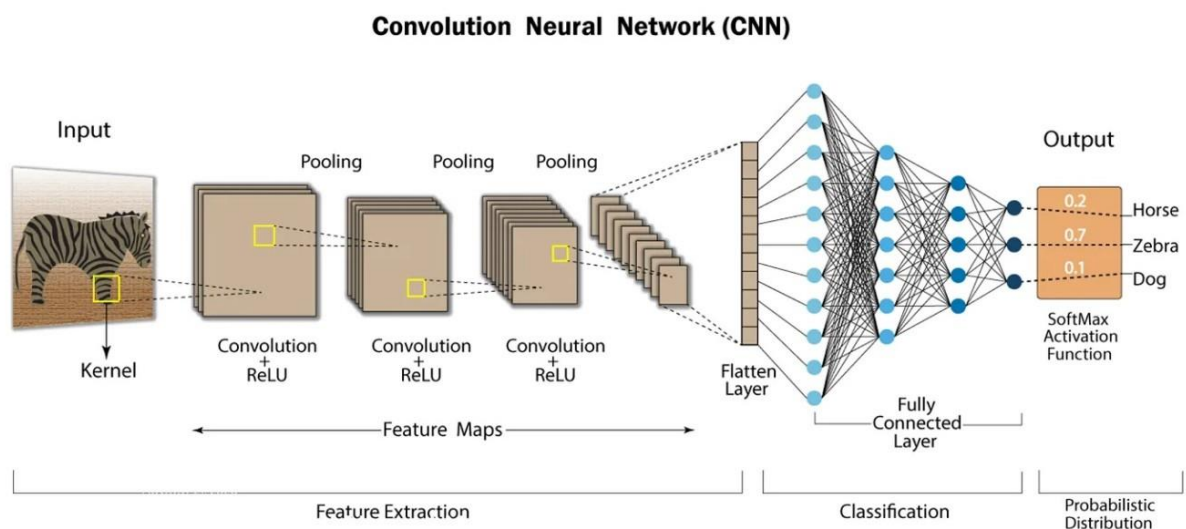
Convolutional Neural Network (CNN) หรือเครือข่ายประสาทเทียมแบบคอนโวลูชัน เป็นหนึ่งในเทคโนโลยีที่ขับเคลื่อนการปฏิบัติของปัญญาประดิษฐ์ โดยเฉพาะอย่างยิ่งในด้านการประมวลผลภาพ CNN ได้รับการพัฒนาขึ้นเพื่อเลียนแบบกระบวนการมองเห็นของมนุษย์ และสามารถเรียนรู้ลักษณะเฉพาะของภาพได้อย่างมีประสิทธิภาพ CNN เป็นแบบจำลองประเภทหนึ่งของ Deep Learning ที่มีโครงสร้างเฉพาะสำหรับการประมวลผลข้อมูลที่มีลำดับหรือโครงสร้างในรูปแบบของกริด (Grid) เช่น ภาพถ่ายจากกล้องถ่ายภาพ ภาพถ่ายจากดาวเทียม และภาพถ่ายจากโทรศัพท์มือถือ เป็นต้น โดย CNN ถูกออกแบบมาเพื่อตรวจจับลักษณะเฉพาะของภาพในรูปแบบต่าง ๆ ไม่ว่าจะเป็นการตรวจจับขอบ การตรวจจับรูปร่าง และการตรวจจับวัตถุ หลักการทำงานของ CNN สามารถอธิบายได้โดยสรุป ดังนี้

๑) Convolutional Layer ทำหน้าที่กรองคุณลักษณะ (Features) ต่าง ๆ ของภาพ โดยใช้ตัวกรอง (Filter) หรือเคอร์เนล (Kernel) เป็นเมทริกซ์ขนาดเล็ก เช่น ขนาด ๓x๓ หรือ ๕x๕ เป็นต้น ที่ถูกใช้งานเพื่อทำการ Convolution กับภาพหรือข้อมูลที่ป้อนเข้าไป ตัวกรองนี้จะเลื่อนผ่านภาพเพื่อตรวจจับคุณลักษณะเฉพาะโดยที่แต่ละตำแหน่งจะทำการคำนวณผลรวมแบบจุด (Dot Product) ระหว่างค่ากริดของข้อมูลกับค่าน้ำหนักของตัวกรอง การเลื่อนตัวกรองไปยังตำแหน่งถัดไปบนภาพ ค่าที่ใช้ในการเลื่อนนี้เรียกว่า Stride ซึ่งหาก stride เป็น ๑ ตัวกรองจะเลื่อนทีละกริด หาก stride เป็น ๒ ตัวกรองจะเลื่อนทีละ ๒ กริด ผลลัพธ์ที่ได้จากการเลื่อนตัวกรองผ่านข้อมูลจะสร้างเป็น Feature Map ที่แสดงคุณลักษณะเฉพาะที่ตัวกรองตรวจจับได้

