

Oracle Warehouse Management Cloud

AIML Guide

Release 26C



Oracle Warehouse Management Cloud
AIML Guide

Release 26C

G55068-02

Copyright © 2026, Oracle and/or its affiliates.

Author: Oracle WMS Cloud Product Team

Contents

Get Help	i
1 Oracle WMS Machine Learning	1
WMS Machine Learning	1
Predicting Order Cycle Time, Processing Time, and Waiting Time	1
Intelligent Cycle Counting	2
Machine Learning Process Flow	3
Types of Algorithms Available in WMS	3
Creating a Machine Learning Model in WMS	5
2 AIML Predictive Fulfillment Dashboard	13
AIML Predictive Fulfillment Dashboard	13
Inbound and Outbound Performance Metrics	14
3 GenAI Features	17
Intelligent Wave Summary	17
4 AI Agents	21
AI Agent Studio Support in WMS	21
AI Agent Studio and Fusion-Specific Configuration	21
AI Agent Configuration in WMS	23
AI Agent: Wave Research Advisor	26
AI Agent-Inventory Expiry Assistant	31
AI Agent: Task Management Assistant	32
5 AI Agentic Apps	35
Logistics Execution Command Center	35
6 Index	45
Index	45

Get Help

There are a number of ways to learn more about your product and interact with Oracle and other users.

Get Help in the Applications

Some application pages have help icons  to give you access to contextual help. If you don't see any help icons on your page, click your user image or name in the global header and select Show Help Icons. If the page has contextual help, help icons will appear.

Get Training

Increase your knowledge of Oracle Cloud by taking courses at [Oracle University](#).

Join Our Community

Use [Cloud Customer Connect](#) to get information from industry experts at Oracle and in the partner community. You can join forums to connect with other customers, post questions, suggest [ideas](#) for product enhancements, and watch events.

Share Your Feedback

We welcome your feedback about Oracle Applications user assistance. If you need clarification, find an error, or just want to tell us what you found helpful, we'd like to hear from you.

You can email your feedback to oracle_fusion_applications_help_ww_grp@oracle.com.

Thanks for helping us improve our user assistance!

1 Oracle WMS Machine Learning

WMS Machine Learning

Accurately predicting how long it takes to process an order, from start to finish, is a growing challenge in today's fast-paced, data-driven environment. Oracle WMS now includes integrated AI and machine learning (ML) capabilities that help solve this problem by analyzing your historical order data and generating smart, reliable predictions.

We currently support machine learning models for:

- **Predicting Order Cycle Time, Waiting Time, and Processing Time**
- **Intelligent Cycle Counting**
- **Market Basket Analysis**

These capabilities help you optimize planning, reduce operational delays, and better manage customer expectations by using insights drawn directly from your own data.

Predicting Order Cycle Time, Processing Time, and Waiting Time

In modern supply chain environments, **predictability is power**. Warehouses often operate with limited visibility into how long an order will take to move from receipt to shipment. This lack of foresight can create planning inefficiencies, staffing issues, customer dissatisfaction, and missed service level agreements (SLAs).

By leveraging **AI/ML-driven predictions for Order Cycle Time, Processing Time, and Waiting Time**, businesses can make smarter, data-backed decisions that drive operational efficiency and customer satisfaction.

Use Case Scenario

A high-volume distribution center is preparing for a seasonal surge in order volume. Operations managers need to forecast how long orders will take to complete, not just on average, but across different order types, times of day, and warehouse zones.

Using Oracle WMS's predictive capabilities:

- **Order Cycle Time** predictions help determine when specific orders are likely to leave the facility, which enables accurate delivery promises and staffing adjustments.
- **Processing Time** forecasts highlight potential internal delays (e.g., picking, packing) and allow managers to preemptively redistribute labor to keep workflows balanced.
- **Waiting Time** predictions identify bottlenecks, such as staging or dock congestion, before they occur, providing the ability to adjust load schedules or sequence orders more efficiently.

Business Value

- **Proactive Labor Planning:** Allocate resources to high-delay areas before issues arise.
- **Customer Experience:** Provide accurate delivery estimates and improve on-time performance.
- **Throughput Optimization:** Identify and eliminate workflow bottlenecks in real-time.
- **SLA & Carrier Compliance:** Minimize late shipments by predicting orders at risk and taking action earlier.
- **Cost Reduction:** Avoid overstaffing and overtime by aligning labor with actual processing needs.

By embedding predictive intelligence into core WMS workflows, businesses can shift from reactive firefighting to proactive decision-making.

Intelligent Cycle Counting

Maintaining accurate inventory is foundational to successful warehouse operations, however, traditional cycle counting methods are time-consuming, labor-intensive, and often lack strategic prioritization. Many warehouses rely on fixed schedules, ABC classifications, or random sampling, which can lead to overcounting low-risk items and undercounting high-risk or high-impact SKUs.

Intelligent Cycle Counting, powered by AI/ML in Oracle WMS, transforms this process by using data-driven predictions to determine **which items are most likely to be inaccurate** AND prioritizes those for counting.

Use Case Scenario

A warehouse handles thousands of SKUs with varying turnover rates, locations, and movement histories. Rather than using static ABC logic, the team enables Intelligent Cycle Counting to dynamically identify inventory at higher risk of inaccuracy such as:

- Items with frequent adjustments or discrepancies
- SKUs recently involved in multiple picks or returns
- Inventory stored in high-traffic or hard-to-reach locations

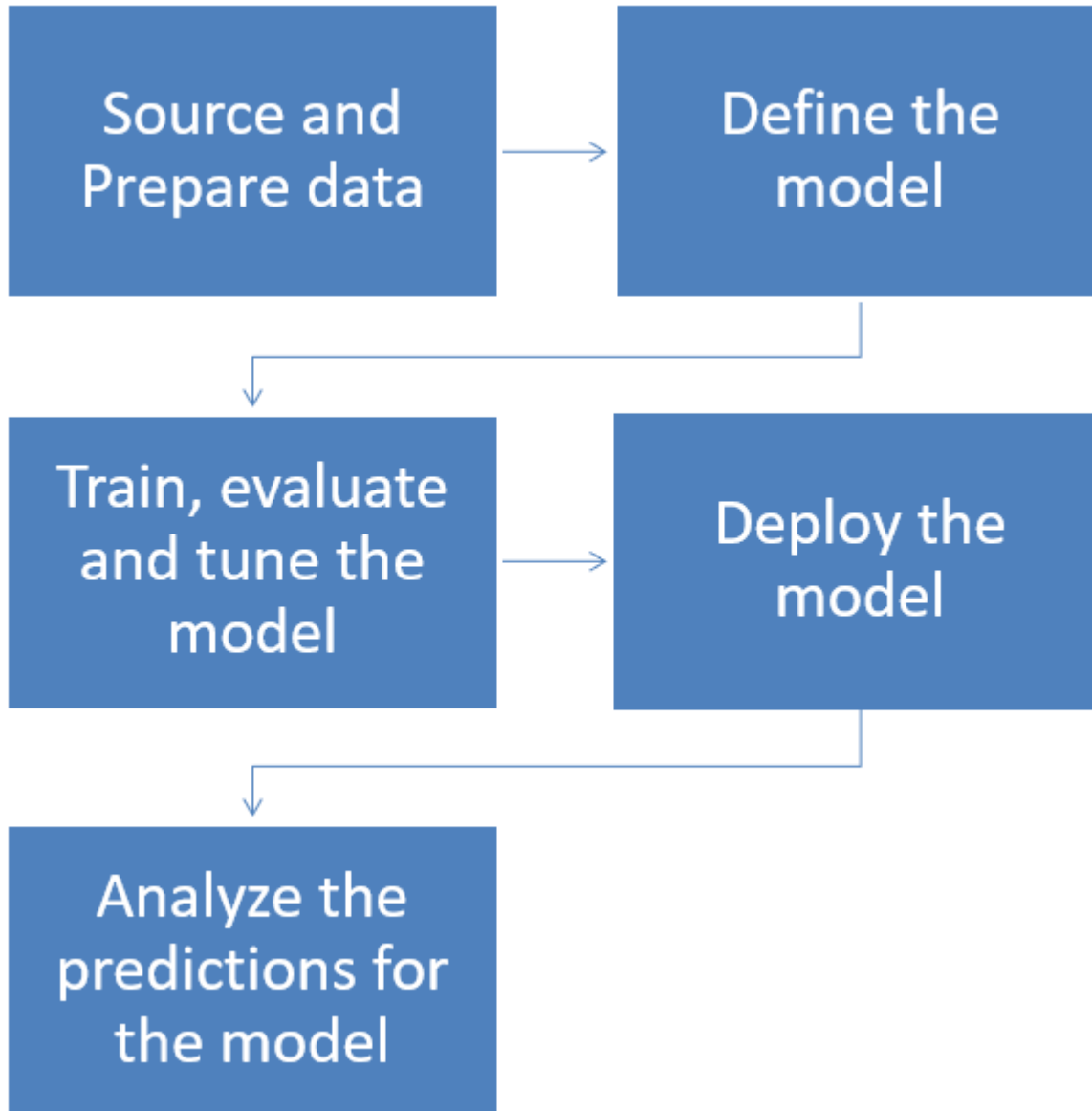
By running a scheduled AI/ML job, the WMS will recommend a prioritized list of items to count based on real-time inventory behavior, historical accuracy trends, and movement data.

