

User Guide

Prism - Predictive Intelligence Platform

Version 1.1.0 | May 2026





Table of Contents

Table of Contents	1
1. Introduction.....	2
1.1 What is Prism?	2
1.2 Key Features	2
1.3 Architecture Overview	2
2. Data Requirements by Use Case.....	3
2.1 Churn Prediction Use Case.....	3
Required Columns	3
Data Format Guidelines	4
2.2 Hospital Readmission Risk Use Case.....	4
Data Format Guidelines	6
2.3 Credit Risk Scoring Use Case	6
Data Format Guidelines	7
2.4 Inventory Optimization Use Case.....	7
3. Prerequisites.....	8
3.1 AWS Account Requirements	8
3.2 Required AWS Services.....	8
Bedrock Quota Requirements	9
SageMaker Instance Quota Requirements.....	9
3.3 CDK Bootstrap Guide	10
Option A: AWS CloudShell (Recommended)	11
Option B: From Your Local Machine.....	12
Verification Checklist	14
Troubleshooting.....	14
3.4 Before You Launch Checklist.....	15
4. Setup Guide	15
4.1 AWS Marketplace Deployment.....	15
Step 1: Subscribe to Prism	15
Step 2: Configure the Deployment	15
Step 3: Load CloudFormation Template	16
Step 4: Configure Stack Parameters	16
Step 5: Review and Create.....	17
Step 6: Monitor Deployment.....	17
Step 7: Upload Your Data	18
Step 8: View Results	18
4.2 VPC Deployment Options	18
4.3 Post-Deployment Verification.....	21
5. Data Ingestion	22
5.1 Understanding the Data Ingestion Workflow	22
5.2 Upload Locations by Use Case.....	22



5.3 First Upload: Historical Data	23
5.4 Subsequent Uploads: Inference Only	24
5.5 Triggering the Pipeline Manually	25
5.6 Monitoring Ingestion.....	25
Check Step Functions Status:.....	25
Check Glue Job Status:.....	26
Check SageMaker Pipeline Status:	26
Creating a SageMaker Domain (Optional):.....	26
Check Batch Transform Status:.....	27
5.7 Data Validation	27
5.8 Retraining on New Data	27
6. Workflows & Use Cases.....	28
6.1 Overview	28
6.2 Churn Prediction Use Case.....	28
6.3 Hospital Readmission Risk Use Case.....	29
6.4 Credit Risk Scoring Use Case	30
6.5 Inventory Optimization Use Case	30
6.6 Viewing Predictions.....	31
6.7 Querying Data with Athena	31
7. QuickSight Dashboards	32
7.1 Accessing the Dashboard	32
7.2 Dashboard Components by Use Case	32
7.3 Creating Custom Visuals	34
To create a custom visual:	34
7.4 Sharing Dashboards	35
Share with specific users:	35
Publish to larger audience:	35
7.5 Natural Language Querying with QuickSight Q Topics.....	35
8. Monitoring	37
8.1 CloudWatch Monitoring	37
8.2 Logs	37
Lambda Logs:	37
Glue Job Logs:	37
SageMaker Logs:	38
9. Security Configuration	39
9.1 Encryption	39
9.2 IAM Security	39
9.3 Network Security	39
9.4 Access Control.....	40
10. Cost Management	40



10.1 Cost Components	40
10.2 Cost Monitoring	40
11. Troubleshooting	41
11.1 Common Issues	41
Data Not Processing	41
Glue Job Failures	42
QuickSight Dashboard Empty	42
SageMaker Training Failures	43
11.2 Getting Support.....	43
12. Upgrades & Updates	43
12.1 Update Process	43
12.2 Rollback Procedure	43
13. Backup & Recovery	44
13.1 Data Backup.....	44
S3 Versioning:	44
DynamoDB Point-in-Time Recovery:	44
13.2 Iceberg Time Travel	44
14. Uninstallation	45
14.1 Before Uninstalling.....	45
14.2 Delete CloudFormation Stacks	45
Delete the Installer Stack:.....	45
Delete the PrismPlatform Stack:	45
14.3 Verify Complete Removal	46
15. Support.....	46
15.1 Support Channels	46
15.2 Contacting Support.....	46
10Pearls Support:	46
AWS Marketplace Support:	46
15.3 Information to Include.....	46
15.4 Self-Service Resources	46
Release Notes.....	47
Version 1.0.0 (April 2026).....	47
Initial Release.....	47
Version 1.1.0 (May 2026)	47
New Use Cases.....	47
Infrastructure.....	47
Appendix A: Prism System Configurations	47
A.1 Core Configuration.....	47
A.2 Lambda Configuration.....	48
A.3 Training Dataset Size Tiers	48



A.4 Compute Configuration	48
A.5 Scheduling Configuration	49
A.6 Model Quality Configuration	49
A.7 VPC Configuration.....	49
A.8 S3 Configuration	50
A.9 Glue Configuration.....	50
A.10 SageMaker Configuration.....	51
A.11 QuickSight Configuration.....	51

1. Introduction

1.1 What is Prism?

Prism (Predictive Intelligence Platform) is a customer-deployed, event-driven data Lakehouse and MLOps platform delivered as AWS CDK infrastructure. It provides end-to-end data processing, machine learning model training, and business intelligence capabilities within your own AWS account.

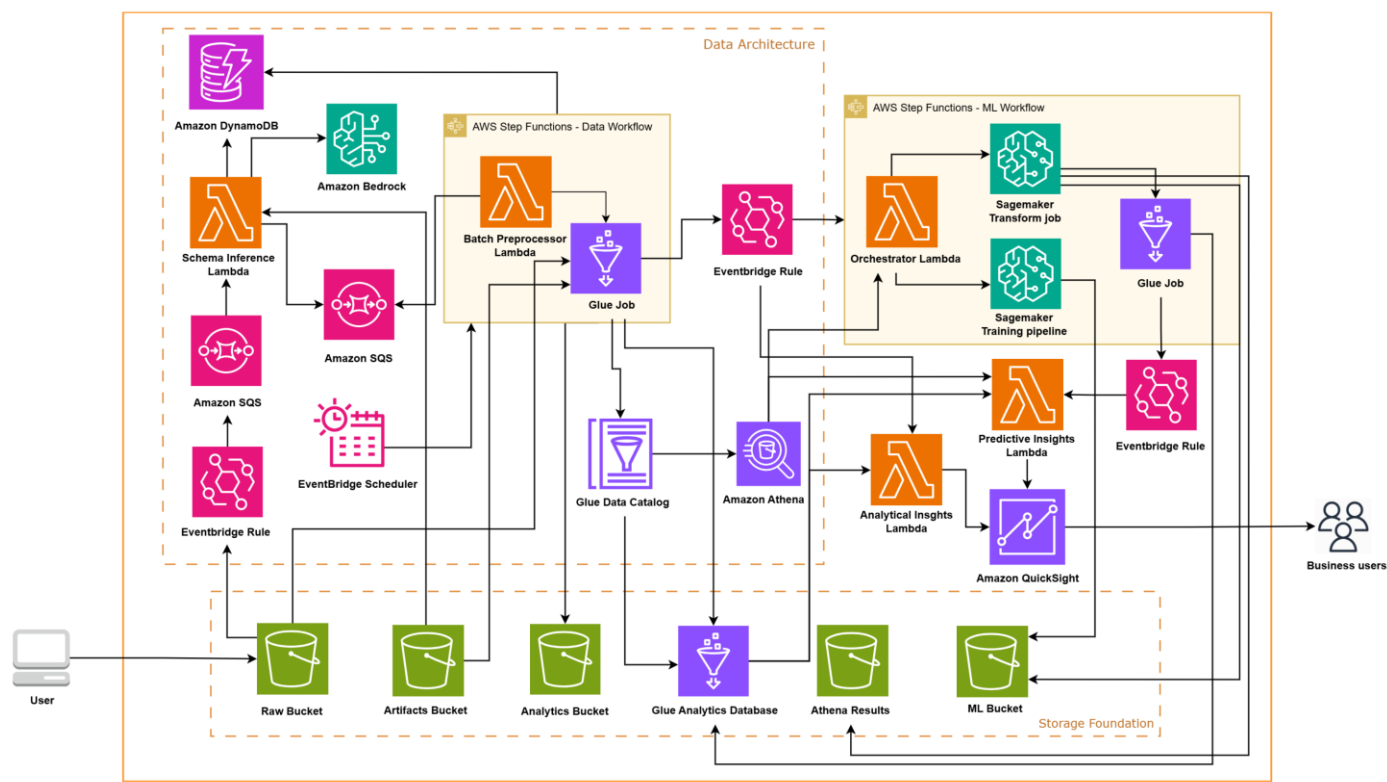
1.2 Key Features

- **Automated Data Processing** – Simply upload your data to S3, and the platform automatically processes it through the entire pipeline without manual intervention.
- **Smart Data Organization** – Your data is stored in a modern data lakehouse format (Apache Iceberg) that supports safe updates, tracks all changes over time, and lets you query historical data whenever needed.
- **Intelligent Schema Mapping** – The platform uses Amazon Bedrock AI to automatically understand your data structure and map it to the required format, saving you time and reducing errors.
- **Reliable Workflow Management** – All data processing runs through AWS Step Functions, which automatically handles retries and ensures your workflows complete successfully even if temporary issues occur.
- **Automated Machine Learning** – Once your data is ready, the platform automatically trains, evaluates, and deploys ML models using Amazon SageMaker—no data science expertise required.
- **Ready-to-Use Dashboards** – Beautiful QuickSight dashboards are automatically created for you, showing predictions, trends, and actionable insights from your data.



1.3 Architecture Overview

The following diagram illustrates the end-to-end architecture of Prism:



Prism: Predictive Intelligence Platform (Default VPC)

2. Data Requirements by Use Case

Prism supports multiple predictive analytics use cases. Each use case has specific data requirements tailored to itself.