๒) Activation Function เป็นฟังก์ชันทางคณิตศาสตร์ที่ใช้ในแต่ละนิวรอน (Neurons) ของเครือข่ายประสาทเทียม (Neural Network) รวมถึง CNN เพื่อเพิ่มความซับซ้อนและสามารถเรียนรู้ลักษณะและรูปแบบที่ซับซ้อนในข้อมูลได้ ฟังก์ชันนี้จะถูกนำไปใช้หลังจากการคำนวณ Dot Product ในแต่ละชั้นข้อมูล (Layer) ของเครือข่าย Activation Function จะช่วยเพิ่มความไม่เป็นเชิงเส้น (Non-Linearity) ของชุดข้อมูล เนื่องจากข้อมูลในโลกแห่งความเป็นจริงมีความซับซ้อนและอาจไม่ได้มีความสัมพันธ์เชิงเส้นต่อกัน ดังนั้นการใช้ Activation Function ที่ไม่เป็นเชิงเส้นช่วยให้โมเดลสามารถเรียนรู้ความสัมพันธ์ที่ซับซ้อนได้ โดยการควบคุมค่าผลลัพธ์ของนิวรอน (Neurons) ให้อยู่ในช่วงที่เหมาะสม เช่น ระหว่าง ๐ ถึง ๑ หรือ -๑ ถึง ๑ ซึ่งช่วยให้การฝึกฝนเครือข่ายมีความเสถียรมากขึ้น ประเภทของ Activation Function ที่ใช้บ่อยใน CNN ได้แก่ ReLU, Leaky ReLU, Sigmoid, Tanh และ Softmax ซึ่งแต่ละฟังก์ชันมีข้อดีและข้อเสียที่แตกต่างกันไปขึ้นอยู่กับการใช้งาน

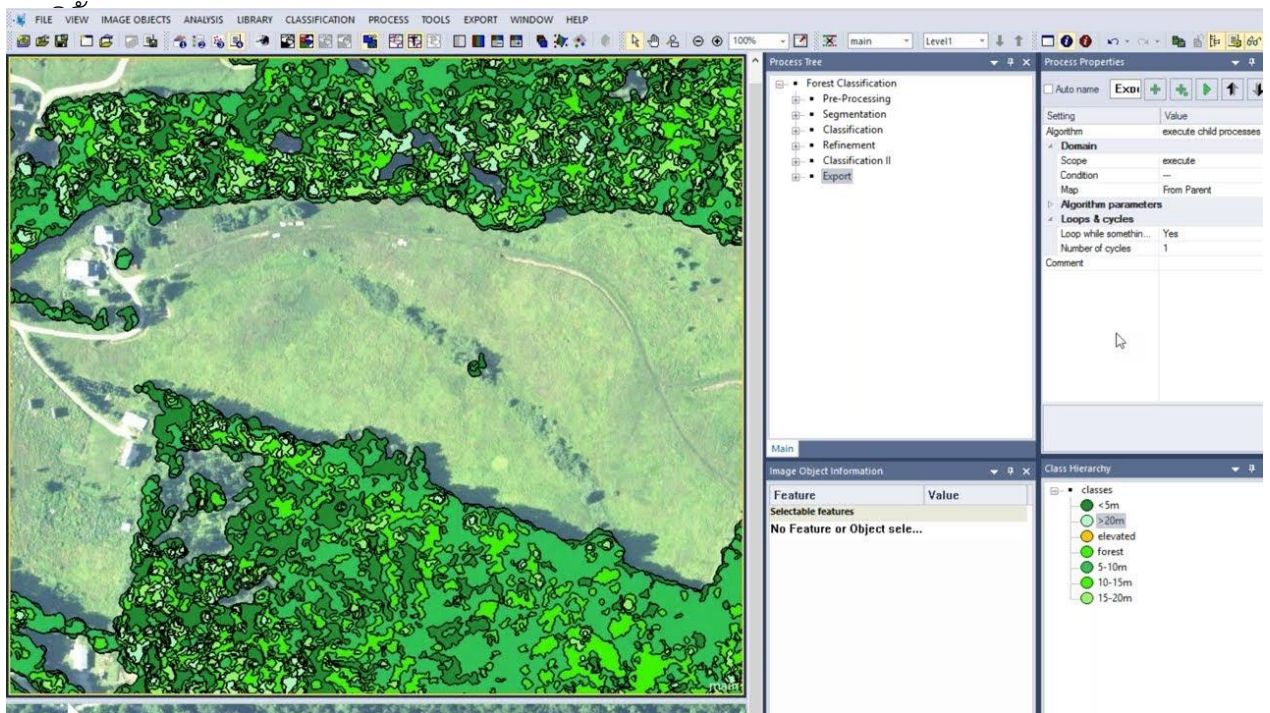
๓) Pooling Layer เป็นชั้นข้อมูล (Layer) ที่ใช้ใน CNN เพื่อทำการลดขนาดของแผนที่คุณลักษณะ (Feature Map) ที่ได้จากชั้นข้อมูล Convolution โดยการรวมค่าจากกลุ่มของพิกเซลใกล้เคียงเข้าด้วยกัน (Pooling) ซึ่งมีวัตถุประสงค์หลักเพื่อลดจำนวนพารามิเตอร์ และการคำนวณในเครือข่าย พร้อมทั้งรักษาข้อมูลที่สำคัญที่สุดและลดความเสี่ยงของการ Overfitting โดยการเลือกค่าสูงสุด (Max Pooling) หรือค่าเฉลี่ย (Average Pooling) จากพื้นที่ย่อยของภาพ

๔) Fully Connected Layer (FC Layer) ทำหน้าที่เชื่อมต่อทุกนิวรอน (Neurons) ในชั้นก่อนหน้ากับทุกนิวรอน (Neurons) ในชั้นถัดไป หลังจากที่ได้ข้อมูลผ่านชั้น Convolution และ Pooling โดยข้อมูลจะถูกปรับให้เป็น Vector หนึ่งมิติ (flattening) และถูกป้อนเข้าสู่ FC Layer จากนั้นผลลัพธ์ของแต่ละนิวรอน (Neurons) จะถูกคำนวณโดยการคูณค่าน้ำหนัก (Weights) กับค่าข้อมูลอินพุต (Input Data) แล้วรวมกับค่าคงที่ (Biases) จากนั้นค่าผลลัพธ์จะถูกส่งผ่าน Activation Function เพื่อเพิ่มความไม่เป็นเชิงเส้นระหว่างการฝึกฝนโมเดล น้ำหนักและค่าคงที่ ของ FC Layer จะถูกปรับปรุงผ่านกระบวนการ Backpropagation เพื่อลดความคลาดเคลื่อนของผลลัพธ์ (Error) การเชื่อมต่อแบบนี้ช่วยให้ข้อมูลสามารถถูกประมวลผลอย่างเต็มที่ และทำให้โมเดลสามารถเรียนรู้ความสัมพันธ์ที่ซับซ้อนได้ ทำให้สามารถรวบรวมลักษณะที่ตรงจวบได้และใช้สำหรับการจำแนกหรือการทำนาย FC Layer มักใช้ในส่วนท้ายของ CNN เพื่อทำการจำแนกประเภท (Classification) อย่างไรก็ตาม หากจำนวนพารามิเตอร์ที่ใส่เข้าไปในแบบจำลองมีจำนวนมาก อาจทำให้มีความเสี่ยงต่อการ Overfitting