Business Value

- **Higher Accuracy with Less Effort:** Count fewer items while improving overall inventory accuracy.
- **Data-Driven Prioritization:** Focus counting efforts where risk is highest, not just based on item value.
- **Labor Optimization:** Reduce unnecessary cycle counts, freeing up staff for other critical tasks.
- **Shrinkage Reduction:** Catch and correct inventory issues early
- **Audit Confidence:** Strengthen audit readiness with continuous, intelligent verification.

Intelligent Cycle Counting shifts inventory management from **scheduled routines** to **smart, predictive actions**, helping businesses stay lean, accurate, and responsive without adding manual overhead.

Machine Learning Process Flow



Types of Algorithms Available in WMS

Oracle WMS currently supports a variety of **machine learning algorithms**, each designed to handle specific types of prediction problems. While each algorithm requires different parameters and follows a unique training methodology, many of them can be applied across multiple supported metrics.

These algorithms form the foundation for predictive capabilities such as forecasting **Order Cycle Time**, **Waiting Time**, **Processing Time**, and powering features like **Intelligent Cycle Counting** and **Market Basket Analysis**.

Random Forest

Random Forest is a widely used algorithm for solving both **regression** and **classification** problems (see glossary for definitions). It operates by building multiple decision trees and combining their outputs for improved accuracy and robustness.

Key Parameters:

- **n_estimators**

The number of decision trees in the model. More trees can increase accuracy but also add processing time.

- **Default:** 100
- **Allowed Range:** 1–200

- **max_depth**

Specifies how deep each decision tree can grow. Deeper trees can capture more complexity in the data but also risk overfitting.

- **Allowed Range:** 1–20
- **Note:** No default value is set. Entering a value outside the range will trigger an error.

- **max_depth**

- This value should not be changed.

Training Tip: Run several iterations to train and evaluate the model. This helps fine-tune parameters and improves overall accuracy by capturing variations in results.

Feed Forward Neural Networks

Feed Forward Neural Networks are inspired by the human brain and excel at tasks that involve **pattern recognition** such as identifying images, text, or sounds. In WMS, they are typically used for **classification** and **clustering** tasks.

Key Parameters:

- **hidden_layer_sizes**

Determines the architecture of the network—the number of hidden layers and the number of neurons (nodes) in each.

- **max_iter**

Indicates the number of training iterations, or **epochs**.

Epoch: One complete pass through the entire training dataset.

Gradient Boosting

Gradient Boosting is a powerful and flexible algorithm that can be used for both **classification** and **regression**. It builds models in a sequential way, where each new model corrects errors made by the previous one.

Key Parameters:

- **max_iter**

The number of boosting stages or epochs to run.

- **max_depth**

Controls the depth of individual decision trees (same behavior as in Random Forest).

- **learning_rate**

Determines how quickly the model adapts to the data.

- A **high learning rate** may overshoot the optimal solution.
- A **low learning rate** may take longer but offers more control and stability.

Note: Results will vary based on the model and parameters given.

Creating a Machine Learning Model in WMS

To build and use a machine learning model in WMS, follow the structured process below:

Step 1: Create a Scheduled Job

Initiate the process by scheduling a job to prepare the necessary training data.

Step 2: Review the Training Data

Use the **AIML Training Data** screen to review and validate the dataset generated by the scheduled job.

Step 3: Define the Model

Configure your model using parameters and filters. Assign the model to an **AI/ML Training Template**, where you can also set a default template for reuse. This can be done in AIML Training Template UI.

Step 4: Run Predictions

Execute the model using the **AIML Prediction Run UI** to generate predictions. The system will capture and display the results from this run.

Step 5: Deploy and Refine the Model

Deploy the model to begin using it in production. You can also retrain and adjust the model as needed based on results and performance insights.

Available AI/ML Screens in WMS

- **AIML Training Data:** View the structured training dataset prepared via scheduled job.
- **AIML Training Template:** Create, configure, and manage training templates; set defaults for ease of reuse.
- **AIML Training Date Rule Hdr: Create custom seasonality based date ranged as need for you warehouse needs. This only applies for Market Basket Analysis.**
- **AIML Models UI:** View models created from your training templates, including parameters and status.
- **AIML Model's Training Logs UI: View relevant logging information related to model creation.**
- **AIML Prediction Run UI:** Run models and view detailed prediction results.
- **AIML Predictive Dashboard:** Visualize recent prediction outcomes (currently supports Order Cycle Time only).

Step One: Create a Scheduled Job

Each metric is configured to work with specific AI/ML algorithms and defines the type of input data required, as well as the expected output.

Currently, we support AI/ML functionality for the following metrics:

- **Order Cycle Time**
- **Order Waiting Time**
- **Order Picking Time**
- **Intelligent Cycle Count**

Using the Scheduled Jobs UI

You can schedule the generation of training data for your AI/ML models directly through the **Scheduled Jobs UI**.

In the **CRUD pane**, users will find a dropdown menu where they can select the appropriate job type. To initiate the data generation process, choose:

"Generate Training Data for [Metric]"

This action creates a scheduled job that prepares the data needed for model training.

Job Number *

Job Type *

- Generate IHT by billing location type
- Generate Inventory Balance Snapshot
- Generate Inventory History Extract
- Generate Inventory Summary
- Generate LPN Modes
- Generate Market Basket Analysis Run for Putaway
- Generate OB LPN Billing Report
- Generate OB Load Files
- Generate Order Files
- Generate Parcel Manifest Files
- Generate Prediction Run for Intelligent Cycle Counting
- Generate Prediction Run for Order Cycle/Waiting/Processing time
- Generate Verify Shipment Alert
- Generate training data for Intelligent Cycle Counting
- Generate training data for Market Basket Analysis
- Generate training data for Order Cycle/Waiting/Processing time

Selecting one of the above job types from the dropdown will expand additional **Job Parameters**, which must be configured before scheduling the job:

- **Username** (*mandatory*): Enter a valid WMS login. This identifies the user associated with the job.
- **From Date**: The starting date for the data to be included in training.
- **To Date**: The ending date for the data range.
- **Past Number of Days**: Alternatively, you can specify a rolling window of past days to pull data from (e.g., last 30 days).

Note: Either the **From/To Date range** (mm/dd/yy) or the **Past Number of Days** (must be in format -x) must be provided. These values determine the time window used to generate training data for the selected metric.

Note: Each of the jobs will not be enabled by default. If you attempt to save without any value for username, you will receive an error message.

The output (model) of the job (different for each metric) will be captured in the **AIML Models UI**. One for each job run.

- Each record thus generated is defined by a unique composite key (business key).
- The job will add, or update records based on the key.