2.1 Churn Prediction Use Case

Churn Prediction is specifically designed for telecom industry datasets. To ensure accurate predictions, your data must include specific customer and usage information.

Required Columns

Your telecom churn dataset must contain the following columns:

Column Name	Type	Description
year	Integer	Calendar year for the record (used for time analysis/partitioning).
month	Integer	Calendar month for the record (1–12).
customer_id	String	Unique identifier for each customer.



churn_date	Date	Date when customer churned (blank/null if not churned).
city	String	Customer city (or primary location).
plan_type	String	Customer plan type (e.g., prepaid/postpaid, bundle name).
has_phone_service	Boolean	Whether the customer has phone/voice service.
has_internet_service	Boolean	Whether the customer has internet/data service.
internet_service_type	String	Internet service technology/type (if applicable).
contract_duration_category	String	Contract duration bucket (e.g., month-to-month, 12-month, 24-month).
contract_start_date	Date	Contract start date.
contract_end_date	Date	Contract end date.
customer_join_date	Date	Date the customer joined (used for tenure calculations).
monthly_total_bill_amount	Float	Total billed amount for the month.
total_lifetime_charges	Float	Total charges accumulated over customer lifetime.
payment_method	String	Payment method used by the customer.
monthly_data_usage	Float	Monthly data usage (e.g., GB/MB based on your convention).
monthly_support_calls	Integer	Count of support calls/tickets in the month.
monthly_total_recharge_amount	Float	Total prepaid recharge amount in the month (if applicable).
churn_label	Integer	Target label (1 = churned, 0 = not churned).

Data Format Guidelines

- **File Format:** CSV (comma-separated values).
- **Encoding:** UTF-8.
- **Header Row:** First row must contain column names.
- **Missing Values:** Use empty strings or "NA" for missing values.
- **Decimal Separator:** Use period (.) for decimal numbers.
- **Column Names:** Ensure there are no repeated columns, and the column names are descriptive and semantically similar to the column names listed above.
- **Optional Columns:** Your dataset may include additional columns beyond the required ones. The platform will process them accordingly.

2.2 Hospital Readmission Risk Use Case

Prism analyses hospital admission data, vitals history, and social determinants to flag high-risk patients at discharge. This use case helps healthcare providers identify patients at risk of 30-day hospital readmission.

Required Columns



Your readmission risk dataset must contain the following columns:

Column Name	Type	Description
patient_id	String	Unique identifier for each patient.
encounter_id	String	Unique identifier for each hospital encounter/admission.
gender	String	Patient gender.
age_at_admission	Integer	Patient age at time of admission.
race	String	Patient race.
admission_start	DateTime	Date and time of admission.
admission_stop	DateTime	Date and time of discharge.
admission_type	String	Type of admission (e.g., Emergency, Elective, Urgent).
admission_reason	String	Primary reason for admission.
ed_admission_flag	Integer	Emergency department admission flag (1 = yes, 0 = no).
charlson_score	Integer	Charlson Comorbidity Index score.
heart_failure_flag	Integer	Heart failure diagnosis (1 = yes, 0 = no).
copd_flag	Integer	COPD diagnosis (1 = yes, 0 = no).
diabetes_t2_flag	Integer	Type 2 diabetes diagnosis (1 = yes, 0 = no).
ckd_flag	Integer	Chronic kidney disease diagnosis (1 = yes, 0 = no).
hypertension_flag	Integer	Hypertension diagnosis (1 = yes, 0 = no).
depression_flag	Integer	Depression diagnosis (1 = yes, 0 = no).
sbp_last	Float	Last recorded systolic blood pressure (mmHg).
dbp_last	Float	Last recorded diastolic blood pressure (mmHg).
hr_last	Float	Last recorded heart rate (bpm).
spo2_last	Float	Last recorded oxygen saturation (%).
temp_last	Float	Last recorded temperature.
rr_last	Float	Last recorded respiratory rate.
creatinine_last	Float	Last recorded creatinine level.
hemoglobin_last	Float	Last recorded hemoglobin level.
wbc_last	Float	Last recorded white blood cell count.
platelets_last	Float	Last recorded platelet count.
sodium_last	Float	Last recorded sodium level.
glucose_last	Float	Last recorded glucose level.
albumin_last	Float	Last recorded albumin level.
n_medications	Integer	Number of medications prescribed.
mechanical_ventilation_flag	Integer	Mechanical ventilation used (1 = yes, 0 = no).
hemodialysis_flag	Integer	Hemodialysis performed (1 = yes, 0 = no).
marital_status	String	Patient marital status.
patient_zip	String	Patient ZIP code.
payer_name	String	Insurance payer/type.
n_observations	Integer	Number of clinical observations recorded.



Data Format Guidelines

- File Format: CSV (comma-separated values).
- Encoding: UTF-8.
- Header Row: First row must contain column names.
- Missing Values: Use empty strings or "NA" for missing values.
- Decimal Separator: Use period (.) for decimal numbers.
- Column Names: Ensure there are no repeated columns, and the column names are descriptive and semantically similar to the column names listed above.
- Optional Columns: Your dataset may include additional columns beyond the required ones. The platform will process them accordingly.
- Date/Time Format: Use ISO 8601 format (YYYY-MM-DD HH:MM:SS) for datetime fields.

2.3 Credit Risk Scoring Use Case

Prism enriches bureau data with transactional behaviour and alternative signals for dynamic risk scores. This use case helps financial institutions assess creditworthiness and predict loan default probability.

Required Columns

Your credit risk dataset must contain the following columns:

Column Name	Type	Description
case_id	Integer	Unique identifier for each credit application.
application_date	Date	Date of loan application.
decision_date	Date	Date of credit decision.
date_of_birth	Date	Applicant's date of birth.
gender	String	Applicant's gender.
income_monthly_declared	Float	Declared monthly income.
employment_status	String	Current employment status.
region_code	String	Geographic region code.
applicant_state	String	State of residence.
requested_amount	Float	Requested loan amount.
loan_purpose	String	Purpose of the loan.
total_outstanding_debt	Float	Total existing debt.
max_dpd_12m	Integer	Maximum days past due in last 12 months.
max_dpd_active	Integer	Maximum DPD on active accounts.
max_dpd_closed	Integer	Maximum DPD on closed accounts.
num_hard_enquiries_6m	Integer	Hard credit inquiries in last 6 months.
dpd_hist_range	Float	Historical DPD range.
bureau_score_external	Float	External credit bureau score (optional but recommended).
num_active_credit_lines	Integer	Number of active credit accounts (optional but recommended).
debt_to_income_ratio	Float	Debt-to-income ratio (optional but recommended).



target	Integer	Default label (1 = default, 0 = no default).
--------	---------	--

Data Format Guidelines

- File Format: CSV (comma-separated values).
- Encoding: UTF-8.
- Header Row: First row must contain column names.
- Missing Values: Use empty strings or "NA" for missing values.
- Decimal Separator: Use period (.) for decimal numbers.
- Column Names: Ensure there are no repeated columns, and the column names are descriptive and semantically similar to the column names listed above.
- Optional Columns: Your dataset may include additional columns beyond the required ones. The platform will process them accordingly.
- Date Format: Use YYYY-MM-DD format for date fields.

2.4 Inventory Optimization Use Case

Prism analyses historical sales, stock levels, and supply chain signals to generate demand forecasts and recommend reorder quantities, helping operations teams reduce stockouts and excess inventory.

Column Name	Type	Description
date	Date	Reporting date (daily granularity, YYYY-MM-DD)
Store_id	String	Unique identifier for the store or warehouse location
Product_id	String	Unique SKU or product identifier
category	String	Product category (e.g., Electronics, Apparel)
Inventory_level	Integer	Closing stock units on hand at end of the day
Units_ordered	Integer	Units received from supplier on this date
Units_sold	Integer	Units sold during this period — this is the forecast target
Demand_forecast	Float	Upstream demand forecast (e.g., from a legacy system); used as a feature
price	Float	Selling price per unit
discount	Float	Discount rate applied (0–100)
Holiday_promotion	Integer	Whether a holiday or promotion was active: 0 = No, 1 = Yes
Weather_condition	String	Weather condition label for the period (e.g., Sunny, Rainy)
seasonality	String	Season label — must be one of: Spring, Summer, Autumn, Winter
Competitor_pricing	Float	Competitor's price for the equivalent product
region	String	Geographic region identifier (e.g., North, South-East)



Data Format Guidelines

- File Format: CSV (comma-separated values).
- Encoding: UTF-8.
- Header Row: First row must contain column names.
- Missing Values: Use empty strings or "NA" for missing values.
- Decimal Separator: Use period (.) for decimal numbers.
- Column Names: Ensure there are no repeated columns, and the column names are descriptive and semantically similar to the column names listed above.
- Optional Columns: Your dataset may include additional columns beyond the required ones. The platform will process them accordingly.
- Date Format: Use YYYY-MM-DD format for date fields.

3. Prerequisites

3.1 AWS Account Requirements

Requirement	Details
AWS Account	Active AWS account with billing enabled.
IAM Permissions	Administrator access or equivalent permissions.
AWS Region	Region with all required services.

3.2 Required AWS Services

The following AWS services must be available and enabled in your account:

Service	Purpose	Cost Type
Amazon S3	Data storage	Pay-per-use
AWS Glue	ETL processing & data catalog	Pay-per-use
Amazon Athena	SQL query engine	Pay-per-query
Amazon SageMaker	ML model training & inference	Pay-per-use
Amazon QuickSight	Business intelligence dashboards	Subscription
Amazon Bedrock	AI schema inference	Pay-per-use
AWS Lambda	Serverless compute	Pay-per-use
Amazon EventBridge	Event routing	Pay-per-use
Amazon SQS	Message queuing	Pay-per-use
Amazon DynamoDB	Schema registry	Pay-per-use



AWS Step Functions	Workflow orchestration	Pay-per-use
--------------------	------------------------	-------------

Bedrock Quota Requirements

Before deploying, ensure Amazon Bedrock Nova Pro and Meta Llama 4 Maverick V1 have sufficient token quota in your account:

1. Open the AWS Console and navigate to **Amazon Bedrock → Service Quotas**.
2. Search for "**Nova Pro**" and look for **"Cross-region model inference tokens per minute for Amazon Nova Pro"** in the list. Additionally, search for **"Meta Llama 4 Maverick V1"** and look for **"Cross-region model inference tokens per minute for Meta Llama 4 Maverick V1"**.
3. Check the **TPM (Tokens Per Minute)** value. It must be **2,000,000 (2M) or higher**.
4. If it shows 0 or less than 2M, click **Request quota increase** and request a minimum of **2,000,000 TPM**.
5. Quota increase requests are typically approved within a few business days.