ภาพแสดงโครงสร้างและขั้นตอนการทำงานของแบบจำลอง CNN (Shahriar N., ๒๐๑๙)

ในการดำเนินโครงการฯ จะใช้เครื่องมือ CNN จากโปรแกรมสำเร็จรูป eCognition Developer ซึ่งเป็นผลิตภัณฑ์จากบริษัท Trimble Inc. และจำเป็นต้องได้รับใบอนุญาต (License) ในการใช้งานโปรแกรม โดยเฉพาะชุดเครื่องมือ Deep Learning ที่เป็นเครื่องเฉพาะสำหรับผู้ใช้งานระดับสูง (Advance User) โดยสามารถเรียกใช้งานเครื่องมือด้วยการสร้าง Process Tree เพื่อการสร้าง Rule Set สำหรับการจำแนกเนื้อที่ยืนต้นปาล์มน้ำมัน จากการเลือกอัลกอริทึม (Algorithm) ในการแบ่งส่วนภาพ (Segmentation) จากคุณลักษณะเชิงคลื่น (Spectral Characteristics) ของภาพ จากนั้น จึงฝึกหัดแบบจำลองด้วยการกำหนดประเภทชั้นข้อมูล (Class Assignment)



ภาพตัวอย่างการจำแนกการใช้ประโยชน์ที่ดินด้วยโปรแกรม eCognition Developer

๔. ซอฟต์แวร์ (Software) ที่ใช้ในการดำเนินงาน

๔.๑ Sentinel Application Platform (SNAP)

โปรแกรม Sentinel Application Platform (SNAP) โดย ESA (European Space Agency) ถูกพัฒนาขึ้นเพื่อเป็นเครื่องมือในการประมวลผลและวิเคราะห์ข้อมูลจากดาวเทียม Sentinel ของโครงการ Copernicus รวมถึงดาวเทียมอื่นๆ SNAP มีวัตถุประสงค์หลักในการสนับสนุนการวิเคราะห์ภาพถ่ายดาวเทียมเพื่อการวิจัยและการใช้งานในเชิงพาณิชย์ สามารถประมวลผลภาพ ปรับปรุงคุณภาพของข้อมูล และวิเคราะห์ข้อมูลแบบหลายมิติ ซึ่งเหมาะสมกับงานหลากหลายประเภทตั้งแต่การติดตามสภาพแวดล้อม การศึกษาการเปลี่ยนแปลงของพื้นดิน ไปจนถึงการสำรวจการใช้ประโยชน์ที่ดิน

จุดเด่นของ SNAP คือความสามารถในการรองรับข้อมูลจากหลายแหล่ง การประมวลผลแบบโมดูลาร์ (Modular) ที่ช่วยให้ผู้ใช้สามารถปรับแต่งการวิเคราะห์ตามความต้องการ โดยใช้โมดูล (Module) หรือปลั๊กอิน (plugins) ต่าง ๆ และการใช้งานที่ไม่ซับซ้อนมาก ทำให้เหมาะสำหรับนักวิจัยและผู้ใช้งานทั่วไป อย่างไรก็ตาม ข้อจำกัดของ SNAP อาจอยู่ที่ความต้องการทรัพยากรคอมพิวเตอร์ที่มีประสิทธิภาพสูงเพื่อประมวลผลข้อมูลขนาดใหญ่ และการที่ผู้ใช้ใหม่อาจต้องใช้เวลาในการเรียนรู้การใช้งานที่ครบถ้วน โปรแกรม SNAP มีฟังก์ชันและเครื่องมือหลายอย่างที่น่าสนใจ โดยเฉพาะในการจำแนกการใช้ประโยชน์ที่ดิน เช่น Land

Use and Land Cover (LULC) Mapping ซึ่งเป็นชุดเครื่องมือสำหรับการสร้างแผนที่การใช้ประโยชน์ที่ดิน และการเปลี่ยนแปลงของพื้นที่ ซึ่งใช้ข้อมูลจากดาวเทียมหลายชนิดเพื่อให้ได้ผลลัพธ์ที่แม่นยำ และเครื่องมือ Change Detection ซึ่งเป็นเครื่องมือสำหรับการตรวจสอบและวิเคราะห์การเปลี่ยนแปลงของพื้นที่ดิน ในช่วงเวลาที่ต่างกัน ซึ่งสามารถใช้ในการติดตามการเปลี่ยนแปลงการใช้ประโยชน์ที่ดินได้อย่างมีประสิทธิภาพ เป็นต้น

