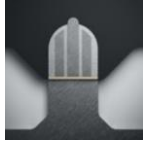


IMLA – A machine learning platform to democratize ML modeling in ML-APC

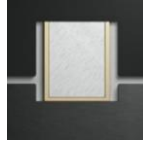
Duo(Jason) Zhang/Intel Corporation
Runshen Xu/Intel Corporation
Ali Ahmadi/Intel Corporation
Gobind Bisht/Intel Corporation
Rajan Sawhney/Intel Corporation

Why ML-APC

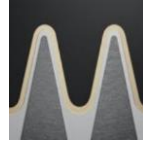
Manufacturing Roadmap



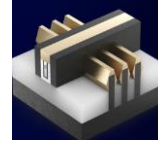
Strained Silicon



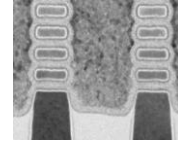
Hi-K Metal Gate



First FinFET

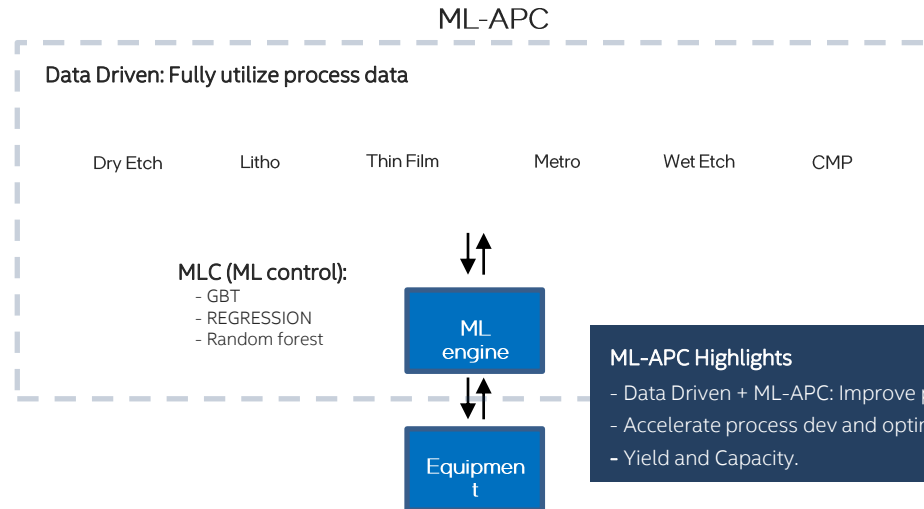
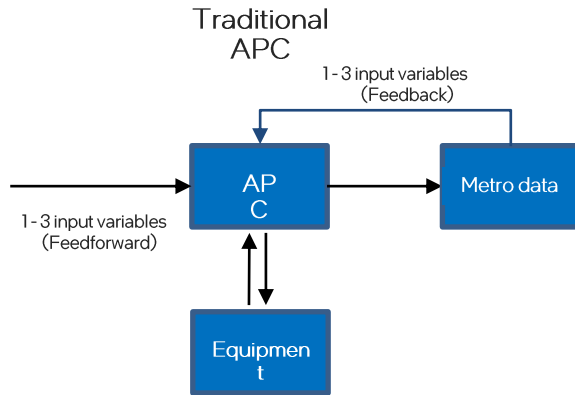