Warning: Deploying before this quota is approved will cause the platform's AI schema analysis feature to fail at runtime.

SageMaker Instance Quota Requirements

Before deploying, ensure your account has sufficient SageMaker instance quotas for processing, training, and transform jobs:

Instance Type	Quota Name to Search	Minimum Required (Each use case)
m1.m5.large	"m1.m5.large for processing job usage"	1
m1.m5.xlarge	"m1.m5.xlarge for processing job usage"	1
m1.m5.2xlarge	"m1.m5.2xlarge for processing job usage"	1
m1.m5.4xlarge	"m1.m5.4xlarge for processing job usage"	1
m1.m5.large	"m1.m5.large for transform job usage"	1
m1.m5.xlarge	"m1.m5.xlarge for transform job usage"	1
m1.m5.2xlarge	"m1.m5.2xlarge for transform job usage"	1



m1.m5.4xlarge	"m1.m5.4xlarge for transform job usage"	1
m1.m5.xlarge	"m1.m5.xlarge for training job usage"	3
m1.m5.2xlarge	"m1.m5.2xlarge for training job usage"	3
m1.m5.4xlarge	"m1.m5.4xlarge for training job usage"	3

The Default Quotas are as follows:

Service	Quota	Default	Action
Lambda	Concurrent executions	1,000	Request increase for high volume
SQS	Messages in queue	Unlimited	-
S3	Buckets per account	100	Request increase if needed
Glue	Max DPU's per account	100	Request increase for parallel jobs
Bedrock	Nova Pro and Meta Llama 4 Maverick ("Cross-region model inference tokens per minute for Amazon Nova Pro" and "Llama 4 Maverick V1") TPM (Tokens Per Minute)	Varies	Must be $\geq 2,000,000$ TPM
QuickSight	SPICE capacity	10GB	Purchase additional capacity
Step Functions	State transitions/second	2,000	Usually sufficient

If you want to increase the quota usage:

1. Go to **Service Quotas Console**.
2. Select the AWS service.
3. Find the quota you need to increase.
4. Click **Request quota increase**.
5. Enter desired value and submit.

3.3 CDK Bootstrap Guide

CDK bootstrap creates a CloudFormation stack called **CDKToolkit** in your AWS account that provisions:



Resource	Purpose
S3 Bucket	Stores deployment assets (Lambda code, Glue scripts, etc.)
ECR Repository	Stores container images (if needed)
IAM Roles (5)	Allow CDK to deploy resources on your behalf
SSM Parameter	Tracks the bootstrap version

These resources are used automatically by the Prism installer during deployment.

Prerequisites

- An active AWS account with billing enabled
- A user or role with **AdministratorAccess** (or equivalent permissions for CloudFormation, S3, ECR, IAM, and SSM)
- Access to a terminal (Command Prompt, PowerShell, Terminal, or CloudShell)

Option A: AWS CloudShell (Recommended)

AWS CloudShell is a browser-based terminal available directly in the AWS Management Console. The AWS CLI, Node.js, and the CDK CLI are pre-installed there is nothing to set up on your local machine.

You can run the bootstrap command from a CloudShell session in any region. The target region is specified in the command itself.

Step	Action
1	Sign in to the AWS Management Console with an IAM user or role that has the required permissions.
2	Click the CloudShell icon (>_) in the top navigation bar, or search for "CloudShell" in the services search bar. Wait a few seconds for the session to initialize.
3	Copy and paste the following commands into the CloudShell terminal: <pre>ACCOUNT_ID=\$(aws sts get-caller-identity --query "Account" --output text) cdk bootstrap aws://\$ACCOUNT_ID/us-east-1</pre> The first command retrieves your 12-digit AWS account ID. The second command performs the bootstrap. This typically takes one to two minutes.



4	Run the following command to confirm the CDKToolkit stack was created successfully: <pre>aws cloudformation describe-stacks \ --stack-name CDKToolkit \ --region us-east-1 \ --query "Stacks[0].StackStatus" \ --output text</pre> A result of CREATE_COMPLETE or UPDATE_COMPLETE confirms the bootstrap was successful. You may now proceed with the Prism deployment.
---	---

Option B: From Your Local Machine

Use this option if you prefer to work from your own computer. This requires installing the AWS CLI, Node.js, and the CDK CLI.

Step	Action	Commands
1	Install the AWS CLI	macOS: <pre>curl "https://awscli.amazonaws.com/AWSCLIV2.pkg" -o "AWSCLIV2.pkg" sudo installer -pkg AWSCLIV2.pkg -target /</pre> Windows: Download and run the installer from https://awscli.amazonaws.com/AWSCLIV2.msi Linux: <pre>curl "https://awscli.amazonaws.com/awscli-exe-linux-x86_64.zip" -o "awscliv2.zip" unzip awscliv2.zip sudo ./aws/install</pre>
2	Configure Your AWS Credentials	Run: <pre>aws configure</pre> When prompted, enter the following: <pre>AWS Access Key ID [None]: <your-access-key> AWS Secret Access Key [None]: <your-secret-key> Default region name [None]: us-east-1 Default output format [None]: json</pre>



		To find your access keys, go to IAM Console > Users > select your user > Security credentials tab > Access keys. Click Create access key if you do not have one. Verify your credentials: <pre>aws sts get-caller-identity</pre> You should see a response containing your account ID, user ARN, and user ID.
3	Install Node.js	macOS: <pre>brew install node@20</pre> Windows: Download the LTS installer from https://nodejs.org/en/download Linux (Ubuntu/Debian): <pre>curl -fsSL https://deb.nodesource.com/setup_20.x sudo -E bash - sudo apt-get install -y nodejs</pre>
4	Install the CDK CLI	<pre>npm install -g aws-cdk</pre> Verify the installation: <pre>cdk --version</pre>
5	Run CDK Bootstrap	<pre>ACCOUNT_ID=\$(aws sts get-caller-identity --query "Account" --output text) cdk bootstrap aws://\$ACCOUNT_ID/us-east-1</pre> This typically takes one to two minutes.
6	Verify the Bootstrap	<pre>aws cloudformation describe-stacks \ --stack-name CDKToolkit \ --region us-east-1 \ --query "Stacks[0].StackStatus" \ --output text</pre> A result of CREATE_COMPLETE or UPDATE_COMPLETE confirms the bootstrap was successful.



Verification Checklist

After bootstrapping, you can confirm all resources are in place by running the following commands:

What to Check	Command	Expected Result
CDKToolkit stack	<code>aws cloudformation describe-stacks --stack-name CDKToolkit --region us-east-1 --query "Stacks[0].StackStatus" --output text</code>	CREATE_COMPLETE or UPDATE_COMPLETE
Assets bucket	<code>aws s3 ls --region us-east-1 \ grep cdk-hnb659fds-assets</code>	A bucket name containing your account ID
Bootstrap version	<code>aws ssm get-parameter --name "/cdk-bootstrap/hnb659fds/version" --region us-east-1 --query "Parameter.Value" --output text</code>	A version number (e.g., 22)

If all checks return the expected results, your environment is ready for Prism deployment.

Troubleshooting

Issue	Resolution
"Unable to resolve AWS account"	Your AWS credentials are not configured or have expired. Run <code>aws configure</code> and verify with <code>aws sts get-caller-identity</code> .
"Access Denied" during bootstrap	Your IAM identity does not have sufficient permissions. Use AdministratorAccess or ensure the permissions listed in the Required Permissions section are in place.
"CDKToolkit already exists"	No action is needed. Your account is already bootstrapped. The command will update the stack automatically if a newer version is available.
Bootstrap succeeded but deployment fails	Confirm you bootstrapped in us-east-1 . CDK bootstrap is region-specific, and Prism requires the bootstrap resources in us-east-1.
"npm: command not found"	Node.js is not installed or not in your system PATH. Follow Step 3 in Option B to install it.

3.4 Before You Launch Checklist

Complete these steps before launching the deployment:

Step	Action	How to Verify
------	--------	---------------



0	CDK Bootstrap setup	Detailed guide in section 3.3.
1	Activate Amazon QuickSight	QuickSight Console shows your account is active.
2	Verify Bedrock Nova Pro and Meta Llama 4 Maverick V1 quota	Service Quotas → Bedrock → Nova Pro TPM: “Cross-region model inference tokens per minute for Amazon Nova Pro” ≥ 2,000,000 and Meta Llama 4 Maverick V1 TPM: “Cross-region model inference tokens per minute for Meta Llama 4 Maverick V1” ≥ 600,000 (see Section 3.2).
3	Verify SageMaker instance quotas	Service Quotas → SageMaker → Check processing, training, and transform job quotas (see Section 3.2).
4	Request quota increases if needed	Service Quotas Console → Submit requests for Bedrock and SageMaker.
5	(Optional) Create SageMaker Domain	Required only if you want to view ML pipelines in SageMaker AI Studio (see Section 5.6)

4. Setup Guide

4.1 AWS Marketplace Deployment

Follow these steps to deploy Prism from AWS Marketplace:

Step 1: Subscribe to Prism

1. Navigate to AWS Marketplace.
2. Search for "Prism" or "Predictive Intelligence Platform".
3. Click **View purchase options**.
4. Select your preferred **contract renewal option**.
5. Choose the **license** that matches your intended usage.
6. Review the **EULA and Offer** documents before proceeding.
7. Click **Subscribe** to complete your subscription.