๔.๒ ArcGIS Pro

โปรแกรม ArcGIS Pro ถูกพัฒนาโดยบริษัท Esri (Environmental Systems Research Institute) ซึ่งเป็นองค์กรชั้นนำในด้านเทคโนโลยีระบบสารสนเทศภูมิศาสตร์ (GIS) วัตถุประสงค์หลักของ ArcGIS Pro คือเพื่อให้ผู้ใช้สามารถวิเคราะห์ข้อมูลภูมิศาสตร์ ทำแผนที่ และสร้างการนำเสนอข้อมูลเชิงพื้นที่ได้อย่างมีประสิทธิภาพและง่ายดาย โปรแกรมนี้ถูกออกแบบมาให้ใช้งานในหลายภาคส่วน เช่น รัฐบาล การศึกษา การจัดการทรัพยากรธรรมชาติ และการพัฒนาเมือง เพื่อช่วยในการตัดสินใจ และวางแผนบนพื้นฐานของข้อมูลเชิงพื้นที่ จุดเด่นของ ArcGIS Pro ได้แก่ การประมวลผลข้อมูลเชิงภูมิศาสตร์ที่รวดเร็วและมีประสิทธิภาพ การสนับสนุนการทำงานแบบ ๓ มิติ และการเชื่อมต่อกับข้อมูลออนไลน์ได้อย่างราบรื่น เช่น ฐานข้อมูล ArcGIS Online นอกจากนี้ ArcGIS Pro ยังมีอินเทอร์เฟซ (User Interface) ที่ใช้งานง่ายและเป็นมิตรกับผู้ใช้ แต่จุดด้อยของโปรแกรมคือมีค่าใช้จ่ายสูงสำหรับการใช้งานเชิงพาณิชย์และมีการเรียนรู้การใช้งานที่ซับซ้อนสำหรับผู้เริ่มต้น นอกจากนี้ยังต้องการทรัพยากรคอมพิวเตอร์ที่มีประสิทธิภาพสูง เพื่อให้การทำงานกับข้อมูลขนาดใหญ่ราบรื่นและอาจต้องการอินเทอร์เน็ตความเร็วสูงเมื่อต้องการเชื่อมต่อกับ ArcGIS Online

ArcGIS Pro มีความสามารถในการจัดการและประมวลผลกับข้อมูลที่หลากหลายทั้งแบบ Vector และ Raster และมีเครื่องมือเฉพาะที่จำเป็นต่อการจำแนกการใช้ประโยชน์ที่ดิน ดังนี้

๑) Spatial Analyst เป็นชุดเครื่องมือสำหรับการวิเคราะห์เชิงพื้นที่ เช่น การสร้างแผนที่ความสูง แผนที่การไหลของน้ำ และการวิเคราะห์พื้นที่ที่เหมาะสมสำหรับการใช้ประโยชน์ที่ดิน เป็นต้น

๒) Image Classification เป็นชุดเครื่องมือสำหรับการจำแนกประเภทของพื้นดินและวัตถุในภาพถ่ายดาวเทียมโดยใช้วิธีการทางสถิติและการเรียนรู้ของเครื่อง เช่น Maximum Likelihood, Random Forest และการจำแนกประเภทแบบอัตโนมัติ (Automatic Classification) เป็นต้น

๓) ๓D Analyst เป็นเครื่องมือสำหรับการสร้างและวิเคราะห์ข้อมูล ๓ มิติ เช่น การสร้างแบบจำลองพื้นที่ ๓ มิติ การวิเคราะห์ทัศนวิสัย และการวิเคราะห์การเปลี่ยนแปลงเชิงพื้นที่ในมิติที่สาม เป็นต้น

๔) Network Analyst เป็นเครื่องมือสำหรับการวิเคราะห์เครือข่าย เช่น การวิเคราะห์เส้นทางที่ดีที่สุด การวิเคราะห์เครือข่ายการขนส่ง และการจำแนกพื้นที่ที่มีการเข้าถึงบริการต่าง ๆ เป็นต้น

๕) ModelBuilder เป็นเครื่องมือสำหรับการสร้างแบบจำลองการประมวลผลข้อมูลแบบอัตโนมัติ ซึ่งช่วยให้ผู้ใช้งานสามารถสร้างกระบวนการวิเคราะห์ที่ซับซ้อนได้ง่ายขึ้น

๔.๓ eCognition Developer

โปรแกรม eCognition Developer ถูกพัฒนาโดย Trimble Inc. ซึ่งเป็นบริษัทชั้นนำในด้านเทคโนโลยีการจัดการข้อมูลภูมิศาสตร์และการสำรวจ วัตถุประสงค์หลักของ eCognition Developer คือการวิเคราะห์ภาพถ่ายดาวเทียมและข้อมูลภูมิศาสตร์โดยใช้เทคโนโลยีการวิเคราะห์เชิงวัตถุ (Object-Based Image Analysis: OBIA) โปรแกรมนี้ช่วยให้ผู้ใช้สามารถจำแนกประเภท ตรวจสอบ และวิเคราะห์ข้อมูลภูมิศาสตร์ได้อย่างมีประสิทธิภาพ โดยเฉพาะในการทำแผนที่การใช้ประโยชน์ที่ดิน การวิเคราะห์การเปลี่ยนแปลงของพื้นที่

และการประเมินสภาพแวดล้อม จุดเด่นของ eCognition Developer คือการใช้เทคโนโลยี OBIA ที่ช่วยให้การจำแนกประเภทของข้อมูลเชิงพื้นที่ที่มีความแม่นยำสูง และสามารถจัดการกับข้อมูลที่มีความละเอียดสูงได้อย่างมีประสิทธิภาพ รวมถึงความสามารถในการจัดการ และวิเคราะห์ข้อมูลที่มีหลายมิติ (Multi-Dimensional) เช่น ข้อมูลภาพถ่ายจากหลายแหล่งและหลายช่วงเวลา โปรแกรมนี้ยังมีความยืดหยุ่นสูงในการสร้างและปรับแต่งกระบวนการวิเคราะห์อย่างละเอียด อย่างไรก็ตาม ข้อจำกัดของ eCognition Developer อาจรวมถึงความซับซ้อนในการใช้งาน ที่อาจทำให้ผู้ใช้ใหม่ต้องใช้เวลาเรียนรู้ นอกจากนี้ โปรแกรมนี้ยังมีค่าใช้จ่ายที่ค่อนข้างสูง ซึ่งอาจเป็นอุปสรรคสำหรับผู้ใช้งานทั่วไป ฟังก์ชันที่น่าสนใจใน eCognition Developer โดยเฉพาะต่อการจำแนกการใช้ประโยชน์ที่ดิน ได้แก่