For every run of the scheduler, the job will be logged under Scheduled Job Detail UI. Logs should be provided for each job run. If the job fails, logs will be updated accordingly with the Failed job status.

Automating AIML Metric Prediction Runs

Users will have the ability to generate this metric through a scheduled job. This will predict the number of orders during a user specified time. Users can configure the scheduled job in Scheduled Jobs UI and will create the data required to build the model.

Generating Prediction Runs

The above job type is available in the CRUD panes of the Scheduled Job UI.

- Selection of the above job type from the drop down will further expand the Job parameters:
 - ***Username***
 - ***Training Run Number***
 - ***Orders*** -
 - ***Wave number*** - (optional)
 - ***Order Type(s)***

The result of the model is available in Prediction Run UI when a user clicks on the corresponding details button in Prediction Run UI. By selecting the run number users will be able to see the results of the specific run. Output results will be stored separately per KPI/Metric and per prediction run and will be different for each metric.

The output data will also record if it was the result of a "default" template for that KPI/Metric.

Note: The detail drill-down from the Prediction Run UI will show the data in a data grid.

Step Two: Viewing the Training Data

The process to create one of these models is defined below:

- **Step One:** Create a Scheduled Job.
- **Step Two:** View the Training Data.
- **Step Three:** Define the model with parameters/filters and select it in the AIML Training Template.
- **Step Four:** Run the prediction and view results in the AIML Prediction Run UI.
- **Step Five:** Deploy, retrain models, and interpret results.

These screens allow you to view training data:

- **AIML Training Data** – Allows users to view the prepared data generated by a scheduled job.
- **AIML Training Template** – Where users define templates and optionally set a default.
- **AIML Models UI** – Displays all models built from training templates.
- **AIML Prediction Run UI** – Shows the results from prediction executions.
- **AIML Predictive Dashboard** – Displays the latest prediction results (currently available for Order Cycle Time only).

When the job runs, it processes the **Username** and selected **AIML Template**, and upon successful completion, generates a model. This job also triggers the **Train** action on the corresponding AIML Training Template defined in the job parameters, allowing users to generate models on a recurring basis.

The output (model) from each job run is captured in the **AIML Models UI**, where a new model is logged for every execution.

A new screen called **AIML Training Data** allows users to view the results of the scheduled job for:

- **Order Cycle Time**
- **Order Waiting Time**
- **Order Processing Time**
- **Intelligent Cycle Count**
- **Market Basket Analysis**

Users can select one of these metrics from a dropdown. Once selected, a data grid appears showing the relevant training data. This data represents the output of the "Generate Training Data for..." job.

Each metric will have its own unique data grid:

- Switching the dropdown will load a new grid with fields specific to the selected metric.
- Each sub-screen supports **search functionality**.
- Users can **delete** records but cannot **create or edit** them.
- **CSV download** is available.
- Users can also **filter by date range** using "From Date" and "To Date".

Two new sub-screens—**Intelligent Cycle Count** and **Order Cycle/Waiting/Processing Time**—are now available in the **AIML Prediction Run Details UI**. This is where the final prediction results are stored and displayed.

Step Three: Define Model with Parameters and Filters

A new UI screen, **AIML Training Template**, has been introduced. This screen allows users to configure the parameters and details required to generate an AI/ML model for a specific metric.

The following fields are **mandatory** when creating a training template:

- **Template Name**
- **AIML Metric**
- **Date Range**
- **AIML Algorithm**

Note:

- The template name **cannot contain a forward slash (/)**.
- Once saved, the **template name cannot be edited**.

In this screen, users can **train and generate AI/ML models** for future predictions. Similar to the **Wave Template UI**, the **AIML Training Template** screen allows users to configure specific parameters and settings required to generate a model for each supported metric.

The following KPIs are available for model creation:

- **Order Cycle Time** (regression problem)
- **Order Waiting Time** (regression problem)
- **Order Picking Time** (regression problem)
- **Intelligent Cycle Count**

Action Buttons Available in the Training Template UI

- **Train Template**

This button is enabled **only** for metrics that support model generation (e.g., **Order Cycle Time** and **Order Picking Time**).

Once the template is saved, users can invoke the **'Train Template'** action button. Upon clicking, the system will prompt the user to confirm whether they wish to proceed. If confirmed, the model training process begins, and the resulting model is captured and displayed in the **AI/ML Models UI**.

The **'Predict'** action button is available only for specific metrics that do **not** require a trained model. For these KPIs, selecting **Predict** will directly generate output based on the available data—no model creation is necessary.

Users can define the **date range** for which they want to generate AI/ML KPI data. However, if the selected range does not contain sufficient data, the system will return the following error:

"Insufficient Training Data to Generate Model"

This screen supports full **CRUD operations**. When editing an existing template, all fields except the **Template Name** are editable. Once the required fields are properly configured with valid values, the template can be saved successfully.

Step Four: Predict and Execute Capturing Results in AIML Prediction Run UI

Prediction Run UI

In Prediction Run UI, each prediction run is defined by the training run that has been used and other parameters that will vary by the metric. For example, Order Cycle Time, it is the order filter criteria.

Each time a metric invokes the 'Predict' button, the results will be stored here. Additionally, each scheduled job run for the job type '**Generate Prediction Run**' will result in an entry here.

The following fields are in the Prediction Run UI:

Field Name:	Description:
Prediction Run Nbr	System generated number on invoking Predict action or on running a 'Generate Prediction Run' job.
Training Run Nbr	Reference to the training Run from the AIML Training Models
Prediction Parameters	Reference to the parameters used to execute the model against new data to generate predictions
Default flag	Copied from training run; there can be multiple runs for same metric with this enabled. the latest one is the relevant one.
Status	created, training, completed, failed, cancelled

There is a drill down detail button available when a user selects any given record. This shows the output data for that run. It is specific to the metric of the run that the user selects. The output will be one of the following KPIs:

- - - Intelligent Cycle Count
 - Order Cycle time
 - Order Waiting Time
 - Order processing Time

Users are also able to use search/filters for Prediction Run Number, Training Run Number, Status and Default Flag. Each drill down screen also supports CSV download and appropriate filters as needed.

Order Cycle Time, Order Waiting Time, Order Processing Time, and Intelligent Cycle Count

When selecting a record in the AI/ML Models screen for any of the metrics, a user must click the 'Predict' button to carry out the prediction. A popup will be displayed to further finetune the models.

Note: If no field is populated, then a user will be shown an error.

- On clicking 'OK', the execution i.e. the prediction for the respective metrics will commence for the set of orders selected. (On 'Cancel', the action is cancelled.)
- On completing the execution, a record is generated in the Prediction Run (Inquiry) UI which records the parameters and details used to run that prediction.
- Also, the execution will generate the output i.e. the *Order Cycle Time* or *Order Waiting Time* or *Order Processing Time* depending on the metric of the record selected for prediction, viewed in a separate UI screen ex: 'AIML Order Cycle Wait Processing Time'

AIML Models Inquiry

Every training template execution, whether manual or scheduled will result in an entry here. For metrics that produce a model, the model generated will be stored within this screen.

The following fields are supported within this UI:

Field Name	Description
Facility	User's context facility
Training Run Number	System generated number on running a AIML training template
AIML Training Template	name of the training template ran for the specific run number

Field Name	Description
Training Start Date	Start date of the training data
Training End Date	End date of the training data
AIML Algorithm	Algorithm for which the model was generated
Training Filters	Copied from template
Algorithm Parameters	Copied from template
Default Flag	Same as from template
Status	Created, training, completed, tailed, cancelled
Message Text	- Shows the latest message from the training run - Will be a reference number or link - updated when training completes. some metrics won't have a model (like basket analysis) - This should be copied back to the template "Last Trained Model"
AIML Model Reference	Shows the performance of the model
Performance Metrics	Mean absolute error, r^2
Create Timestamp	Date and time of the model/record. Creation
Create User	User who generated the model/record (one who ran the AI/ML training template)
Metric	Metric being training, comes from template

Action Buttons in AIML Models Inquiry:

- - - **Logs**
 - **Predict**
 - **Prediction Run**

Each of the drill down screens supports CSV download and appropriate filters as needed.