SuperFin

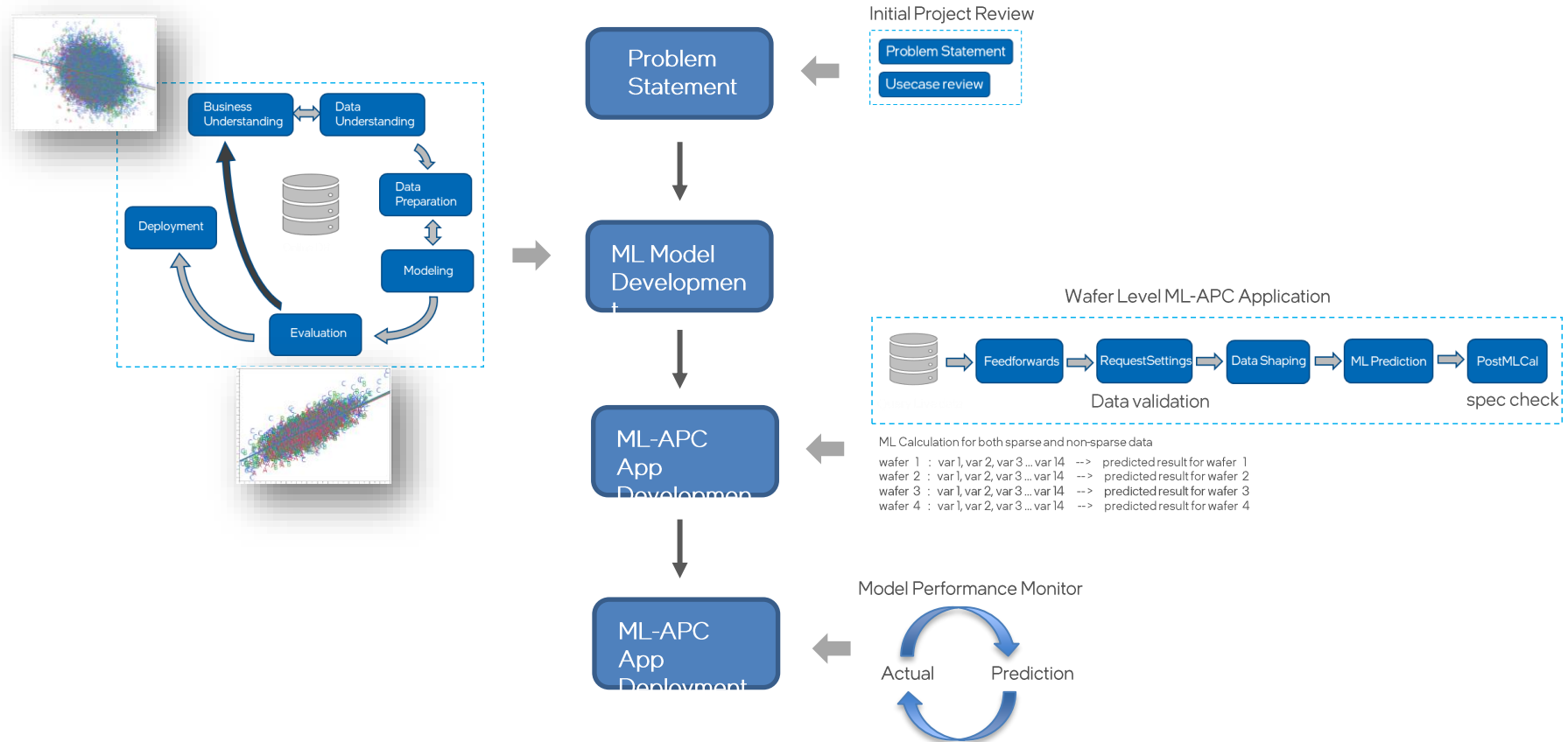


RibbonFET

Leading Edge Process : smaller process margin



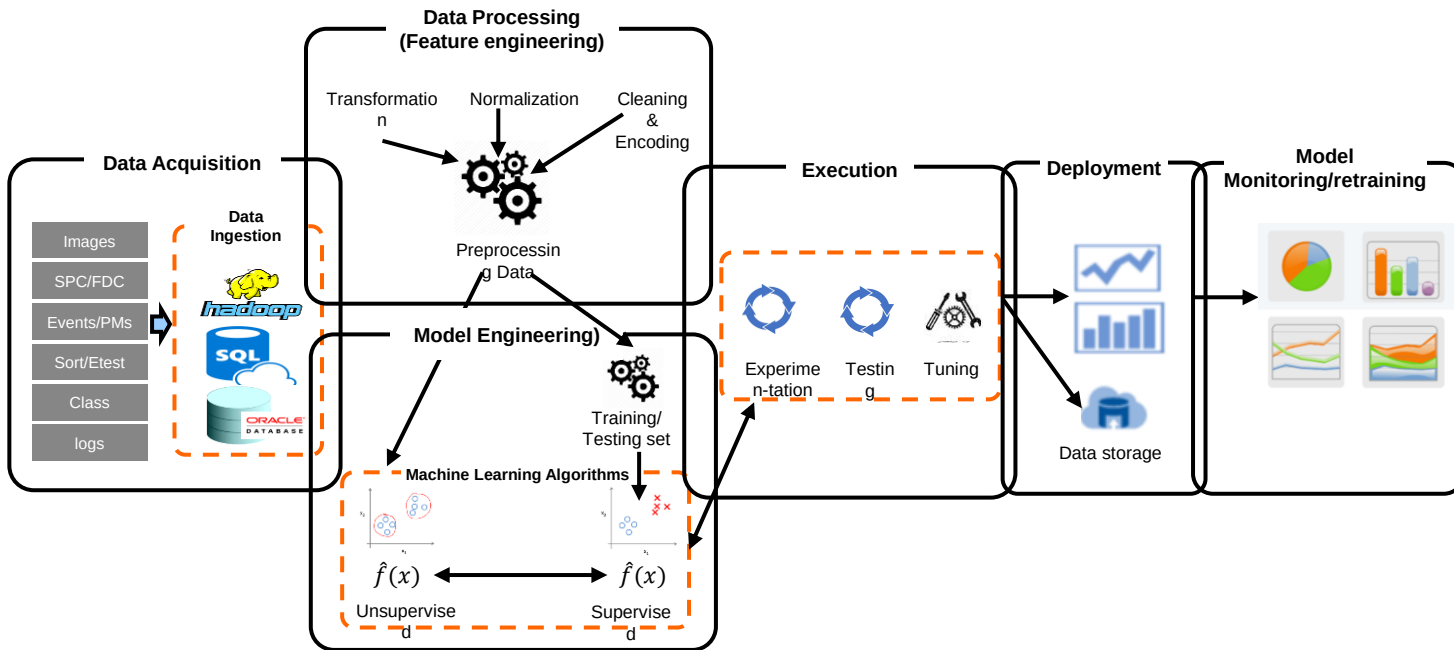
Wafer Level ML-APC Workflow



IMLA: Custom MLOps Platform

- **What**: a set of practices that aims to ***develop, deploy*** and ***maintain*** machine learning models in production *reliably* and *efficiently*
- **Why**: we are seeing a huge surge in using machine learning for decision-making, and business improvement. MLOps means **faster go-to-market** times and **lower operational costs**.
- **How**: an **in-house web-based** solution that automates and streamlines ML workflow and allows data scientists and non-experts to develop and deploy ML models
 - Integrated with Intel POR data sources:
 - Easy integration for custom data pre-processing algorithm
 - Easy integration with both Intel In-House and Open Source ML Algorithm
 - Integrated to Intel online and offline systems