๑) Segmentation เป็นเครื่องมือสำหรับการแบ่งส่วนภาพ (Image Segmentation) ที่ช่วยในการแยกภาพออกเป็นวัตถุหรือพื้นที่ที่มีลักษณะคล้ายกัน ซึ่งทำให้สามารถจำแนกประเภทได้อย่างแม่นยำและละเอียด

๒) Rule-Based Classification เป็นชุดเครื่องมือสำหรับการสร้างกฎเกณฑ์ที่ใช้ในการจำแนกประเภทของข้อมูลเชิงภูมิศาสตร์ เช่น การกำหนดกฎที่อิงจากลักษณะของวัตถุในภาพ เช่น สี รูปร่าง และขนาด เป็นต้น

๓) Object-Based Analysis เป็นเครื่องมือสำหรับการวิเคราะห์ข้อมูลเชิงวัตถุที่ช่วยให้การจำแนกประเภทและการวิเคราะห์เชิงพื้นที่ที่สามารถทำได้ตามลักษณะของวัตถุในภาพ ทำให้การวิเคราะห์การใช้ประโยชน์ที่ดินมีความแม่นยำและละเอียดมากขึ้น

๔) Change Detection เป็นฟังก์ชันสำหรับการตรวจสอบและวิเคราะห์การเปลี่ยนแปลงของพื้นที่ในช่วงเวลาที่แตกต่างกัน เช่น การตรวจสอบการเปลี่ยนแปลงการใช้ประโยชน์ที่ดินหรือการเปลี่ยนแปลงในสิ่งแวดล้อม เป็นต้น

๕) Integration with GIS เป็นความสามารถในการเชื่อมต่อและนำเข้าข้อมูลจากระบบ GIS อื่น ๆ รวมถึงการส่งออกข้อมูลที่วิเคราะห์แล้วไปยังโปรแกรม GIS เพื่อการทำงานร่วมกันอย่างราบรื่น

๕. ขั้นตอนการดำเนินงาน

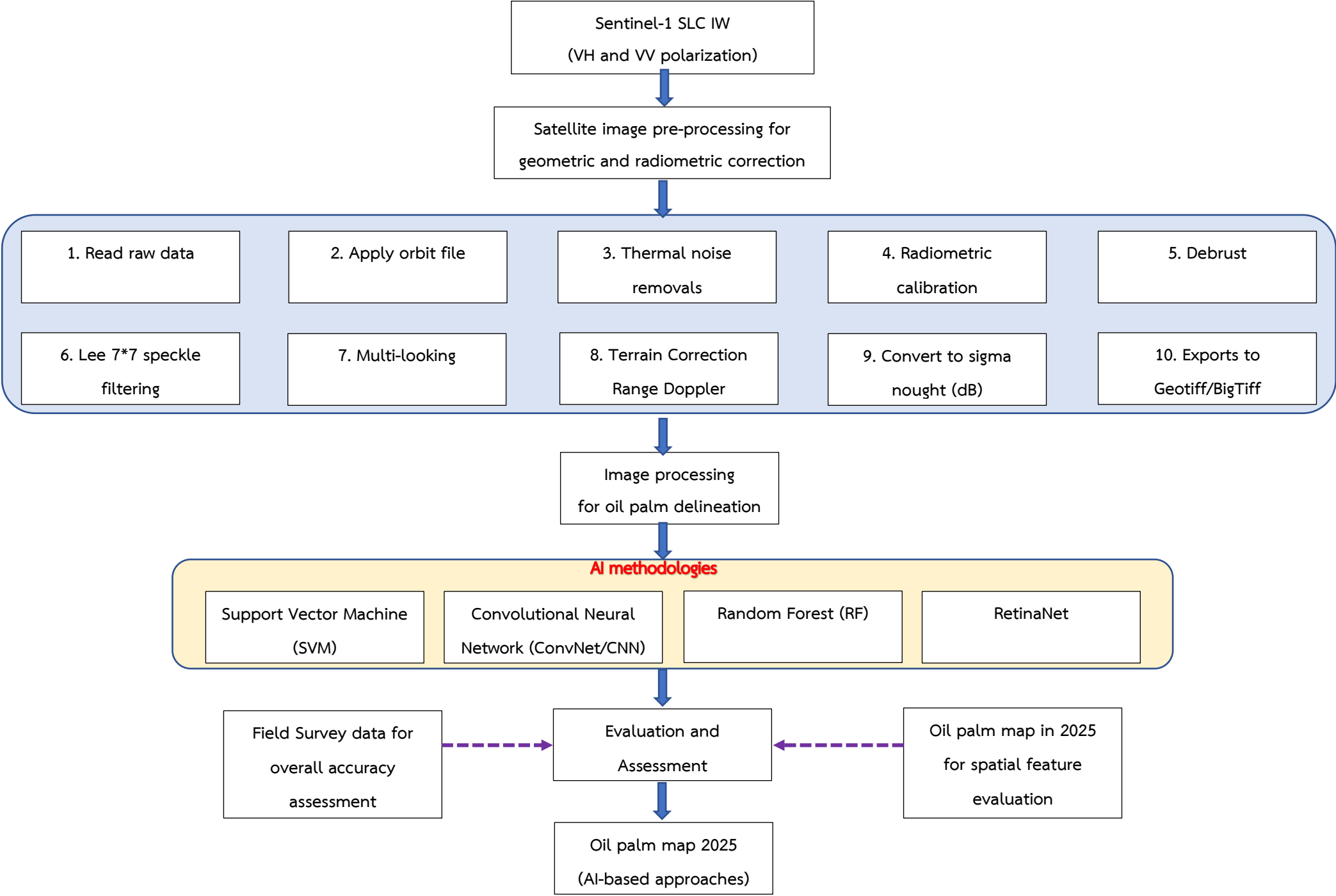
เมื่อภาพถ่ายดาวเทียม Sentinel-๑ ระดับ Single Look Complex โหมด Wide Swath (IW) ถูกดาวน์โหลดจากผู้ให้บริการแล้ว จะถูกนำมาเข้าสู่กระบวนการการปรับแก้ภาพเชิงเรขาคณิต (Geometric Correction) และเชิงคลื่นรังสี (Radiometric Correction) ๑๐ ขั้นตอน จากนั้นจะเข้าสู่กระบวนการประมวลผลภาพด้วยการฝึกหัด (Training) แบบจำลองทั้ง ๔ วิธี ประกอบด้วย RF, SVM, RetinaNet และ CNN จากชุดข้อมูลตัวอย่างที่ได้จากการสำรวจภาคสนาม เมื่อได้ผลการจำแนกเบื้องต้นปาล์มน้ำมันและการใช้ประโยชน์ที่ดินประเภทอื่นแล้ว จึงทำการประเมินผลการจำแนกด้วยการตรวจสอบความถูกต้องทั้งหมด (Overall Accuracy: OA) ด้วยจุดตัวอย่างจากการสำรวจภาคสนาม โดยมีเงื่อนไขว่า ผลการจำแนกที่ได้จะต้องมีความถูกต้องของผู้ผลิต (Producer's Accuracy) และความถูกต้องของผู้ใช้งาน (User's Accuracy) มากกว่าร้อยละ ๔๐ และ ๖๐ ตามลำดับ หากไม่เป็นไปตามเงื่อนไขที่กำหนดไว้ จะต้องมีการปรับแก้ค่าพารามิเตอร์ต่าง ๆ ในแบบจำลองและชุดข้อมูลตัวอย่าง เพื่อให้ได้ผลลัพธ์ตามเกณฑ์ที่กำหนดไว้ นอกจากนี้ ยังมีการเปรียบเทียบผลลัพธ์กับฐานข้อมูลเนื้อที่ยืนต้นปาล์มน้ำมัน ปี ๒๕๖๗ ที่จัดทำโดยศูนย์สารสนเทศการเกษตร เพื่อประเมินคุณลักษณะเชิงพื้นที่ของผลลัพธ์