Step Five: Interpreting Results

After successfully creating and viewing your models in AIML Models UI, it is important to note that your results may not yet be optimal. It could potentially take several iterations of the template to generate better results. Oftentimes it may be best to change the type of algorithm along with the types of parameters. Your results could also vary based on the type of industry you are in.

2 AIML Predictive Fulfillment Dashboard

AIML Predictive Fulfillment Dashboard

The **AI/ML Predictive Dashboard** allows you to make well informed decisions based on past order data to view future Order Cycle/Processing/Waiting Time predictions. This information comes from the most recent prediction from your model in AIML Prediction Run.

The predictive dashboard displays the average order cycle time in hours. We have included various KPIs that gives insight on how your orders have and will move out of the warehouse. By clicking on the order status, it will further filter and display the order type by status.

Your orders above and below your set target threshold will be shown in a datagrid that displays the following information:

- Order Number
- Status, Order Type
- Predicted Time
- Expected Shipping Date
- Required Shipping Date

The ratio of orders above and below your target threshold will also be shown in a pie chart above the data grid.

Steps to Enable

1. Enabled by default.
2. Accessible in the Redwood UI, via the **Username** drop-down.

Please note that only the Default Template will display data in the Predictive Dashboard. If no default template is configured, the predictive dashboard will not return any results.

You can access the AI/ML Predictive Fulfillment Dashboard by clicking the on the username.

Oracle Market Basket Analysis

We've introduced **Market Basket Analysis (MBA)** in Oracle Warehouse Management to help you gain deeper insights into your inventory and customer buying patterns. This AI-powered feature identifies **items frequently ordered together**, enabling smarter decisions when it comes to **inventory placement and putaway strategies**.

By understanding these patterns, you can optimize slotting, reduce picker travel time, and increase overall warehouse efficiency.

What It Does

Market Basket Analysis leverages **Oracle's advanced AI/ML algorithms** to analyze historical order data. It detects relationships between products and groups them into **"frequent item sets"** based on how often they are ordered together and the likelihood that this trend will continue.

The result is a clear, actionable report that helps you:

- Identify product pairings or groupings.
- Improve inventory layout by slotting complementary items together.
- Streamline picking operations based on actual order behavior.

There are **two key metrics** when running a Market Basket Analysis, support and confidence. The system ranks items based on support and then confidence, this means that if supports are equal between items, then they will be based on confidence after.

Support measures how frequently an itemset appears in the dataset. It's calculated as the proportion of transactions that contain a specific item or combination of items out of all transactions.

Confidence indicates the likelihood that item B is purchased when item A is purchased. It's calculated as the ratio of transactions containing both A and B to the number of transactions containing just A.

- Create a scheduled job to generate the training data for Market Basket Analysis in the **Scheduled Jobs** UI.

Note: You must enter a valid username, or else you'll receive an error.

- Navigate to the AI/ML Training Data to verify that the training data contains the correct data.
- Create a Market Basket Analysis Training Template in the **AI/ML Training Template** UI

Note: To run the Market Basket Analysis, you must select the "Apriori" algorithm.

- After creating your Market Basket template with all of the mandatory fields completed, click **Train Template**.
- Navigate to the **AI/ML Model** UI and select the most recent Market Basket Analysis entry. In the "Model Reference" field, click on the hyperlink.
- You can now view the results of your successful Market Basket Analysis.

Note: *association_max_rule_length* should only be changed if users want to see associations past an A & B association. For example, if *association_max_rule_length* = 4, users will see relationships between items A/B/C/D, however this may impact performance.

Inbound and Outbound Performance Metrics

You can access Inbound and Outbound Performance metrics via tabs in **Redwood Desktop Experience → Dashboard**.

From the **Inbound Performance Metric** tab you can view:

- LPNs Received in the last 7 days
- Open IB Shipments by Status

- Today's Appointments
- Last 7 days average dock to stock cycle time

From the **Outbound Performance Metric** tab you can view:

- Open Tasks
- Orders to be Shipped Today
- Outbound LPNs by Status
- Order Fill Rate

Accessing the Performance Metrics

1. To access the Inbound and Outbound Performance Metrics, click the user drop-down menu in WMS, and select the **Try the new Redwood Experience** option. A new tab will open up with the new redwood experience. The Home page gives you quick access to useful links like User Guides and Release Content.
2. Click the **Dashboard** tab to access the available Performance Metrics.

3 GenAI Features

Intelligent Wave Summary

In the **Wave Inquiry Summary** UI, you have a comprehensive and efficient way to review and analyze waves, bringing advanced AI capabilities to your warehouse management system.

AI Wave Summary Button

The **AI Wave Summary** button, in the wave inquiry header, provides quick access to valuable insights.

Click this button to open up a pop-up window, displaying an AI-generated summary of the wave execution, including key details and potential issues. The button is enabled only for specific wave statuses: 'Completed', 'Completed, Not Fully Allocated', 'Completed, Nothing Allocated', and 'Failed', ensuring relevant information is readily available.

The screenshot shows the 'Wave Inquiry' header with a search bar and 'Create Timestamp' and 'Mod Timestamp' buttons. Below the header is a table with columns: Facility, Wave Template, and Allocation Method. A 'Refresh' button and a menu icon are also visible. The menu is open, showing options: Details, Wave Logs, **AI Wave Summary** (highlighted with a red box), and Clear AI Wave Summary. The table contains two rows of data:

Facility	Wave Template	Allocation Method
QATST01	WJ WAVE-CUBE	LIFO
QATST01	WJ WAVE-CUBE	LIFO

Customized Wave Summaries

Different pop-up summaries are generated based on wave types and statuses, providing tailored insights. Summaries for 'Completed' and 'Failed' statuses offer a clear understanding of the wave execution for:

- Picking Waves and Wave-based Replenishment
- Replenishment Waves
- Lean Time Replenishment Waves

Note:

- AI Summary features are not supported for all languages that the WMS supports. For languages not supported by AI, the summary will default to English.
- AI Summary features are enabled by default in North America and Europe geographies. Customers elsewhere will have to raise an SR if they want this feature as this requires authorization of data communication to one of the two geographies

Wave Inquiry and Wave Inquiry Detail UI

From the **Wave Inquiry** and **Wave Inquiry Detail** UIs:

- You can view all records related to a wave run, ensuring a comprehensive overview.
- Action buttons, such as 'Substitute Item', are readily accessible, allowing users to take immediate action.
- Advanced search capabilities enable users to filter and find specific wave inquiries efficiently.

Permission-Based Access

- Group Permission for 'Wave Inquiry / AI Wave Summary' determines user access.
- By default, this permission is disabled, allowing administrators to grant access to specific user groups.

Wave Inquiry UI - Replenishment Wave

The **Wave Inquiry** UI for **Replenishment Wave** includes the following new fields:

- **Total Wave Time:** This field will display the overall duration of the wave, providing a clear understanding of the time taken for the entire process.
- **Task Creation Time:** Warehouse managers can now pinpoint the exact moment when tasks are generated, allowing for better task management and resource allocation.
- **Allocation Time:** Knowing the allocation timing is crucial for optimizing inventory management. This field will reveal when the allocation process occurs, helping you identify potential bottlenecks.
- **Number of Tasks Created:** A clear indication of the workload generated by each wave, enabling better workforce planning.
- **Total Quantity Allocated:** Provides insight into the volume of inventory allocated, aiding in stock management and ensuring efficient resource utilization.

Search by Replenishment Trigger Modes

You can search by Replenishment trigger modes via the following search fields in the Wave Inquiry UI:

- **Minimum Capacity:** Quickly find waves triggered by minimum capacity thresholds, ensuring you can replenish stock promptly.
- **Reactive Replenishment:** Identify waves generated through reactive replenishment strategies, allowing for immediate analysis and response.
- **Percentage of Max:** Locate waves based on percentage-based replenishment triggers, enabling precise inventory management.
- **Order Based Replenishment:** Focus on waves created for order fulfillment, helping you prioritize customer orders efficiently.
- **Movement Request Based Replenishment:** Search for waves related to movement requests, providing a comprehensive view of inventory adjustments.