Step 2: Configure the Deployment

1. Choose the latest software version.
2. Review and copy the usage instructions - you may need these when filling in the deployment parameters.
3. Click **Launch with CloudFormation**.
4. You will be redirected to the CloudFormation console.



Step 3: Load CloudFormation Template

1. Keep the default selections. "**Choose an existing template**" and "**S3 URL**" should already be filled in. Do not change these. Click **Next**.

Step 4: Configure Stack Parameters

Now you should see the **Specify stack details** page. Configure the stack parameters below. Use the usage instructions copied earlier as a reference or refer to the following parameter descriptions for guidance on each field.

Provide a Stack Name

- Recommended stack name is prism.

Compute Configuration

- **Training Dataset Size:** Select your dataset size tier first before choosing any instance types. The stack will fail if instance types are too small for your data. Options: SM for datasets under 500MB, MD for 500MB–2GB, LG for 2GB–5GB.

Tier	Data Size	Processing	Training	Inference
SM	Under 500 MB	m1.m5.xlarge	m1.m5.xlarge	m1.m5.xlarge
MD	500 MB - 2 GB	m1.m5.xlarge	m1.m5.2xlarge	m1.m5.2xlarge
LG	2 GB - 5 GB	m1.m5.2xlarge	m1.m5.4xlarge	m1.m5.4xlarge

Warning: The stack will fail if instances are too small for the selected tier. Choose appropriately based on your dataset size.

- **Max Glue ETL Workers:** Number of parallel workers for the ETL job. Default is 5. Increase up to 15 if you have a larger dataset (2-5 GBs) and want faster processing.

Schedule & Model Quality

- **Pipeline Schedule:** How frequently the pipeline runs automatically to ingest new data. Default rate is 1 hour. Adjust based on how often your data updates.
- **Minimum AUC Threshold:** The minimum model quality score (AUC) required before a newly trained model is deployed. Default is 0.65. A higher value enforces stricter model quality but may prevent deployment if data quality is low.

Environment Configuration

- **QuickSight Service Role Name:** The IAM role QuickSight uses to access your data. The default value `aws-quicksight-service-role-v0` is created automatically when you first activate QuickSight. Only change this if your account uses a custom QuickSight role.
- **QuickSight Username:** Your QuickSight user identifier. To find it, open QuickSight, click the profile icon in the top right, and select **Manage QuickSight**. Can be either your username or the email address associated with your QuickSight account.



- **Admin Email for SNS Notifications:** Email address that will receive operational alerts and pipeline notifications. Replace the placeholder with a valid email address. You will receive a confirmation email from AWS SNS after deployment. You must confirm the subscription to receive alerts.

Installer EC2 Configuration

- **AMI ID:** Pre-configured by 10Pearls. Do not change.
- **Installer EC2 Instance Type:** Pre-configured by 10Pearls. Do not change.
- **Installer VPC CIDR Block:** CIDR block for the temporary installer **VPC**. Default is 10.0.0.0/16. No changes are required.

Assets Configuration

- **Seller Assets S3 Bucket Name:** Pre-configured by 10Pearls. Do not change.
- **Customer Assets S3 Bucket Name:** Pre-configured by 10Pearls. Do not change.

Step 5: Review and Create

1. Review all parameters.
2. Acknowledge IAM resource creation by checking the required boxes.
3. Click **Create stack**.

Step 6: Monitor Deployment

1. The stack will show **CREATE_IN_PROGRESS** until the installer signals completion (20-40 minutes).
2. To monitor progress in real-time:
 - Go to **CloudFormation Console** → **Outputs** tab.
 - Find the **InstallerInstanceId**.
 - Open **CloudShell** and connect to the instance via SSM:

```
aws ssm start-session --target <instance-id> --region us-east-1
```

- Stream logs with:

```
sudo journalctl -u prism-deployer.service -f
```

3. The installer instance **self-terminates on success**.
4. If the instance remains running after 60 minutes, there may be an error - check the deployment logs.
5. Once complete, the stack status will show **CREATE_COMPLETE**.
6. Check the **Outputs** tab for important resource information (bucket names, dashboard URLs, etc.).

Step 7: Upload Your Data

1. Navigate to S3 in the AWS Console and find the raw data bucket: prism-raw-{account}-{region}.
2. Navigate to the folder for your use case and upload your CSV file (refer to **Section 5.2** for folder paths and file requirements).



3. The pipeline triggers automatically based on the schedule selected (default 1 hour) during deployment. To trigger it immediately, refer to **Section 5.5**.

Warning: Do not start multiple executions in parallel. Wait for the first execution to complete before executing the state machine again.

Step 8: View Results

1. Once processing completes, open **QuickSight**.
2. Navigate to **Dashboards** to view your results. Refer to Section 5.3 to understand what to expect after your first upload and how long each stage takes.

4.2 VPC Deployment Options

Prism supports three VPC deployment modes:

USE_DEFAULT (Simplest Option)

Uses the **default VPC** in your AWS account.

Behaviour:

- Deploys resources into the AWS default VPC.
- Uses public subnets.
- No VPC endpoints are created.

When to use:

- Quickest setup with minimal configuration.
- Lowest complexity.
- Recommended if you are unfamiliar with VPC networking.

To deploy with default VPC:

1. During CloudFormation stack creation, set **VpcDeploymentMode** to `USE_DEFAULT`.
2. No additional VPC configuration is required.

Pros: Simplest setup, lowest cost, no VPC expertise required.

Cons: Less network isolation uses public subnets.

CREATE_NEW (Recommended for Production)

Creates a new VPC with private subnets and VPC endpoints for AWS services. This provides better isolation and security.

Required parameters:

- **VpcCidr:** CIDR block for the new VPC (example: 10.0.0.0/16).
- **MaxAzs:** Number of Availability Zones (keep in us-east-1).
- **NatGateways:** Number of NAT Gateways (0 or 1 allowed).

Behavior:



- A new VPC with public and private subnets is created.
- If `NatGateways = 1`, NAT Gateway is created.
- If `NatGateways = 0`, required **interface endpoints** are created for the following AWS services:
 - Amazon SQS.
 - Amazon Bedrock Runtime.
 - Amazon SageMaker Runtime.
 - Amazon SageMaker API.
 - AWS Step Functions.
 - AWS Glue.
 - Amazon CloudWatch Logs.
 - Amazon EventBridge.
 - Amazon Athena.

To deploy with a new VPC:

1. Set **VpcDeploymentMode** to `CREATE_NEW`.
2. Configure the following parameters:
 - **VpcCidr**: CIDR block for the new VPC (default: 10.0.0.0/16).
 - **NatGateways**: Set to 0 to use VPC Endpoints (recommended for cost savings), or 1 for NAT Gateway.
 - **MaxAzs**: Maximum availability zones (default: 1).

Recommendation: If you choose **CREATE_NEW**, keep **NatGateways = 0** unless you specifically require NAT Gateways. This avoids unnecessary networking costs.

Pros: Better security, private subnets, VPC endpoints.

Cons: Slightly higher cost if using NAT Gateway.

USE_EXISTING

Uses an existing VPC in your account. Ideal when you have specific networking requirements or compliance needs.

Prerequisites for USE_EXISTING:

Before selecting this mode, ensure your VPC meets the following requirements:

1. **DNS Settings must be enabled:** These are required for internal service communication.
 - Enable **DNS resolution**.
 - Enable **DNS hostnames**.
2. **Private Subnets:** You must have private subnets available in the VPC.
3. **Route Tables:** Route tables must be properly configured for your subnets.
4. **Internet Access:** If you want to use an existing NAT Gateway, set **NatGateways** to 1 and ensure it is configured for your subnets.

To deploy with an existing VPC:

1. Set **VpcDeploymentMode** to `USE_EXISTING`.
2. Provide the following required parameters:



- **VpcId**: Your existing VPC ID (e.g., vpc-0a1b2c3d).
- **AvailabilityZones**: Comma-separated list of availability zones (e.g., us-east-1a,us-east-1b).
- **PrivateSubnetIds**: Comma-separated list of private subnet IDs (e.g., subnet-0a1b2c3d,subnet-0e5f6a7b).
- **RouteTableIds**: Comma-separated list of route table IDs.

3. Set **NatGateways** to 1 if you want to use your own NAT Gateway instead of VPC endpoints.

Required Security Group Configuration:

If you provide existing security groups for **SageMaker**, **Glue**, or **Lambda**, they must allow **self-referencing traffic** so jobs can communicate within the service.

Example configuration:

```
Allow inbound traffic:
Source: same security group
Protocol: All traffic
```

Equivalent behavior to:

```
SecurityGroupIngress:
Source: self
Port: all
Protocol: all
```

This allows distributed processing jobs (such as SageMaker training and Glue processing) to communicate internally.

Using Existing Security Groups:

If you enable existing security groups for any of the following services:

Service	Parameter
Glue	use_existing_security_groups_glue = true
Lambda	use_existing_security_groups_lambda = true
SageMaker	use_existing_security_groups_sagemaker = true

You must also provide the corresponding security group IDs (e.g., sg-0a1b2c3d).

If this option is **false**, the platform will automatically create and manage security groups for you.

NAT Gateway Behavior:

- **NatGateways = 0**
 - The platform creates required **VPC endpoints** for AWS services.
 - This is the **lowest cost configuration**.
- **NatGateways = 1**
 - The deployment assumes a **NAT Gateway already exists**.
 - The NAT Gateway must be correctly attached to the provided VPC and route tables.



- The platform will route outbound traffic through this NAT Gateway instead of creating endpoints.

Pros: Use your existing network infrastructure, comply with organization policies.

Cons: Requires more configuration and VPC expertise.