ในกระบวนการแปลวิเคราะห์ภาพถ่ายจากดาวเทียม Sentinel-๑ ร่วมกับเทคโนโลยี AI นั้น มีความจำเป็นอย่างยิ่งที่จะต้องมีการปรับแก้ข้อมูลต่างๆ (Pre-Processing) ก่อนเข้าสู่กระบวนการวิเคราะห์ เพื่อให้ได้ข้อมูลที่

มีคุณภาพก่อนการวิเคราะห์ในขั้นตอนต่อไป ประกอบด้วย ขั้นตอนการ Read Raw Data เป็นการอ่านข้อมูล และ Metadata ในการบินถ่ายภาพในระวาง และในวันดังกล่าว, การใส่รายละเอียดไฟล์ของวงโคจร (Apply Orbit File), การขจัดสิ่งปนเปื้อนที่อยู่ในย่านความยาวคลื่น Thermal Infrared (Thermal Noise Reduction), การ Deburst (เป็นการรวมภาพ ๓ Bursts เข้าด้วยกันให้เป็นภาพเดียว), การปรับแก้เชิงรังสี (Radiometric Calibration) และการทำการปรับแก้เนื่องมาจากลักษณะภูมิประเทศ (Terrain Correction) โดยเป็นการใส่ค่าความสูงให้กับข้อมูลภาพ ซึ่งโดยปกตินิยมใช้ Shuttle Radar Topography Mission (SRTM) ที่ National Aeronautics and Space Administration (NASA) เป็นผู้จัดทำ สามารถแสดงขั้นตอนการจัดทำ ข้อมูลเนื้อที่ยินตันปาล์มน้ำมันได้ดังภาพด้านล่าง



แผนภูมิแสดงขั้นตอนการจำแนกปาล์มน้ำมันด้วยเทคโนโลยีปัญญาประดิษฐ์ (Artificial Intelligence: AI) ปีงบประมาณ พ.ศ. ๒๕๖๘



๖. การประเมินผลการจำแนกเนื้อที่ยืนยันต้นปาล์มน้ำมัน

๖.๑ ค่าความถูกต้องทั้งหมด (Overall Accuracy: OA)