This search feature is also available as an optional column in the **Wave Inquiry** UI, ensuring a clutter-free interface. You can enable and customize your search criteria, adding this field when needed for tailored analysis.

Clear AI Wave Summary Button

From the **Wave Inquiry UI**, the 'Clear AI Summary' button allows you to manually clear the wave summary information stored in the database, offering the ability to refresh and update wave details whenever needed.

The **Clear AI Wave Summary** button is useful for:

- **Wave Summary Generation:** After a wave run, the system automatically generates a summary, including order selection, sequencing, and allocation strategies.
- **Wave Failure and Resolution:** If a wave fails, the manager can troubleshoot and resume the wave.
- **Clearing Summary Data:** Upon resuming a failed wave, the 'Clear AI Summary' button removes the previous summary data from the database, reflecting the wave's new 'Completed' status.
- **Regenerating Summary:** Users can click the 'AI Wave Summary' button to create an updated summary, providing a fresh analysis of the wave's performance.

Steps to Enable

To enable AI Wave Summary in the Wave Inquiry UI:

1. Enable the Group Permission '**Wave Inquiry / AI Wave Summary.**'

To access **AI Wave Summary**:

1. From Redwood Desktop, click Ask Oracle.
2. Select the record(s) and click the **More Actions** ellipses button (...).
3. Click **AI Wave Summary**.
4. An AI Summary pop-up will appear with the AI Wave Summary:

AI Wave Summary

Wave Execution Summary

Run Nbr: WVQATSTPC096240

Wave Type: Picking Wave

Wave Template: WJ WAVE-CUBE

Create Timestamp: 09/17/2025 04:00:36 PM

Status: Completed

Nbr Ord Headers Selected: 1

Nbr Ord Details Selected: 1

Nbr Ord Headers Allocated: 1

Nbr Ord Details Allocated: 1

Nbr Ord Headers Allocated: 1

Nbr Ord Details Allocated: 1

Allocation Time: 0.01

Nbr Alloc Mode Seq Executed: 1

Nbr SKUs Selected: 1

Nbr SKUs Allocated: 1

Nbr Tasks Created: 0

Load Creation Time: 0

Nbr of OB LPNs: 1

Cubing Time: 0.01

Key Events

- The wave was started, initiating the allocation process.
- Wave allocation sequence began, indicating the start of item allocation.
- Allocation focused on units/pack for a specific category.
- A small number of items and locations were involved in the allocation, with a limited inventory.
- Units were successfully allocated, but no LPNs were generated.
- Allocation run was completed, and cubing started.

4 AI Agents

AI Agent Studio Support in WMS

We are introducing integrated support for Oracle Fusions AI Agent platform as the agentic AI platform for WMS. We've released a set of capabilities both in AI Studio as well as in WMS to enable much closer and seamless integration.

Note: AI Studio is delivered as part of the core Fusion platform and therefore customers also need to subscribe to a core Fusion product in addition to WMS, to be able to access agentic AI features in WMS.

- AI Studio now supports the business object tool for non-Fusion Oracle products using the “Other data source” feature. This allows the user context to be passed seamlessly between AI Studio and WMS and for standard WMS business objects to be exposed, rather than having to hardcode API calls using the generic REST tool.
- AI Agent chat can be invoked from the WMS Redwood desktop and mobile UI's
- WMS Agents can be configured in WMS with control for specific WMS groups
- Agents can now be embedded in WMS Redwood screens, starting with the Wave Research Advisor in Wave inquiry
- Pre-seeded WMS Agents shipped with Fusion AI Studio.

26A includes three WMS agents that have shipped with AI Agent studio:

- Wave Research Advisor
- Inventory Expiry Assistant
- Task Management Assistant

These pre-seeded agents are designed for limited use cases. Starting the chat with ‘Hi’ will get the agent to summarize its capabilities. The capabilities of these agents may be expanded by Oracle in later releases. In addition, customers can also extend these agents using Agent studio. They can also create their own custom agents (may require purchasing a custom agent SKU) and integrate them with WMS.

NoteFor information on using AI Agent Studio, see [How do I use AI Agent Studio](#)

AI Agent Studio and Fusion-Specific Configuration

WMS 26A AI Agents

WMS is shipping 3 AI Agents for Fusion AI Studio. These are pre-seeded in Fusion and automatically available as of the Feb 6 CWB Fusion patch bundle or later. All are available as runnable agents and two are available as templates. Runnable agents are not editable but are easier to setup and start using out of the box.

Agent	Runnable	Template
Wave Research Advisor	Yes	No

Task Management Assistant	Yes	Yes
Inventory Expiry Agent	Yes	Yes

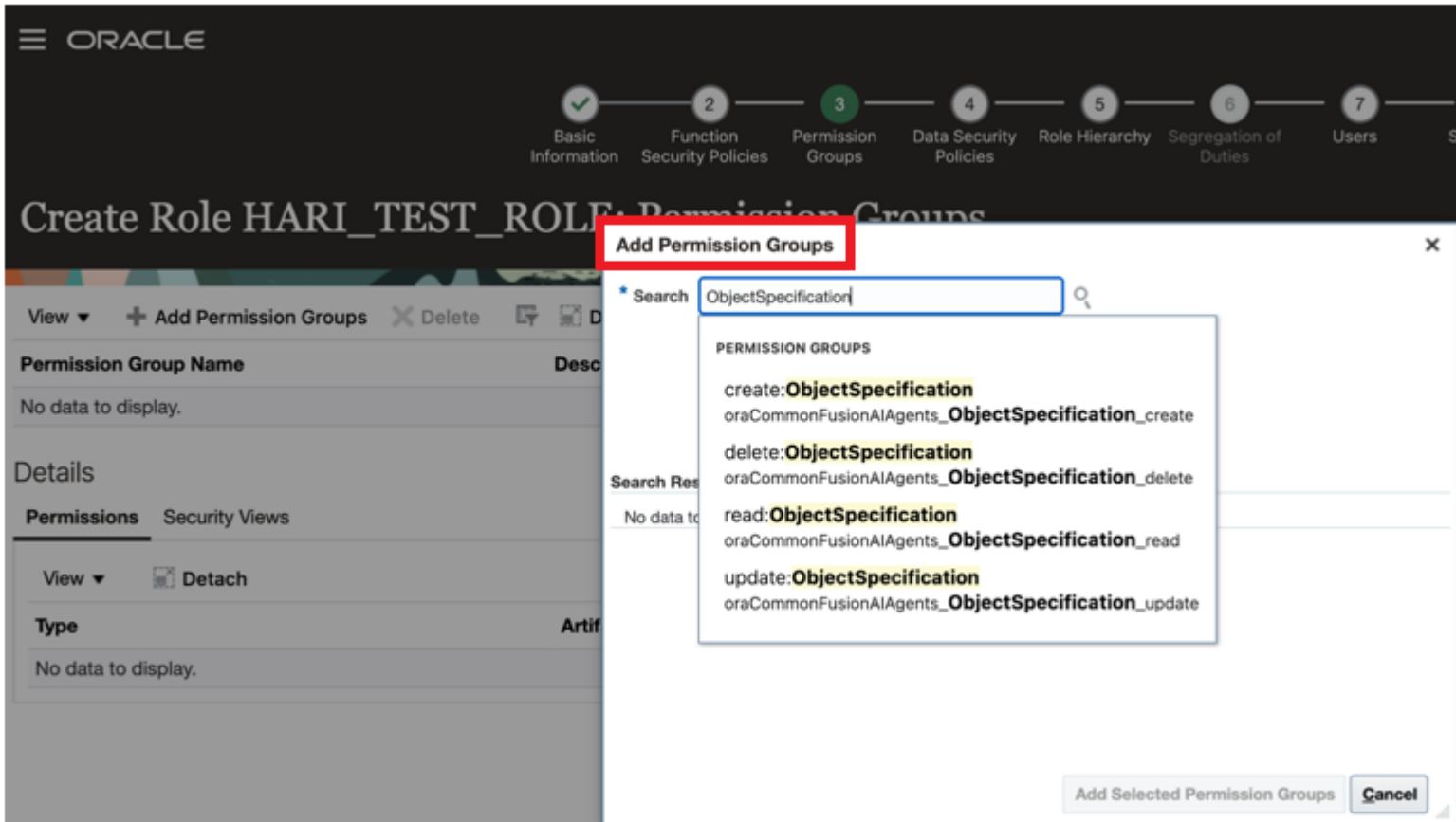
Other Data Source Setup - Credentials Tab

Fusion AI Studio interacts with WMS via a new AI Studio feature called “Other data source”, which is available in the “Credentials” tab as shown in the screenshot. Fusion is seeded with a blank entry named WMSENV. IDCS in both Fusion and WMS must be configured . And this WMSENV entry must be edited using the instructions in AI Studio-WMS-IDCS_configuration.pdf.