VPC Recommendation

If you are **not familiar with VPC networking**, the simplest options are:

- **USE_DEFAULT** – No configuration needed.
- **CREATE_NEW** – Platform handles VPC creation.

In both cases, set:

```
nat_gateways = 0
```

This enables **VPC endpoints instead of NAT Gateways**, which significantly reduces networking costs.

This configuration works out of the box and prevents unnecessary networking costs.

4.3 Post-Deployment Verification

After deployment completes, verify the following:

Check S3 Buckets:

1. Go to **S3 Console**
2. Verify these buckets exist:
 - prism-raw-{account}-{region} – for raw data uploads. See Section 5.2 for the correct folder path for your use case(s).
 - prism-analytics-{account}-{region} – for processed data.
 - prism-ml-{account}-{region} – for ML artifacts.

Check Glue Resources:

1. Go to **AWS Glue Console**.
2. Verify the database prism_db exists.
3. Check that Glue jobs are created.

Check Step Functions:

1. Go to **Step Functions Console**.
2. Verify the state machine is created.
3. Check recent executions (if any).

Check QuickSight:

1. Open **QuickSight Console**.
2. Navigate to **Datasets** and verify datasets are created.



3. Navigate to **Dashboards** and verify dashboard is available.

5. Data Ingestion

5.1 Understanding the Data Ingestion Workflow

All Prism use cases share the same automated ingestion pipeline — once you upload a CSV to the designated S3 folder, the platform handles everything from schema inference through model training to dashboard population without manual intervention.

All use cases support the same two upload modes; choose based on how your data is structured and how frequently you expect to run inference.

Mode 1 - Combined Training and Inference: Upload a single file containing both historical labeled records and recent unlabeled records. The platform trains and scores in one pipeline run.

Mode 2 - Separate Training and Inference: Upload your historical labeled dataset first to train the model, then upload new unlabeled records separately for scoring. You can repeat the inference upload as frequently as needed without retraining.

Regardless of use case or mode, every upload automatically triggers:

- Schema detection via Amazon Bedrock AI
- ETL transformation into Apache Iceberg format via AWS Glue
- SageMaker pipeline execution (training and/or batch inference depending on upload type)
- QuickSight dashboard refresh

Note: The pipeline runs on a scheduled basis (default: every 1 hour). See Section 5.5 to trigger it immediately after uploading.

5.2 Upload Locations by Use Case

Upload your CSV files to the folder corresponding to your use case in the raw S3 bucket created during deployment:

Churn Prediction	<code>s3://prism-raw-{account}-{region}/churn-data/</code>
Hospital Readmission Risk	<code>s3://prism-raw-{account}-{region}/readmission_risk-data/</code>



Credit Risk Scoring	<code>s3://prism-raw-{account}-{region}/credit_risk-data/</code>
Inventory Optimization	<code>s3://prism-raw-{account}-{region}/inventory_optimization-data/</code>

Uploading to the wrong folder will cause the pipeline to either skip your file or process it against the wrong schema. Always verify the path before uploading.

5.3 First Upload: Historical Data

Upload your full historical labeled dataset. The platform trains the model on this data. As part of the same pipeline run, it also automatically generates an initial set of predictions to populate the dashboard — no second upload is needed to see your first results. The slice used for that initial inference is determined by the platform based on the use case:

Use Case	Initial Inference Slice
Churn Prediction	Most recent month of records in the training file
Hospital Readmission Risk	Encounters discharged within the last 30 days of the training file
Credit Risk Scoring	Most recent 2% of records in the training file
Inventory Optimization	Most recent 30 days of records in the training file

For credit risk, you can alternatively include recent unlabelled applications directly in this file — the platform will train on records with known target values and score the rest in the same run. See Section 6.4 for details.

The target column must be present and populated for all historical records **except for Readmission Risk use case** where it is derived by the platform.

Use Case	Duration	Description
Churn Prediction	churn_label	1 = churned, 0 = not churned
Hospital Readmission Risk (derived by Prism)	readmitted_30d	1 = readmitted within 30 days, 0 = not readmitted
Credit Risk Scoring	target	1 = default, 0 = no default
Inventory Optimization	Units_sold	Numerical value > 0



Processing Timeline:

Stage	Duration	Description
Schema Inference	~1 minute	Amazon Bedrock analyses your data structure
Glue ETL	~2-5 minutes	Transforms data to Iceberg format (varies by size: 1MB-2GB typical)
SageMaker Training Pipeline	~20-50 minutes	Model training, evaluation, and deployment

What to Expect:

1. Once Gluejob completes, your **historical analytics QuickSight dashboard** will be available within 1-5 minutes depending on the data size.
2. Once training completes, the platform automatically scores the initial inference slice and your **prediction insights dashboard** is populated.
3. Monitor training progress in **Sagemaker AI Console** → **SageMaker Studio** → **Pipelines** (a domain is required to access SageMaker Studio, refer to section 5.6).
4. Wait for the SageMaker training pipeline to complete successfully before uploading inference data.

5.4 Subsequent Uploads: Inference Only

Once the model is registered from the first upload, all subsequent uploads go directly to batch transform — the platform does not retrain automatically. The registered model is used for all future scoring, so it is important to verify the model's AUC meets your quality threshold before relying on it for production inference. Monitor this in SageMaker Model Registry before uploading inference files.

You can repeat this upload as frequently as needed against the same registered model.

Processing Timeline:

Stage	Duration	Description
Schema Inference	~1 minute	Validates data structure
Glue ETL	~2-5 minutes	Processes inference data
SageMaker Batch Transform	Varies	Generates predictions (varies by dataset size)
Dashboard Update	~1-5 minutes	QuickSight dashboard updated to include predictions sheet

What to Expect:

1. Once Glue completes, it automatically triggers a SageMaker Batch Transform job.
2. Monitor progress in **SageMaker AI Console** → **Batch Transform jobs**.



3. When inference finishes, a second Glue job loads the predictions into the Iceberg table.

4. Your **QuickSight prediction insights dashboard** is automatically updated.

5.5 Triggering the Pipeline Manually

If you prefer not to wait for the next scheduled run, you can trigger the pipeline immediately by starting the workflow directly. This method is more straightforward and avoids the need to modify complex scheduling rules:

Step	Action
1	Open the AWS Step Functions Console.
2	Navigate to "State machines" in the left-hand menu.
3	Search for the pipeline (typically named with a prism-processing- prefix).
4	Select the state machine and click the Start execution button.
5	Confirm the execution: You can optionally provide a unique name for the run, or simply click Start execution again to use the default settings.
6	Monitor Progress: The console will display a visual graph of the workflow. The run is complete when the status reaches Succeeded .

5.6 Monitoring Ingestion

Check Step Functions Status:

The platform uses two Step Functions state machines:

Step	Action
1	Go to Step Functions Console .
2	Find and monitor the following state machines: • prism-processing-workflow – Main orchestration workflow for data processing. • prism-ml-workflow – ML training and inference workflow.
3	View the latest execution.
4	Check execution status and details.

Check Glue Job Status:

The platform uses multiple Glue jobs:

Step	Action
------	--------



1	Go to AWS Glue Console .
2	Navigate to ETL Jobs .
3	Monitor the following jobs: ● prism-etl-{usecase}-job – Main ETL job for processing churn data. ● prism-{usecase}-ingest-predictions-job – Job for loading inference results.
4	View job run history and status.

Check SageMaker Pipeline Status:

Step	Action
1	Go to SageMaker AI Console .
2	Open SageMaker Studio (a domain is required to access SageMaker Studio, refer to the section below)
3	Navigate to Pipelines .
4	Select your pipeline (it will have a “prism” prefix).
5	View execution status and logs.

Creating a SageMaker Domain (Optional):

To view ML pipelines with a visual interface in SageMaker AI Studio:

Step	Action
1	Go to SageMaker AI Console .
2	Click Domains in the left navigation.
3	Click Create domain .
4	Choose Quick setup for the fastest configuration.
6	Click Submit .
7	Wait for domain creation to complete (5-10 minutes).
8	Once created, click Open Studio to access the visual pipeline interface.

Note: Creating a SageMaker Domain is optional and not required for the platform to function. It only provides a visual interface for monitoring ML pipelines.



Check Batch Transform Status:

Step	Action
1	Go to SageMaker AI Console .
2	Navigate to Batch Transform jobs .
3	View job status and completion time.

Check SageMaker Training and Processing Jobs:

Step	Action
1	Go to SageMaker AI Console .
2	Navigate to Training → Training jobs to view model training status.
3	Navigate to Processing → Processing jobs to view data processing status.

5.7 Data Validation

The platform performs automatic validation including:

- Column presence verification.
- Data type checking.
- Value range validation.
- Duplicate detection.

If validation fails, check CloudWatch logs for detailed error messages.

Tip: Review Section 2 (Data Requirements) before uploading to ensure your data has all required columns in the correct format.

5.8 Retraining on New Data

There is no automatic retraining. If you want to retrain the model, for example, after accumulating significant new labeled data, delete the existing model from SageMaker and re-upload your dataset for the specific use case:

Step	Action
1	Go to SageMaker AI Console → SageMaker Studio → Models → My Models → prism- <use-case>-model-group, for example: <i>prism-churn-model-group</i> . Note: A domain is required to access SageMaker Studio, refer to section 5.6.
2	Go to SageMaker AI Console → Deployments & inference → Deployable models and delete the deployable model with the name prism- <use-case>-model-v1, for example: <i>prism-churn-model-v1</i> .
3	Upload your new labeled dataset and trigger the pipeline manually via Step Functions (see Section 5.5).



6. Workflows & Use Cases

6.1 Overview

All Prism use cases follow the same automated pipeline architecture. Once you upload data to S3, the platform handles the entire process from schema inference through model training to dashboard updates. The key differences between the use cases lie in data structure requirements, target definitions, and how predictions should be interpreted.