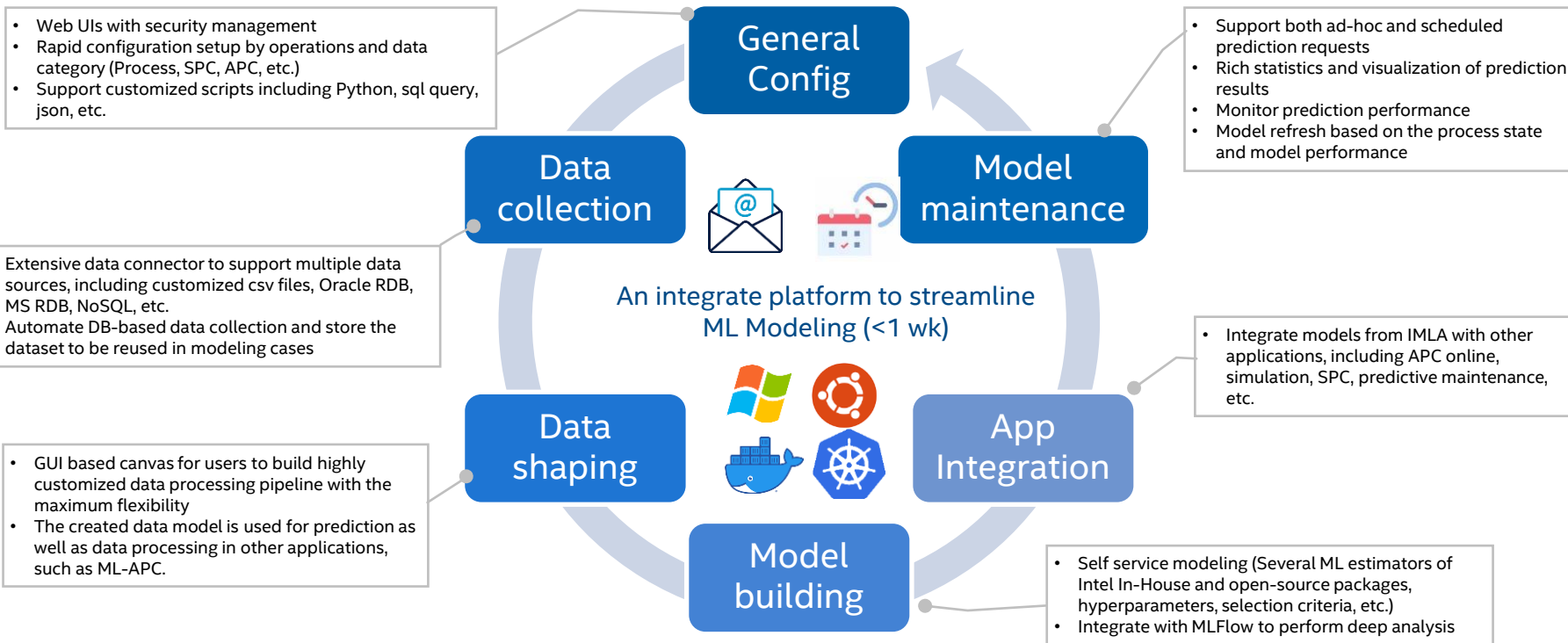
Machine Learning Lifecycle



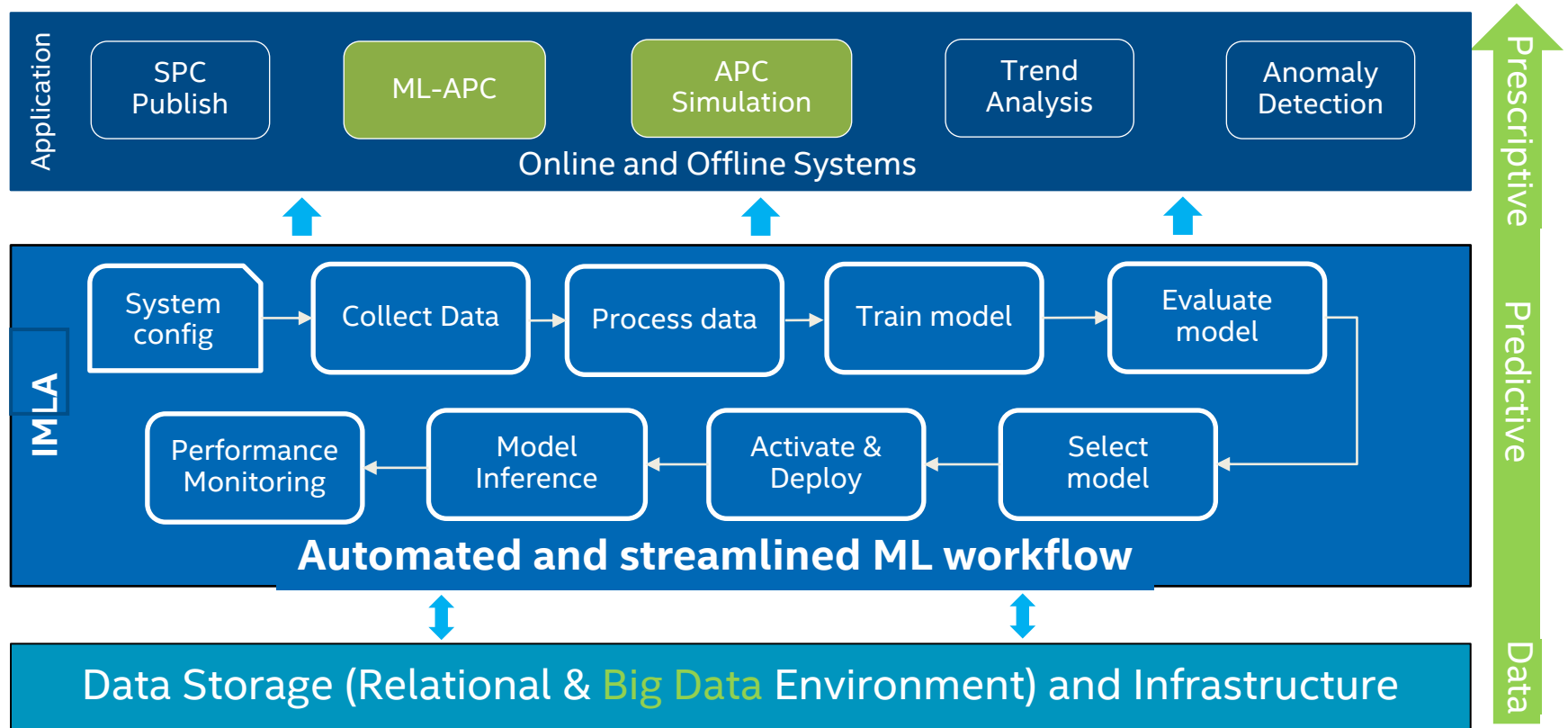
Machine Learning Lifecycle

Task (Proportion of Effort)	Subtasks
Problem understanding (5% - 10%)	a) Determine the objective ; b) Define success criteria c) Assess constrains
Data understanding (10% - 25%)	a) Assess data situation; b) Obtain data access c) Explore data
Data preparation (20% - 40%)	a) Clean data ; b) Feature engineering
Modeling (15% - 20%)	a) Select model approach; b) Build models c) Tune models
Evaluations of Results (5% - 10%)	a) Select Model ; b) Validate Model c) Explain Model
Deployment (5% - 10%)	a) Deploy Model; b) Maintain and Monitor
Model monitoring and retraining (10% - 15%)	a) Monitor model and send notification; b)re-train outdated models