ค่าความถูกต้องทั้งหมด (Overall Accuracy: OA) เป็นค่าที่ใช้วัดความถูกต้องของแบบจำลองในการจำแนกประเภท (Classification) ของกริดที่อยู่บนภาพถ่ายดาวเทียม โดยคำนวณจากจำนวนของการจำแนกประเภทที่ถูกต้องทั้งหมดหารด้วยจำนวนของตัวอย่างทั้งหมดในชุดข้อมูลทดสอบ ค่า OA เป็นแนวคิดที่มีมานานในสาขาการวิเคราะห์ข้อมูล การจำแนกประเภท และการเรียนรู้ของเครื่อง (Machine Learning) จึงเป็นหนึ่งในมาตรฐานทั่วไปที่ใช้ประเมินคุณภาพของโมเดลการจำแนก การคำนวณค่า OA สามารถคำนวณได้จากสมการด้านล่าง

$$OA = \frac{\sum_{i=1}^k n_{ii}}{N}$$

โดย k คือจำนวนประเภททั้งหมด (classes)

n_{ii} คือจำนวนของตัวอย่างที่ถูกจำแนกถูกต้องในประเภทที่ i

N คือจำนวนของตัวอย่างทั้งหมด

ข้อดีของการคำนวณค่า OA คือ มีสมการที่ง่ายต่อการคำนวณ สามารถประเมินประสิทธิภาพทั่วไปของแบบจำลอง สามารถอธิบายปรากฏการณ์ที่เกิดขึ้นในการจำแนกประเภทได้โดยง่าย อย่างไรก็ตาม ข้อจำกัดที่สำคัญของค่า OA คือ เป็นค่าที่ไม่ได้คำนึงถึงการกระจายของข้อมูล ในกรณีที่มีการกระจายของประเภทข้อมูลที่ต้องการจำแนกไม่เท่ากัน ค่า OA อาจทำให้เกิดการประเมินที่ผิดพลาดได้ นอกจากนี้ ค่า OA ยังไม่สะท้อนความผิดพลาดที่สำคัญ เนื่องจากไม่สามารถบอกถึงรายละเอียด ของการทำนายผิดที่สำคัญในบางประเภทข้อมูลได้ ดังนั้นการประเมินประสิทธิภาพของแบบจำลอง การจำแนกการใช้ประโยชน์ที่ดิน ค่า OA สามารถใช้เป็นตัวชี้วัดเบื้องต้นถึงประสิทธิภาพของแบบจำลองการจำแนกในภาพรวม แต่หากต้องการวิเคราะห์ที่ละเอียดขึ้น อาจต้องพิจารณาค่าเฉพาะเจาะจงเพิ่มเติม เช่น Producer's Accuracy, User's Accuracy, และ Kappa Statistic ซึ่งจะช่วยให้เห็นภาพชัดเจนยิ่งขึ้นถึงความสามารถในการจำแนกประเภทที่ดินต่างๆ และความผิดพลาดที่อาจเกิดขึ้นในแต่ละประเภท

๖.๒ ความถูกต้องของผู้ผลิต (Producer's Accuracy)

ความถูกต้องของผู้ผลิต (Producer's Accuracy) หรือ Omission error หมายถึง ความแม่นยำในการจัดกลุ่มหรือจำแนกประเภทของพื้นที่ใดพื้นที่หนึ่งที่ผู้ผลิตหรือผู้สร้างแผนที่คาดหวังให้เป็นไปตามความจริง ซึ่งสามารถคำนวณได้จากจำนวนของตัวอย่างที่ถูกจัดประเภทอย่างถูกต้องในกลุ่มนั้นๆ เทียบกับจำนวนตัวอย่างทั้งหมดที่ควรจะเป็นของกลุ่มนั้นในแผนที่อ้างอิง (Reference Data)

$$Producer's Accuracy = \frac{\text{จำนวนตัวอย่างที่จำแนกได้ถูกต้องในกลุ่มนั้น}}{\text{จำนวนตัวอย่างทั้งหมดที่อยู่ในกลุ่มนั้นในข้อมูลอ้างอิง}} \times 100\%$$

ความถูกต้องของผู้ผลิตเป็นตัวชี้วัดที่สำคัญ เนื่องจากแสดงถึงความสามารถของกระบวนการจัดประเภท ในการแยกแยะ หรือระบุพื้นที่ตามกลุ่มที่กำหนดไว้อย่างถูกต้อง โดยจะช่วยให้ผู้ผลิตสามารถระบุ

ได้ว่าประเภทของพื้นที่ใดมีปัญหาหรือความคลาดเคลื่อนในการจำแนก และสามารถนำข้อมูลนี้ไปปรับปรุงกระบวนการหรือโมเดลการจำแนกให้มีความแม่นยำยิ่งขึ้น

๖.๓ ความถูกต้องของผู้ใช้ (User's Accuracy)

ความถูกต้องของผู้ใช้ (User's Accuracy) หรือ Commission error หมายถึง ความแม่นยำของข้อมูล ที่ได้รับการจัดประเภทจากมุมมองของผู้ใช้ข้อมูล ซึ่งแสดงถึงสัดส่วนของตัวอย่างที่ถูกจัดประเภทในกลุ่มใดกลุ่มหนึ่งที่มีความถูกต้องเมื่อเปรียบเทียบกับข้อมูลอ้างอิง ความถูกต้องของผู้ใช้จึงสะท้อนถึงโอกาสที่พื้นที่ใดพื้นที่หนึ่งที่ถูกจัดประเภทให้เป็นประเภทหนึ่ง ๆ นั้น จะเป็นประเภทที่ถูกต้องตามข้อมูลอ้างอิง (Reference Data)

$$User's Accuracy = \frac{\text{จำนวนตัวอย่างที่จำแนกได้ถูกต้องในกลุ่มนั้น}}{\text{จำนวนตัวอย่างทั้งหมดที่จัดให้เป็นกลุ่มนั้น}} \times 100\%$$