Common Pipeline Steps:

Step	Action
1	Data Upload – You upload data to the S3 raw bucket.
2	Event Trigger – S3 upload triggers EventBridge event.
3	Schema Inference – Amazon Bedrock analyses and maps your data schema.
4	Data Processing – AWS Glue transforms data into Iceberg format.
5	Feature Engineering – Relevant features are extracted for ML.
6	Model Training – SageMaker trains the prediction model.
7	Batch Inference – Model generates predictions for all customers upon data upload.
8	Dashboard Update – QuickSight displays predictions and insights.

6.2 Churn Prediction Use Case

Data Granularity and Structure:

Your churn data should be structured with monthly aggregation at the customer level. Each row represents a single customer's activity and status for a specific month, containing all relevant behavioral signals and service usage information for that period.

Understanding the Churn Target:

The churn target indicates whether a customer was considered "churned" during a specific month. The platform does not prescribe how you define churn—this is determined by your business logic. For example,



you might define churn as "one month of no payment and inactivity" or use a different criterion based on your operational needs.

What matters is consistency: the target definition you use during training directly influences how predictions should be interpreted. If your churn target equals 1 when a customer shows "one month of no payment and inactivity," then a prediction of 1 for a customer in a given month means they are highly likely to become inactive and miss payments in the following month.

How the Platform Handles Training:

To predict future churn, the platform shifts the target labels forward by one month during training. This ensures the model learns to predict next-month churn based on current-month features. Additionally, customers who have already churned are excluded from subsequent months to prevent data leakage and maintain prediction accuracy.

When analyzing predictions and dashboard insights, always interpret results in the context of your specific churn definition to ensure alignment with your business objectives.

6.3 Hospital Readmission Risk Use Case

Data Granularity and Structure:

Your readmission data should be structured at the patient-encounter level, with one row per hospital admission for a patient. Each row should include patient demographics, admission details, clinical vitals, lab results, comorbidities, and social determinants captured during the encounter.

Understanding the Readmission Target:

The readmission target indicates whether a patient will be readmitted to the hospital within 30 days of discharge. The platform derives this variable from historical data by identifying whether each patient had a subsequent admission (encounter) within 30 days following their discharge date.

How the Platform Handles Training:

The platform trains on historical encounter data where the readmission outcome is already known. Features are extracted from the admission period (demographics, vitals, labs, procedures, medications) to predict the post-discharge readmission risk. The model learns patterns associated with readmission events based on this 30-day readmission definition.

When reviewing predictions and clinical insights in the dashboard, interpret risk scores in the context of this readmission definition to support appropriate care coordination and intervention strategies.

6.4 Credit Risk Scoring Use Case



Data Granularity and Structure:

Your credit risk data should be structured at the loan application level, with one row per credit application (identified by `case_id`). Each row can contain multiple months of historical signals in a flat file format, combining credit bureau information, transactional behavior, and alternative data sources.

Understanding the Default Target:

The default target indicates whether a borrower failed to repay their loan, typically defined as non-payment after a specific number of days past due (DPD). While the standard threshold is 90 days past due, your organization may use a different definition (e.g., 60 DPD, 120 DPD).

The platform does not enforce a specific default definition—what matters is that you apply your definition consistently. When the model predicts a high default probability for an application, it means that applicant is likely to meet your specific default criteria based on the patterns learned during training.

How the Platform Handles Training and Inference:

The platform supports two operational modes:

Mode 1: Combined Training and Inference

Upload a single file containing both historical applications (with known outcomes - target column populated) and recent applications (requiring risk scores - target column NULL). The platform automatically trains on historical data and generates predictions for recent applications.

Mode 2: Separate Training and Inference

- First upload: Provide historical loan applications with known default outcomes. The platform trains the model and uses the 2% most recent records for initial inference to populate the dashboard.
- Subsequent uploads: Provide new loan applications with null target values. The platform applies the trained model to generate risk scores without retraining.

When analyzing credit risk predictions and portfolio analytics in the dashboard, interpret default probabilities in the context of your specific default definition and DPD threshold to support appropriate lending decisions and risk management strategies.

6.5 Inventory Optimization Use Case

Data Granularity and Structure

Data must be uploaded at daily granularity. Each row represents one product-store-date combination. The composite key `date + store_id + product_id` must be unique within a file. Historical data spanning at least 60 days is recommended; the platform sets aside the most recent 28 days as the inference slice and uses the remainder for training.

Understanding the Forecast Target

The model is a regression task that predicts `units_sold` — the number of units expected to be sold for a given product at a given store on a future date. It is not a binary classifier (there is no stockout flag or



probability score). The output column in the predictions table is `predicted_units_sold`. Operations teams can use this value alongside `inventory_level` and `units_ordered` to determine whether a reorder is needed before the next period.

How the Platform Handles Training

The Glue ETL job automatically engineers the following features before training: temporal signals (`day_of_week`, `month`, `year`, `is_weekend`), rolling sales averages (7-day and 30-day), lag features (prior-day and prior-week units sold), and inventory health ratios (`days_of_stock_remaining`, `inventory_turnover_rate`, `sales_velocity_change_pct`). The model uses LightGBM (gradient-boosted regression). Feature selection retains the top 40% of features by LightGBM importance score. Training uses the full history older than the most recent 56 days (28-day validation window + 28-day test/inference window). Records are not excluded for zero-stock periods; all rows in the upload are used.

6.6 Viewing Predictions

After processing completes:

Step	Action
1	Open QuickSight Console .
2	Navigate to Dashboards .
3	Select the appropriate Prism dashboard for your use case: • Prism Customer Churn Dashboard (Churn Prediction) • Prism Readmission Risk Dashboard (Readmission Risk) • Prism Credit Risk Dashboard (Credit Risk Scoring)
4	View the available sheets containing analytics and predictions.

6.7 Querying Data with Athena

You can run custom SQL queries against your data:

Step	Action
1	Go to Amazon Athena Console .
2	Select the workgroup prism-workgroup .
3	Select database prism_analytics_db .
4	Run queries against available tables based on your use case: • Churn: <code>churn_historical</code>, <code>churn_predictions</code> • Readmission Risk: <code>readmission_risk_historical</code>, <code>readmission_risk_predictions</code> • Credit Risk: <code>credit_risk_historical</code>, <code>credit_risk_predictions</code>



7. QuickSight Dashboards

7.1 Accessing the Dashboard

Step	Action
1	Navigate to QuickSight Console .
2	Click Dashboards in the left navigation.
3	Select the dashboard for your use case: • Prism Customer Churn Dashboard • Prism Readmission Risk Dashboard • Prism Credit Risk Dashboard • Prism Inventory Optimization Dashboard

Each dashboard contains multiple sheets with analytics and prediction insights specific to the use case.

7.2 Dashboard Components by Use Case

7.2.1 Churn Prediction Dashboard

The churn prediction dashboard includes:

Churn Risk Overview

- Total customers analyzed.
- Customers at high risk (>70% probability).
- Customers at medium risk (40-70% probability).
- Customers at low risk (<40% probability).

Risk Distribution Chart

- Visual breakdown of churn risk categories.
- Trend over time (if historical data available).

Top Risk Factors

- Key features contributing to churn.
- Feature importance visualization.

Customer List

- Sortable table of all customers.
- Risk score for each customer.
- Filter by risk level, contract type, etc.

7.2.2 Hospital Readmission Risk Dashboard



The hospital readmission risk dashboard includes:

Readmission Risk Overview

- Total encounters analyzed
- Patients at high risk (>70% probability)
- Patients at medium risk (40-70% probability)
- Patients at low risk (<40% probability)

Risk Distribution Chart

- Visual breakdown of readmission risk categories
- Trend over time by admission type and diagnosis

Top Risk Factors

- Key clinical and social determinants contributing to readmission
- Feature importance visualization

Patient Encounter List

- Sortable table of all encounters
- Risk score for each patient
- Filter by risk level, admission type, comorbidities, payer, etc.

7.2.3 Credit Risk Dashboard

The credit risk dashboard includes:

Default Risk Overview

- Total applications analyzed
- Applications at high risk (>70% probability)
- Applications at medium risk (40-70% probability)
- Applications at low risk (<40% probability)

Risk Distribution Chart

- Visual breakdown of default risk categories
- Trend over time by loan purpose and applicant segment

Top Risk Factors

- Key credit bureau, transactional, and alternative signals contributing to default
- Feature importance visualization

Application List

- Sortable table of all loan applications
- Risk score for each application
- Filter by risk level, loan purpose, income range, employment status, etc.



7.2.4 Inventory Optimization Dashboard

Inventory Overview (KPI tiles)

- Total active SKU-store combinations in the current period
- SKUs where days_of_stock_remaining < lead_time_days (high stockout risk)
- SKUs where days_of_stock_remaining > 60 (excess inventory)
- Model forecast accuracy — MAPE across all SKU-store combinations in the inference window

Demand Forecast Chart

- Line chart: predicted_units_sold vs actual units_sold over the most recent 90 days
- Filterable by category, store_id, product_id, region

Inventory Health List

- Sortable table of all SKU-store combinations
- Columns: store_id, product_id, category, inventory_level, days_of_stock_remaining, predicted_units_sold, inventory_turnover_rate
- Filterable by risk level (high / normal / excess), store_id, category

7.3 Creating Custom Visuals

While the default visuals on the Dashboard are generated automatically for required columns, you can also create your own custom visuals directly from the QuickSight interface, and it's simple to do.

To create a custom visual:

Step	Action
1	Open the Analyses tab from the QuickSight side panel and choose the appropriate analysis: • Prism Customer Churn Analysis • Prism Hospital Readmission Risk Analysis • Prism Credit Risk Analysis • Prism Inventory Optimization Analysis
2	Click Add Visual from the Insert tab and choose a chart type from the visuals panel onto the sheet.
3	Select your desired dimensions and measures from the available fields.
4	Customize the chart type, colours, and labels as needed.
5	Click Done editing to save your visual to the sheet.
6	Publish the Analysis as Dashboard.