Configuring Fusion Permissions

In Fusion 26A, please follow the steps below to add this role and permission. This is in addition to the roles/permissions required for Fusion AI Studio. In this example, the additional role is being created and added to the SCMOOPERATIONS user which already has the roles needed for AI Studio agents. In 26B this additional setup will not be required.

Add Permission Groups



Permission Group Added

View ▾ + Add Permission Groups × Delete Detach		
Permission Group Name	Description	Artifacts
create:ObjectSpecification	create:ObjectSpecification	ObjectSpecification
read:ObjectSpecification	read:ObjectSpecification	ObjectSpecification
update:ObjectSpecification	update:ObjectSpecification	ObjectSpecification

Role Added to User

Create Role HARI_TEST_ROLE: Users

Roles and Privileges

Roles and Permission Groups

+ Add User

× Delete

Detach

User Login

SCMOPERATIONS

Agents and WMS Business Objects

1. In the “Business objects” tab, there should be a seeded business object called WMS26A. This business object is linked to the WMSENV data source
2. In the first “AI Agent Studio” tab, searching for WMS should show the two templates listed in the table above. Agent teams can be created from templates and users can edit them.
3. In the “Agent teams” tab, searching should list all three runnable agents listed in the table above. There will be other WMS agent teams, but these are not meant to be used directly. They are used by the Wave Research Advisor

Configuring WMS

1. View the runnable agents and note down the agent team codes.
2. In the WMS UI, find the AI Agents screen (only available in Redwood) and set these up.
3. From the Groups screen, open the agents drill-down screen and add the agents that you wish to allow a group to access.
4. Users will see the agents assigned to their groups in the AI chat window in WMS.

For additional details, see [AI Agent Configuration in WMS](#)

AI Agent Configuration in WMS

The following are configuration steps for enabling AI Agents in WMS:

Steps to Enable and Configure

To connect with the Fusion AI Agent Studio app in WMS, do the following:

1. On the SAAS App Configuration UI, click **Create**.

Note: You must access the SAAS App Configuration UI from the Redwood UI in WMS.

2. From the **App Name** drop-down list, select **AI-AGENT-STUDIO**.
3. Enter **ALL** of the following mandatory fields detailed in the table below.

Note: Prerequisite – make sure FA IDCS Oauth2 with grant-type JWT confidential app is configured. For more information about **Fusion Identity Cloud Service (IDCS)** setup, see sections 3 and 4 of the **AI Studio Integration for WMS** document.

4. Click **Save**.

Refer to the following table for each parameter's meaning and the format required.

Field Name	Description	Details
Base URL	FA AI Studio Base URL	Fusion base URL where AI Studio is hosted. https://xxx.fa.us6.oraclecloud.com
App parm3	IDCS Token URL	https://idcs-xxx.identity.oraclecloud.com:443/oauth2/v1/token
App parm4	Client id of the WMS confidential app in FA IDCS	1. Enter the Client ID. 2. In your OCI Instance follow the path "Identity and Security" --> Domains --> 3. Select the domain that hosts the OCI instance. 4. From the Domains screen, click the "Oracle Cloud Services" tab. 5. Search or browse to your OCI application, then click to open.
App parm5	Fusion IDCS confidential application scope.	Example: urn:opc:resource:fusion:{tenant-pod}:fusion-ai/. It can vary if your Fusion base URL is a custom URL. Please confirm as described in the Fusion IDCS confidential application setup section. (See section 3.1 "Create a Confidential Application" in the AI Studio Integration for WMS document.
App parm8	Private Key for the WMS confidential app in FA IDCS	A valid RSA private key. This is used to sign the access token request. It should be in the format of PKCS#1. Note: App parm8 is encrypted in the database and not displayed in the grid view (but it is editable).

Field Name	Description	Details
App parm9	Public x509 cert	Corresponds to private key entered in parm8 field.

Note: SaaS App configurations can be created for Facility * (one AI Agent configuration per company).

To add the AI Agent to the list of available WMS AI Agents, do the following:

- 1. Once you’ve added the **AI Agent Configuration UI** from Modules, you’ll be able to access it from Redwood Desktop -> **Configuration Tab** or -> searching from **Ask Oracle**.
- 2. From the **AI Agent Configuration UI**, click **Create**.
- 3. Enter your AI Agent **Name** and **Agent Team Code**.
- 4. Click **Save**.

QATST01 / QATSTPC

AI Agent Configuration

AI Agent Configuration

Search

Name

Agent Team Code

Refresh

Copy

Edit

Delete

Company	Name	Agent Team Code
☐ QATSTPC	WMS Inventory Expiration Advisory Workflow	WMS Inventory Expiration Advisory Workflow
☐ QATSTPC	WMS Task Management Analysis Agent Workflow	WMS Task Management Analysis Agent Workflow
☒ QATSTPC	Wave Research Advisor	Wave Research Advisor

Note: Agent Team Code must be unique to the company and match the agent team code in Fusion AI Agent Studio.

Required Configuration Via Groups

Additionally, you need to assign access to agents by assigning them to Groups. To add an AI Agent to a group, do the following:

1. From the **Redwood Groups** UI, select a group for which you want to assign the AI Agent.
2. Click **More Actions** (elipses button to the right of the Refresh button) (...), then **AI Agents** action button.
3. On the AI Agents screen, click **Create**.
4. To assign an AI Agent to the selected group, select an AI Agent from the drop-down list.
5. Click **Save**.

To access the AI Agent in WMS after it is configured:

1. From the User drop-down, select **Try the New Redwood Experience**.
2. Hover your cursor over the **Ask Oracle button**, and an elipses slide-out button will appear. Click the elipses button.
3. From the list in WMS, select the AI agent and enter a prompt to get started.



Note: If only one AI Agent is configured, the system launches the configured AI Agent. From Redwood Mobile, you can click the Ask button to launch the AI Agent pane.

Note: For information on using AI Agent Studio, see [How do I use AI Agent Studio?](#)

AI Agent: Wave Research Advisor

The **Wave Research Advisor AI Agent** is an intelligent, automated solution designed to simplify and enhance your review of warehouse wave runs. By analyzing data from wave logs, metrics, and related entities, the agent delivers immediate, consolidated insights without requiring users to navigate multiple screens or pages.

Using a workflow format, the Wave Research Advisor Agent empowers warehouse supervisors and outbound managers to more efficiently monitor wave performance, quickly identify and address issues, and focus on optimizing operations.

Here's a snapshot of the **Wave Research Advisor Agent's** capabilities:



Hello! I'm your WMS AI Wave Research Agent, here to help you with all your wave-related inquiries, including orders, inventory, cubing, and tasking. Here's how I can assist you:

What I can do:

- Summarize wave details
- Check if an order or LPN was allocated by a wave
- List outbound LPNs created by a wave
- Identify tasks created by a wave
- Answer location-based allocation questions

Sample commands you can try:

Wave Summary:
Give me a summary of wave {run_nbr}

Order Allocation:
For wave {run_nbr}, was order {order_nbr} allocated?

LPN Allocation:
Was LPN {LPN} allocated by wave {run_nbr}?