ความถูกต้องของผู้ใช้เป็นตัวชี้วัด ที่สำคัญสำหรับผู้ใช้อ้างอิงข้อมูลในการประเมินคุณภาพของการจัดประเภทข้อมูลที่ได้รับจากกระบวนการวิเคราะห์ เช่น การสร้างแผนที่ที่ดินหรือแผนที่การใช้ประโยชน์ที่ดิน เนื่องจากความถูกต้องของผู้ใช้ สะท้อนถึงโอกาสที่ข้อมูลประเภทใดประเภทหนึ่งที่ถูกระบุในแผนที่หรือข้อมูลจำแนก จะสอดคล้องกับสถานะจริงในพื้นที่นั้น ๆ หากความถูกต้องของผู้ใช้มีค่าสูง หมายความว่า ผู้ใช้สามารถมั่นใจได้ว่าข้อมูลที่ได้รับการจัดประเภทนั้นมีความแม่นยำและเชื่อถือได้ ในทางกลับกัน หากค่าความถูกต้องของผู้ใช้ต่ำ ผู้ใช้ข้อมูลอาจต้องระมัดระวังในการตัดสินใจหรือใช้ข้อมูลดังกล่าว

๖.๔ ค่า Kappa (Kappa Statistics)

ค่า Kappa (Kappa Statistics) หรือที่เรียกกันว่า Cohen's Kappa เป็นสถิติที่ใช้ในการวัดระดับความเห็นพ้องกัน (Agreement) ระหว่างตัวแปรสองตัวหรือมากกว่า ที่มักใช้ในการประเมินความแม่นยำของการจัดประเภทข้อมูล เช่น การวิเคราะห์ภาพถ่ายดาวเทียม หรือการจัดกลุ่มข้อมูลในระบบสารสนเทศภูมิศาสตร์ (GIS) โดยค่านี้จะคำนวณเพื่อวัดความแม่นยำของการจำแนกประเภท ซึ่งแก้ไขจากการเห็นพ้องโดยบังเอิญ ค่า Kappa ถูกพัฒนาขึ้นเพื่อแก้ไขข้อจำกัดของความถูกต้องทั้งหมด (Overall Accuracy: OA) ซึ่งอาจเกิดจากการเห็นพ้องกันโดยบังเอิญ ค่า Kappa จะคำนึงถึงทั้งความถูกต้องจริงและความน่าจะเป็นของการเห็นพ้องกันโดยบังเอิญ (Random Agreement) ทำให้เป็นการวัดความแม่นยำที่มีความซับซ้อนและน่าเชื่อถือมากยิ่งขึ้น

$$\kappa = \frac{P_o - P_e}{1 - P_e}$$

โดย P_o = ความถูกต้องทั้งหมด หรือสัดส่วนของการจำแนกประเภทที่ถูกต้องทั้งหมด

P_e = ความน่าจะเป็นของการเห็นพ้องกันโดยบังเอิญ ซึ่งคำนวณจากข้อมูลของการจำแนกประเภทที่ได้



ตารางการตีความค่า Kappa เพื่อแสดงการเห็นพ้องกันของตัวแปร

ค่า Kappa	การตีความ
$K < 0$	แสดงถึงการเห็นพ้องกันที่น้อยกว่าการคาดเดาโดยบังเอิญ
$0 \leq K < 0.20$	การเห็นพ้องกันเล็กน้อย
$0.20 \leq K < 0.40$	การเห็นพ้องกันน้อย
$0.40 \leq K < 0.60$	การเห็นพ้องกันปานกลาง
$0.60 \leq K < 0.80$	การเห็นพ้องกันมาก
$0.80 \leq K \leq 1.00$	การเห็นพ้องกันอย่างยอดเยี่ยม

ค่า Kappa เป็นตัวชี้วัดที่สำคัญ เนื่องจากช่วยให้เข้าใจถึงคุณภาพของการจำแนกประเภทข้อมูล โดยไม่เพียงแต่พิจารณาความถูกต้องทั้งหมด แต่ยังรวมถึงการปรับปรุงจากการเห็นพ้องกันที่อาจเกิดขึ้นโดยบังเอิญ ทำให้ค่า Kappa เป็นเครื่องมือที่มีประสิทธิภาพสำหรับการประเมินคุณภาพของกระบวนการวิเคราะห์และจำแนกข้อมูลในหลายสาขา เช่น ภูมิศาสตร์ การแพทย์ และสังคมศาสตร์



บรรณานุกรม

- Lin, T.-Y., Goyal, P., Girshick, R., He, K., & Dollár, P. (๒๐๑๗) 'Focal loss for dense object detection', Proceedings of the IEEE International Conference on Computer Vision (ICCV), pp. ๒๙๘๐-๒๙๘๘. Available at: https://openaccess.thecvf.com/content_ICCV_๒๐๑๗/papers/Lin_Focal_Loss_for_ICCV_๒๐๑๗_paper.pdf (Accessed: ๑๘ December ๒๐๒๓).
- Shahriar, N. (๒๐๑๙). What is Convolutional Neural Network (CNN) in Deep Learning? [online] LinkedIn. Available at: <https://www.linkedin.com/pulse/what-convolutional-neural-network-cnn-deep-learning-nafiz-shahriar/> [Accessed ๒๗ July ๒๐๒๔].
- Singh, Gurkirt & Akrigg, Stephen & Di Maio, Manuele & Fontana, Valentina & Javanmard Alitappeh, Reza & Saha, Suman & Jeddisaravi, Kossar & Yousefi, Farzad & Culley, Jacob & Nicholson, Tom & Omokeowa, Jordan & Khan, Salman & Grazioso, Stanislao & Bradley, Andrew & Di Gironimo, Giuseppe & Cuzzolin, Fabio. (๒๐๒๑). ROAD: The ROad event Awareness Dataset for Autonomous Driving. ๑๐.๔๘๕๕๕๐/arXiv.๒๐๒๑.๑๑๕๕๕.





ส่วนภูมิสารสนเทศการเกษตร
ศูนย์สารสนเทศการเกษตร
สำนักงานเศรษฐกิจการเกษตร

เบอร์โทรศัพท์ 0 2940 7030 ต่อ 7225
E-mail: cai-gis@oae.go.th