This gives you full flexibility to explore the data beyond the default views, whether you want to compare additional columns, try different chart types, or build visuals tailored to your specific analysis needs.

7.4 Sharing Dashboards

Share with specific users:

Step	Action
1	Open the dashboard.
2	Click Share icon in the top right.
3	Enter email addresses of QuickSight users.
4	Set permission level (Viewer or Co-owner).
5	Click Share .

Publish to larger audience:

Step	Action
1	Click Share → Publish dashboard .
2	Configure sharing settings.
3	Copy the dashboard link to distribute.

7.5 Natural Language Querying with QuickSight Q Topics

Prism automatically creates a QuickSight Q Topic as part of the analytics pipeline. A Q Topic enables natural language querying instead of navigating dashboard filters, you can ask questions about your churn data in plain English directly within the QuickSight console.

Example Questions by Use Case

Churn Prediction:

- "How many customers churned last month?"
- "What is the average monthly bill for churned customers?"
- "Show me customers with fiber optic internet who are on month-to-month contracts"
- "What is the total lifetime value of customers by payment method?"

Hospital Readmission Risk:

- "How many patients were readmitted in the last 30 days?"



- "What is the average Charlson score for high-risk patients?"
- "Show me emergency admissions with heart failure diagnosis"
- "What percentage of patients with diabetes were readmitted?"

Credit Risk:

- "How many loan applications have high default risk?"
- "What is the average bureau score for low-risk applications?"
- "Show me applications with debt-to-income ratio above 0.5"
- "What is the default rate by loan purpose?"

Inventory Optimization:

- "Which products are at risk of stockout in the next 7 days?"
- "What is the recommended reorder quantity for product X at store Y?"
- "Show me products with more than 60 days of excess inventory"
- "What is the forecast accuracy by product category?"
- "Which stores have the highest inventory turnover rate this month?"
- "What was the predicted versus actual units sold for SKU X last week?"

To Access Q topics follow the steps:

Step	Action
1	Open the QuickSight Console.
2	Click the Q search bar at the top of the page (or navigate to Q Topics from the left navigation).
3	Select the appropriate Q Topic: • Customer Churn Analysis Topic • Hospital Readmission Analysis Topic • Credit Risk Analysis Topic • Inventory Optimization Topic
4	Type a question in plain English and press Enter.
5	QuickSight Q returns an answer as a visual (chart, table, or number) that you can pin to a dashboard.

8. Monitoring

8.1 CloudWatch Monitoring

View Processing Metrics:



1. Go to **CloudWatch Console**.
2. Navigate to **Metrics**.
3. Find metrics under:
 - AWS/Lambda.
 - AWS/Glue.
 - AWS/States.

8.2 Logs

Lambda Logs:

Step	Action
1	Go to CloudWatch Console .
2	Navigate to Logs → Log groups .
3	Find log groups starting with /aws/lambda/prism-.

Glue Job Logs:

Step	Action
1	Go to AWS Glue console.
2	Click ETL Jobs.
3	Find the Glue jobs for your workflow: • Churn: prism-etl-churn-job, prism-churn-ingest-predictions-job • Readmission Risk: prism-etl-readmission_risk-job, prism-readmission-ingest-predictions-job • Credit Risk: prism-etl-credit_risk-job, prism-credit-risk-ingest-predictions-job • Inventory Optimization: prism-etl-inventory_optimization-job, prism-inventory-optimization-ingest-predictions-job
4	Click on runs -> logs to see execution logs.

Step Functions Logs:

Step	Action
1	Go to Step Functions Console .
2	Select the state machine (prism-processing-workflow or prism-ml-workflow).



3	Click on an execution to view details and logs.
---	---

SageMaker Logs:

You can view SageMaker logs in three ways:

Via SageMaker Pipeline:

Step	Action
1	Go to SageMaker AI Console .
2	Open SageMaker Studio (a domain is required to access SageMaker Studio, refer to section 5.6)
3	Navigate to Pipelines .
4	Select your pipeline execution (with the prefix prism).
5	Click on any step to view its logs.

Via Individual Jobs:

Step	Action
1	Go to SageMaker AI Console .
2	Navigate to Training → Training jobs or Processing → Processing jobs .
3	Select the job.
4	Click View logs to open CloudWatch logs.

Via Cloudwatch Log group:

Job Type	Cloudwatch Log Group Path
Processing Jobs	/aws/sagemaker/ProcessingJobs
Training Jobs	/aws/sagemaker/TrainingJobs
Batch Transform	/aws/sagemaker/TransformJobs

9. Security Configuration

9.1 Encryption

S3 Encryption:



- All S3 buckets use Server-Side Encryption (SSE-S3) by default.
- You can upgrade to SSE-KMS for additional control.

DynamoDB Encryption:

- Tables are encrypted at rest using AWS managed keys.

In Transit:

- All data transfer uses TLS 1.2+.

9.2 IAM Security

Least Privilege:

- All IAM roles follow the principle of least privilege.
- Roles are scoped to specific resources where possible.

Lambda Function Isolation:

- Each Lambda function has its own dedicated IAM execution role.
- Roles are individually scoped to only the permissions that specific function requires.
- Note: The QuickSight role is shared between inference and analytics functions for dashboard operations.

Service Roles:

- Glue jobs use a dedicated service role with ETL-specific permissions.
- SageMaker uses a specific ML role for training and inference.
- Each Step Functions state machine has its own dedicated orchestration role.

9.3 Network Security

VPC Security (when using CREATE_NEW or USE_EXISTING):

- Private subnets for compute resources.
- VPC endpoints for AWS service access.
- Security groups restrict traffic.

S3 Access:

- Bucket policies restrict access to authorized principals.
- Block public access is enabled.

9.4 Access Control

QuickSight Access:

- Users must be invited to view dashboards.
- Row-level security can restrict data access.



Athena Access:

- Workgroup permissions control query access.
- IAM policies restrict database/table access.

10. Cost Management

10.1 Cost Components

The following are some of the major cost components. This is not an exhaustive list:

Service	Cost Driver	Typical Range
S3	Storage + requests	\$0.023/GB/month
Glue	DPU-hours	\$0.44/DPU-hour
Athena	Data scanned	\$5/TB scanned
SageMaker	Instance hours	Varies by instance
Lambda	Invocations + duration	\$0.20/million requests
QuickSight	User licenses	\$9-18/user/month
Step Functions	State transitions	\$0.025/1000 transitions
VPC	NAT Gateway + data transfer	\$0.045/hour + data costs
VPC Endpoints	Per endpoint per AZ	\$0.01/hour per endpoint
EC2 (Installer)	Instance hours	Temporary, terminates after deployment
DynamoDB	Read/write capacity	Pay-per-request

Note: Using `NatGateways = 0` with VPC Endpoints is recommended to reduce networking costs.

10.2 Cost Monitoring

To ensure you are only monitoring costs associated with this application, use the standardized resource tags during budget creation:

1. **Open the Billing Console:** Navigate to the [AWS Budgets](#) dashboard.
2. **Create Budget:** Click **Create budget**, select **Cost budget**, and click **Next**.
3. **Define Budget Amount:** Enter a name (e.g., Prism_Monthly_Budget) and set your desired period and amount.
4. Filter by Prism Tag (Critical Step).
5. Define your notification triggers (e.g., notify when actual costs exceed 80% of the budget) and add your email address.



Use Cost Explorer:

1. Go to **Cost Explorer**.
2. Filter by service or tag to see costs for S3, Glue, SageMaker, QuickSight, etc.
3. Analyse costs over time.

11. Troubleshooting

11.1 Common Issues

Data Not Processing

Symptoms: Uploaded file not processed, no Step Functions execution.

Solutions:

1. Check if EventBridge rules are enabled.
 - Go to **EventBridge Console**.
 - Verify all “prism” prefix rules are in ENABLED state.

☐	Name	Status	Type	Event bus
☐	prism-etl-to-qs-rule	✔ Enabled	Standard	default
☐	prism-glue-to-ml-rule	✔ Enabled	Standard	default
☐	prism-inference-to-qs-rule	✔ Enabled	Standard	default
☐	prism-s3-to-sqs-rule	✔ Enabled	Standard	default
☐	Prism-PlatformStack-M-CodeBuildCompleteRule80B9-VS2rfBQjVRh8	✔ Enabled	Standard	default

Note: The prefix may change depending on the stack name that is configured during installation.

2. Check Lambda invocation logs.
 - Go to **CloudWatch Console** → **Logs**.



- Check `/aws/lambda/prism-schema-reader` for checking the errors. Some common errors could be:
 - Your data failed to fulfill the schema requirements
 - Your license expired
 - Model unavailable: rate limit exceeded

Glue Job Failures

Symptoms: Glue job timed out. Glue job failed.

Solutions:

1. Check Glue job error logs in CloudWatch.
2. Verify input data format matches expected datatypes specified in the schema .

QuickSight Dashboard Empty

Symptoms: Dashboard shows no data or loading forever.

Solutions:

1. Verify Athena can query the data tables.
2. Check QuickSight SPICE dataset refresh status.
3. Refresh the dataset manually:
 - Go to **QuickSight** → **Datasets**
 - Select dataset → **Refresh now**

SageMaker Training Failures

Symptoms: ML model training fails.

Solutions:

1. Check SageMaker training job logs.
2. Verify sufficient instance quota.
3. Check training data volume is adequate.
4. Review hyperparameters for validity.

11.2 Getting Support

If issues persist:

1. Gather relevant logs and error messages.
2. Note the timestamps when the issue occurred.
3. Contact support (see Section 15).



12. Upgrades & Updates

12.1 Update Process

Before updating:

Review release notes for breaking changes.

Update via AWS Marketplace:

1. Go to **AWS Marketplace Subscriptions**.
2. Find your Prism subscription.
3. Click **Update Stack** when a new version is available.
4. Review parameter changes.
5. Execute stack update.