Location-based Allocation:
Was anything allocated in area A, aisle 7?

Cubing Questions:
How many outbound LPNs were created by wave {run_nbr}?

Tasking Questions:
What tasks were created by wave {run_nbr}?

Just type your question or use one of the sample commands above, and I'll be glad to assist!

Ask Oracle

To implement this AI Agent, we have introduced the following in Redwood:

- A new app configuration “AI-AGENT-STUDIO” on the **SAAS App Configuration (SaasAppConfigView)** UI to connect with Fusion AI Agent Studio.
- A new UI screen **AI Agent Configuration (AiAgentConfigView)** to add AI Agents.
- A new action button **AI Agents** on the **Groups** UI to assign AI Agents to a group.
- New **Agent Code Parameter** parameter which allows you to configure launching agents from the Wave Inquiry screen.
- A new ellipses button “Ask” on the Redwood UI and Redwood Mobile to launch AI Agents.

Before you use the AI Agent, you have to complete the key configurations mentioned below in **Steps to Enable and Configure**. Once you complete all configurations, you can launch the AI Agent and use prompts to perform specific tasks.

KEY SCREEN PARAMETERS – WAVE INQUIRY UI

Wave Research Advisor Parameter

You can also launch the agent directly from the **Wave Inquiry** UI in Redwood which will allow you to easily pull in wave information when you are working with the agent.

1. From **Screens -> Screen Parameters**, assign the agent code **Parameter Value (Wave Research Advisor)** for the **ai-agent-code** parameter. Then assign the **Parameter Value Wave Summary {run_nbr}** to the **ai-agent-selected-row-question** parameter.

ORACLE wms

Ship Via

User

Users (8) ×

Order Header

Cubing Rules

Groups (6)

Screens (9)

Modules

← Screens (9) ▶ Screen Parameters

↺ 🔍

Screen	Module Parameter	Parameter
Wave Inquiry (12)	ai-agent-selected-row-question	Text
Wave Inquiry (12)	ai-agent-code	Text

Screen Parameter Value

1. From the **Redwood Groups** UI, select a group for which you want to assign the AI Agent.
2. Click **More Actions** (elipses button to the right of the Refresh button) (...), then **AI Agents** action button.
3. On the **AI Agents** screen, click **Create**.
4. To assign an AI Agent to the selected group, select an AI Agent from the drop-down list.
5. Click **Save**.
6. You'll then be able to launch the Wave Research Advisor agent via "Ask Oracle." From the **Wave Inquiry** UI.

Wave Inquiry

Wave Inquiry

○

Search

Create Timestamp

Mod Timestamp

Run Nbr

Allo

↺ Refresh

...

	Total Wave Time ↕	Task Creation Time ↕	Allocation Time ↕	Nbr Tasks Cre
☐	0.01	0	0	
☐	0.01	0	0	

AI Agent Selected Row Question Parameter

A new screen parameter, **ai-agent-selected-row-question**, is available on the **Wave Inquiry UI**. This parameter allows users (or admins) to define a default question template containing a placeholder ({{{ {run_nbr}}}} or similar) which will then populate the AI Agent Chat window.

If the user selects exactly one row and the ai-agent-selected-row-question parameter is populated:

- The system substitutes the selected row's field inside the brackets. (e.g., {run_nbr}
 - The resulting question is automatically sent as the initial message to the AI agent when the chat opens.
 - If the placeholder does not match any available field, the literal template string is sent as the initial chat message (e.g., "Wave summary {wave_nbr}").

For example:

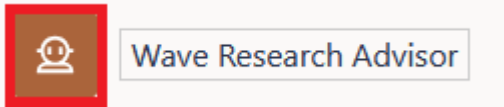
The following is an example of two selected rows with the Wave Summary run number and Task creation run number:

```
Summary {run_nbr}
Task created {run_nbr}
```

SELECT AI AGENT

Once you have your agent configured, you have the option to switch to a different agent from inside the agent chat window.

To switch to a different agent, click the **Select AI Agent** button.



Here are some sample prompts/questions that could be useful when interacting with the **Wave Research Advisor** agent:

Category	Sample Prompt/Question	Example Answers
Orders	For Wave W20 was order O123 allocated and how much quantity?	- Order O123 with order type DS was partially allocated. 27 units out of 50 were allocated including 3 LPNs. 2 tasks were created for this order - Order O123 was selected but not allocated by the wave. It was of order type DF which has the partial allocation flag disabled. Not enough inventory was available to allocate it fully - Order O123 was not selected for this wave, however it is in status allocated and linked to a different wave W25
Orders	Has Order O123 been wave released?	- Order: 0123 - Order Status: Allocated - Allocated by Wave provided: no - Allocated by Wave WXYZ - Total order quantity versus Total Allocated quantity 6 vs 6 - Order Ship dates null
Orders	Is order OrdCJ1812202502 allocated?	- Order: OrdCJ1812202502 - Order Status: Allocated - Allocated by Wave provided: No - Allocated by Wave: WVQATSTPC096734 - Total Order quantity vs Total allocated quantity: 600 EA vs 600 EA - Order Ship date: 12/17/2025
Location based queries	Was anything allocated in area A, aisle 7	- No inventory was allocated from that active location 11 units were allocated from area A, aisle 7 from 3 different locations

Steps to Enable and Configure

See *AI Agent Configuration in WMS* for the configuration steps for enabling the Wave Research Advisor AI Agent in WMS.

AI Agent-Inventory Expiry Assistant

The Inventory Expiry Assistant agent is a powerful AI-driven assistant designed for warehouse supervisors and inventory managers to proactively monitor and manage inventory shelf life. This intelligent tool helps identify, isolate, and secure lot-controlled goods that have expired or are nearing expiration, ensuring only safe and compliant inventory is processed for customer orders.

This agent provides:

- **Real-Time Expiration Monitoring** - Instantly identify expired inventory or items at risk of expiring across Reserve, Active, and Outbound areas using a workflow agents.
- **Interactive Chat** - Users can query the AI Agent with specific requests. For example, to list expired inventory by location, date range, item, or storage area - and get instant, relevant responses.
- **Container Locking** - Through the chat interface, users can apply locks to individual or multiple containers to prevent allocation. The Agent ensures that only containers with correctly configured locks are secured and notifies users if a lock cannot prevent allocation.

With the Inventory Expiry Assistant, you can significantly reduce the risks of shipping expired products, ensure regulatory compliance, and optimize warehouse operations with AI-powered efficiency.

Before you use the AI Agent, you have to complete the key configurations in the **Steps to Enable and Configure** section. Once you complete all configurations, you can launch the AI Agent and use prompts to perform specific tasks.

Here are some sample prompts/questions that could be useful when interacting with the Inventory Expiration Assistant:

Sample Prompts/Questions	Sample Response
Find me LPNs in area A that are expired	• There are 25 LPNs in area A that are expired. Here are the LPN numbers: xxxx • There are no expired LPNS in area A. However 3 LPNs are about to expire in 2 days: xxx
List the Active locations with Expired Inventory (OR)	There are 458 LPNs with the expired inventory. Here are the first few by earliest expiry listed below: Container Number: CSPMM141890011 Item/SKU: EXP011 Expiry Date: 2019-10-31 Location: R0152A1
List the Inbound LPNs with inventory that will expire in the next 3 days	There are no LPNs that are expiring in the next 3 days.
List the Outbound LPNs with expired inventory from last month	There are 395 LPNs with the expired inventory. Here are the first few by earliest expiry listed below: Container Number: CSPMM00000704 Item/SKU: TSTITM58 Expiry Date: 2017-06-02 Location: Container Number: OBOR440404 Item/SKU: TSTITMBE1 Expiry Date: 2018-01-17 Location:

Steps to Enable and Configure

See *AI Agent Configuration in WMS* for the configuration steps for enabling the Inventory Expiry Assistant AI Agent in WMS.