IMLA – Highly Extensible and Horizontally Scalable

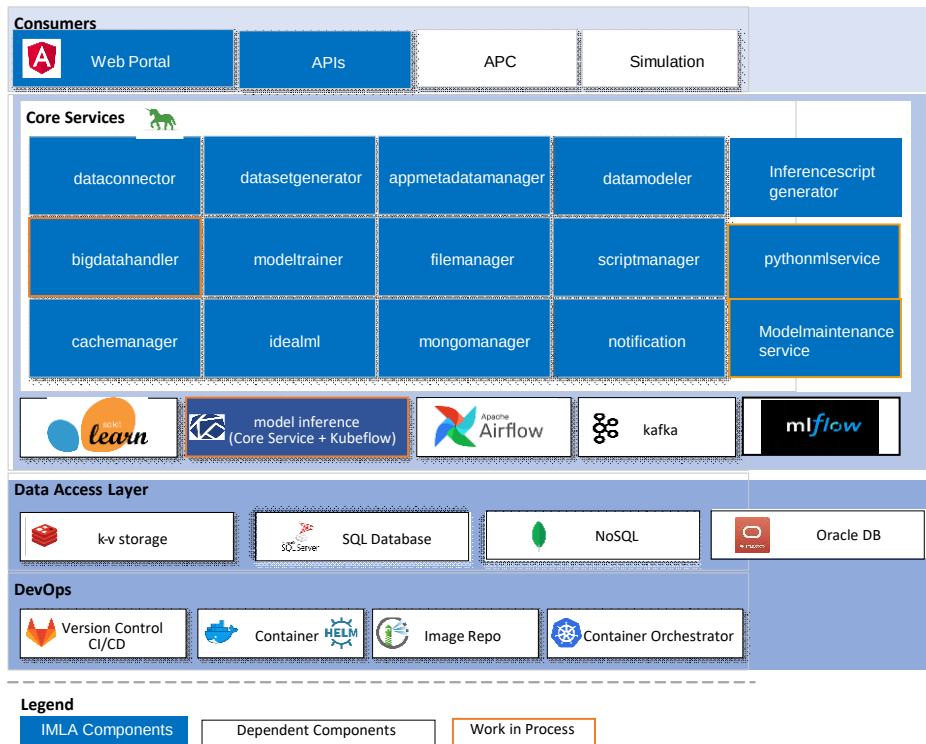


IMLA Architecture



IMLA Architecture

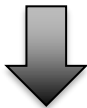
- Supports Intel In-House Algorithm, Sklearn and any open-source ML library
- Airflow for job scheduling and monitoring
- Several layers of caching
- Model versioning and inference scheduling
- Containerized which for automated deployment
- User-friendly web interface



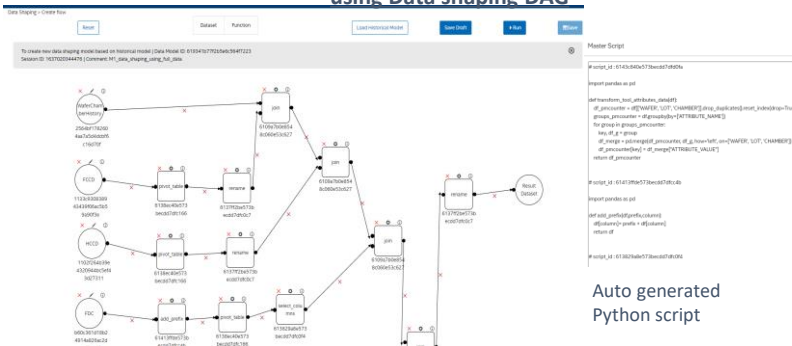
Use Case

Step1: Submit data collection job

Actions	ID	WorkspaceID	DataCategory	Comment	CreatedDate
	894	M1_PS	FDC		2022-01-12 09:39:49 PM PST
	895	M1_PS	SPCContext		2022-01-12 10:28:27 PM PST
	896	M1_PS	WaterChamber-history		2022-01-12 10:28:28 PM PST
	903	M1_PS	SPCContext		2022-01-18 05:42:14 PM PST