12.2 Rollback Procedure

If an update causes issues:

1. Go to **CloudFormation Console**.
2. Select the **Prism installer** stack.
3. Click **Stack actions** → **Roll back**.
4. Wait for rollback to complete.
5. Verify system functionality.

13. Backup & Recovery

13.1 Data Backup

S3 Versioning:

All S3 buckets have versioning enabled by default. To recover deleted files:

1. Go to **S3 Console**.
2. Navigate to the bucket.
3. Click **Show versions** toggle.
4. Find the version you want to restore.
5. Download or copy the previous version.



DynamoDB Point-in-Time Recovery:

Schema registry supports point-in-time recovery:

1. Go to **DynamoDB Console**.
2. Select the schema registry table.
3. Click **Backups** tab.
4. Use **Restore to point-in-time** to restore.

13.2 Iceberg Time Travel

Apache Iceberg supports time travel queries, allowing you to access historical data. You can query data as it existed at a specific point in time using Athena:

```
-- Query data as of a specific snapshot
SELECT * FROM prism_analytics_db.churn_historical
FOR VERSION AS OF <snapshot-id>
```

To find available snapshots along with timestamps:

```
SELECT snapshot_id, committed_at
FROM "prism_analytics_db"."churn_historical$snapshots"
ORDER BY committed_at DESC
```

For other use cases, replace table names accordingly:

- Readmission Risk: readmission_risk_historical\$snapshots
- Credit Risk: credit_risk_historical\$snapshots

14. Uninstallation

14.1 Before Uninstalling

Critical: Review what data you want to retain before uninstalling.

1. Save QuickSight dashboards (export as PDF or republish to another account).
2. Cancel AWS Marketplace subscription to stop billing.

Note: The raw S3 bucket has a retention policy and will not be deleted automatically, so your uploaded data will be preserved.

14.2 Delete CloudFormation Stacks

You need to delete two CloudFormation stacks:



Delete the Installer Stack:

- 1. Go to **CloudFormation Console**.
- 2. Find and select the installer stack (the name you chose during deployment).
- 3. Click **Delete Stack**.
- 4. Confirm deletion.
- 5. Wait for deletion to complete.

Delete the PrismPlatform Stack:

- 1. Go to **CloudFormation Console**.
- 2. Select the “PrismPlatformStack” stack (created by the installer).
- 3. Click **Delete Stack**.
- 4. Confirm deletion.
- 5. Wait for deletion to complete (monitor status).

Note: **PrismPlatformStack** is the parent stack and initiating its deletion will automatically trigger the removal of all associated nested stacks.

14.3 Verify Complete Removal

- 1. Go to **Resource Groups Console**.
- 2. Search for resources with names starting with prism-*
- 3. Verify no resources remain.

Note: Raw S3 Bucket created by the stack will need to be manually removed by the user.

15. Support

15.1 Support Channels

Channel	Use Case	Response Time
Documentation	Self-service help	Immediate
AWS Marketplace Support	Subscription issues	24-48 hours

15.2 Contacting Support



10Pearls Support:

- Email: aws-platform-support@10pearls.com
- Website: <https://10pearls.com>

AWS Marketplace Support:

- Via AWS Console: Support Centre.

15.3 Information to Include

When contacting aws support, provide:

1. AWS Account ID (not credentials).
2. Error messages and timestamps.
3. Steps to reproduce the issue.
4. CloudWatch logs (relevant excerpts).
5. Screenshots of the issue (if applicable).

15.4 Self-Service Resources

Resource	Location
This User Guide	Included with product
EULA	EULA
AWS Marketplace Listing	Per subscription
AWS Documentation	docs.aws.amazon.com

Release Notes

Version 1.0.0 (April 2026)

Initial Release

- Telecom Churn Prediction use case.
- Apache Iceberg data Lakehouse.
- AI-powered schema inference (Amazon Bedrock).
- SageMaker ML pipelines.
- QuickSight programmatic dashboards.
- Event-driven processing architecture.
- VPC deployment options (default, new, existing).



Version 1.1.0 (May 2026)

New Use Cases

- Hospital Readmission Risk workflow: Predict 30-day readmission risk for hospital patients at discharge, powered by clinical vitals, lab results, comorbidities, and social determinants.
- Credit Risk Scoring workflow: Assess borrower default probability using credit bureau data, transactional behaviour, and alternative signals.

Infrastructure

- Bedrock schema inference updated to Meta Llama 4 Maverick; QuickSight Q synonym generation uses Amazon Nova Pro.

Appendix A: Prism System Configurations

This section consolidates all configuration options for Prism.

A.1 Core Configuration

Parameter	Type	Default	Description
QuickSightUsername	String	<i>required</i>	QuickSight user email for dashboard ownership
QuickSightRoleName	String	aws-quicksight-service-role-v0	QuickSight IAM service role
AdminEmail	String	<i>required</i>	Email address for SNS operational alerts

A.2 Lambda Configuration

Setting	Value	Description
Runtime	Python 3.12	Lambda runtime environment
Architecture	x86_64	Processor architecture
Memory	256-512 MB	Memory allocation (varies by function)
Timeout	30-900 seconds	Function timeout (varies by function)
Tracing	Active	X-Ray tracing enabled
VPC	Configured	Runs in VPC when using CREATE_NEW or USE_EXISTING modes

A.3 Training Dataset Size Tiers

Select the appropriate tier based on your dataset size. This automatically configures the correct instance types:



Tier	Data Size	Processing Instance	Training Instance	Inference Instance
SM (Small)	Under 500 MB	ml.m5.xlarge	ml.m5.xlarge	ml.m5.xlarge
MD (Medium)	500 MB - 2 GB	ml.m5.xlarge	ml.m5.2xlarge	ml.m5.2xlarge
LG (Large)	2 GB - 5 GB	ml.m5.2xlarge	ml.m5.4xlarge	ml.m5.4xlarge

Important: Selecting a tier that's too small for your data will cause the deployment to fail. When in doubt, choose the next tier up.

A.4 Compute Configuration

Parameter	Type	Default	Description
GlueWorkerType	String	G.1X	Glue worker size (G.1X, G.2X, G.4X, G.8X)
GlueMaxWorkers	Number	5	Maximum Glue workers for auto-scaling
InferenceInstanceType	String	ml.m5.large	SageMaker batch inference instance (auto-set by tier)
TrainingInstanceType	String	ml.m5.xlarge	SageMaker training instance (auto-set by tier)

A.5 Scheduling Configuration

Parameter	Type	Default	Description
ScheduleExpression	String	rate(1 hour)	EventBridge schedule expression for batch pipeline
SchedulerEnabled	String	true	Enable/disable scheduled processing

A.6 Model Quality Configuration

Parameter	Type	Default	Description
ChurnMinimumAUC	Number	0.80	Minimum AUC threshold for Churn Prediction model deployment (range: 0.50-0.99).
ReadmissionRiskMinimumAUC	Number	0.65	Minimum AUC threshold for Readmission Risk model deployment (range: 0.50-0.99).



CreditRiskMinimumAUC	Number	0.65	Minimum AUC threshold for Credit Risk model deployment (range: 0.50-0.99).
InventoryOptimizationMaximumMAPE	Number	0.7	Maximum acceptable MAPE for the Inventory Optimization model. The model is rejected and not deployed if MAPE on the validation set exceeds this value (default 70%). Lower is better.

A.7 VPC Configuration

Parameter	Type	Default	Description
VpcDeploymentMode	String	USE_DEFAULT	CREATE_NEW, USE_EXISTING, or USE_DEFAULT
VpcCidr	String	10.0.0.0/16	VPC CIDR block (for CREATE_NEW)
NatGateways	Number	0	Number of NAT gateways. Set to 0 to use VPC Endpoints (lower cost). Set to 1 for NAT Gateway.
MaxAzs	Number	1	Maximum availability zones (for CREATE_NEW, keep in us-east-1)
VpcId	String	-	Required for USE_EXISTING mode
AvailabilityZones	String	-	Required for USE_EXISTING mode (comma-separated list)
PrivateSubnetIds	String	-	Required for USE_EXISTING mode (comma-separated list)
RouteTableIds	String	-	Required for USE_EXISTING mode (comma-separated list)

A.8 S3 Configuration

Setting	Default	Description
Versioning	Enabled	All buckets have versioning enabled for data protection
Encryption	SSE-S3	Server-side encryption enabled by default
Lifecycle	Configured	Older versions are transitioned to cheaper storage classes

A.9 Glue Configuration

Setting	Value	Description
---------	-------	-------------



Job Name	prism-etl-{use-case}-job	ETL job for processing churn analytics data
Type	Spark	Spark-based ETL job
Glue Version	5.0	Supports Spark 3.5, Scala 2, Python 3
Language	Python 3	Python runtime
Worker Type	G.1X	4 vCPU and 16GB RAM per worker
Auto-scaling	Enabled	Dynamically scales workers up and down
Maximum Workers	5	Maximum number of workers when auto-scaling
Job Timeout	30 minutes	Maximum job runtime
Number of Retries	1	Automatic retry on failure
Maximum Concurrency	1	Only one job run at a time
Job Bookmark	Disabled	Does not track previously processed data
Data Lake Format	Iceberg	Apache Iceberg table format

A.10 SageMaker Configuration

Setting	Default	Description
Model Framework for Churn	XGBoost	ML algorithm used for prediction
Model Framework for Readmission Risk	LightGBM	ML algorithm used for prediction
Model Framework for Credit Risk	LightGBM	ML algorithm used for prediction
Model Framework for Inventory Optimization	LightGBM	ML algorithm used for prediction

A.11 QuickSight Configuration

Setting	Default	Description
Dataset Refresh	On data update	Datasets refresh when new data is processed
SPICE	Enabled	In-memory acceleration for fast queries
Row-Level Security	Optional	Can be configured for multi-tenant deployments



© 2026 10Pearls. All rights reserved.

For licensing terms, see EULA



<https://www.10pearls.com>



info@10pearls.com