Step2: Custom data processing using Data shaping DAG



Auto generated
Python script



Step4: Schedule Inference to automatically receive latest prediction

Model Inference Completed Notification for Inference ID 4f2e2671847d47dcaf049e6a76a58861



Retention Policy Mail Cloud - Inbox (1 year)

Expires 9/7/2023



target_predicted_09072022_000305UTC_09082022_000305UTC.csv
435 KB

Dear IMLA users,

The schedule inference job has been completed for model case 629e31e6f4a26c0001fb459b. Inference Schedule: daily

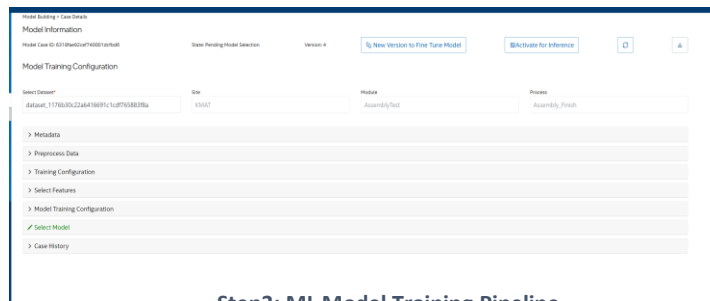
Inference result has been attached.

Log and Warning Message:

Performance monitoring message:

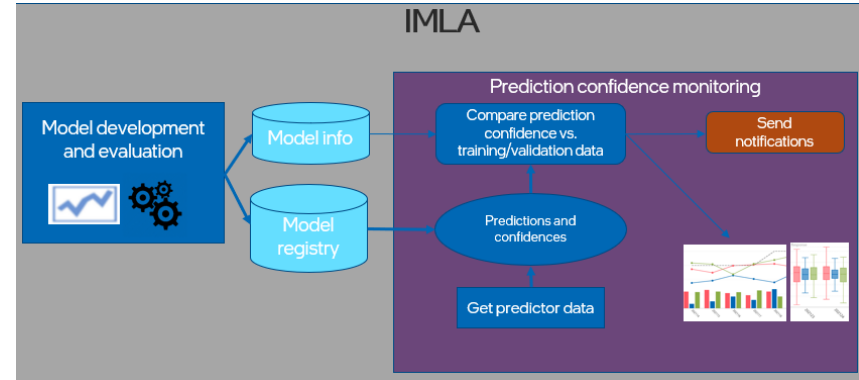
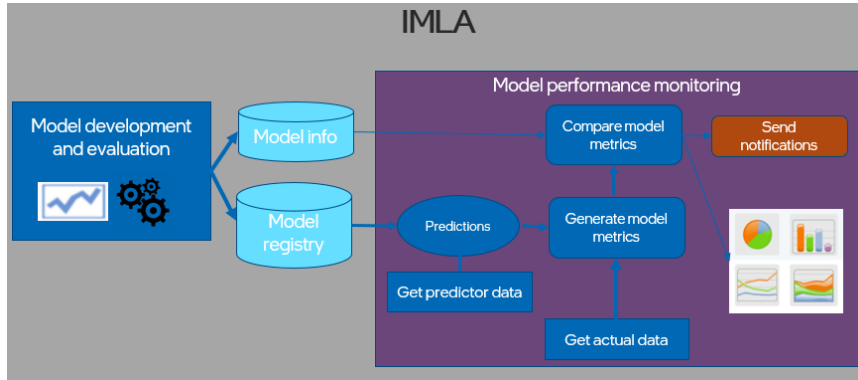


Step5: Integrate models in APC app



Step3: ML Model Training Pipeline

IMLA Performance Monitoring Flow



- IMLA proactively monitors model performance by prediction confidence even if there is a delay in obtaining actual target values
- When actual target is available (metro target e.g.). IMLA will further monitor the model health via computing the performance metric
- Stakeholder will be automatically notified if system has found the model performance deteriorates

Thank you!