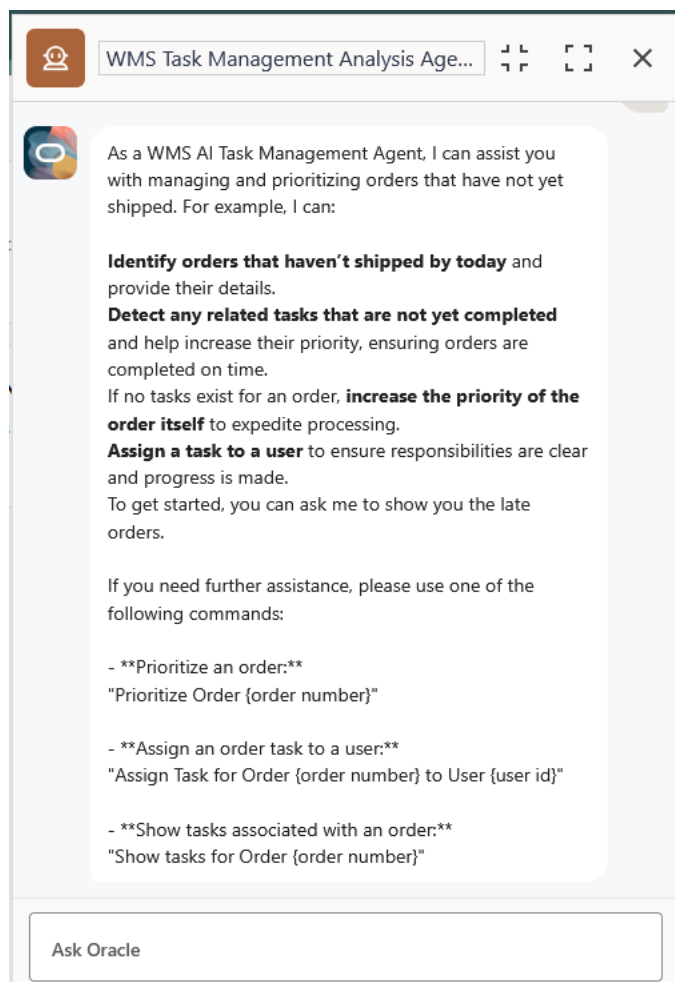
AI Agent: Task Management Assistant

The Task Management Assistant in Oracle WMS is an AI-powered agent that helps you manage tasks efficiently. It enables you to identify and prioritize orders that are at risk, manage associated tasks, and assign tasks to users. With the Task Management Assistant, you can perform various actions such as:

- Search for “Orders at Risk” that have not shipped by the required dates (example: orders that haven’t shipped by today).
- Get list of such orders, including order number, required ship date, customer name and create timestamp.
- View tasks associated with a specific order number.
- Prioritize order tasks in the agent workflow (example: reprioritize tasks for urgent orders).
- If an order is in created status and no tasks are associated with it, reprioritize the order itself.
- Assign a task to a user.

Note: This AI Agent is available only on the Redwood UI and Redwood Mobile.

Here is a snapshot of the **Task Management Assistant’s** capabilities:



Before you use the AI Agent, you have to complete the key configurations mentioned below in **Steps to Enable and Configure**. Once you complete all configurations, you can launch the AI Agent and use prompts to perform specific tasks.

Note: You can configure these configurations only on the Redwood UI.

Here are some sample prompts/questions that could be useful when interacting with the Task Management Assistant:

Sample Prompt/Question	Task Management Assistant's Behavior
Can you share all orders that are not shipping today?	Displays the first 20 late orders.
What orders are missing shipping date?	Displays the first 20 late orders.
Show tasks for Order {order number}	Displays tasks for the order and are not prioritized (priority=0) yet.If no tasks are associated with the order, returns a generic message.
Assign Task for Order {order number} to User {WMS user id}	Assigns all tasks for the order to the specified user.

Prioritize Order {order number}	Reprioritizes the “ready” status tasks for the order.If the order is in “created” status, reprioritizes the order priority itself.
Out-of-scope or non-functional questions	Returns a generic message.

Note: The AI Agent displays a maximum of 20 records.

Steps to Enable and Configure

The following are the key steps involved in configuring and launching the Task Management Assistant:

1. Configure “AI-AGENT-STUDIO” on the **SaaS App Configuration (SaasAppConfigView)** UI with required fields.
2. Create AI Agents (Task Management Assistant) with “Agent Team Code” on the **AI Agent Configuration (AiAgentConfigView)** UI.

Note: Agent Team Code must be unique to the company and match the team code in Fusion AI Agent Studio.

3. Assign the AI Agent to the desired group.
4. On the Redwood UI, hover over Ask Oracle Button and click “Ask” button to launch AI Agent pane. From the list of AI Agents, select your Task Management Assistant.

On the Redwood Mobile, you can click Ask button to launch AI Agent pane.

Note: If only one AI Agent is configured, the system launches the configured AI Agent.

5. Provide your prompts to perform specific requests as mentioned above.

See *AI Agent Configuration in WMS* for detailed configuration steps for enabling the Task Management Assistant AI Agent in WMS.

5 AI Agentic Apps

Logistics Execution Command Center

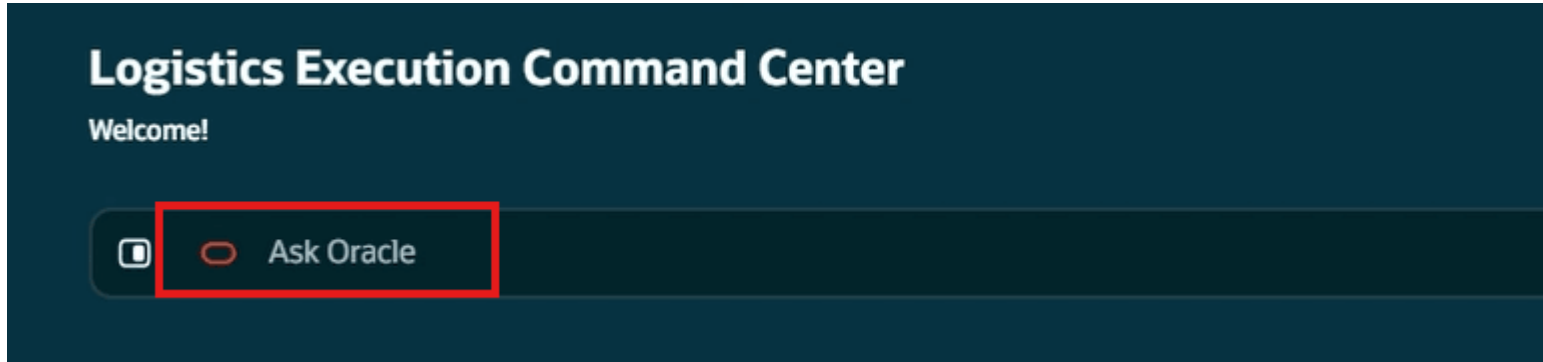
The Logistics Execution Command Center Agentic app streamlines workflows by surfacing urgent issues in a single **Command Summary**, consolidating Warehouse Management (WMS) and Oracle Transportation Management (OTM) data. All user actions are logged and synchronized automatically, optimizing process flow.

The app can allow a planner to use both softwares in a single app, providing useful metrics using interactive UI/UX, highlighting exceptions with the supply chain operations and providing actionable insights into both products. This app identifies the top orders from OTM and WMS that pose the highest short-term risk to supply chain operations if no action is taken.

Traditionally, identifying these exceptions requires multiple custom dashboards, complex queries, and manual configuration. An agentic app can ingest OTM and WMS data directly, automatically detect at-risk orders, and generate actionable insights far faster than conventional methods.

Using "Ask Oracle" in the Logistics Execution Command Center

The Ask Oracle experience introduces the Logistics Execution Command Center, the agentic core of the app. Users can ask questions in natural language to get immediate, cross-system intelligence—reducing time spent navigating menus and accelerating decision-making.



Summary of Orders at Risk and Priority Actions

On launch, logistics managers see a summary of orders and priority actions that reflects the “pulse” of execution across both transportation and warehousing. The most critical action items, including Orders at Risk are surfaced immediately so teams can focus attention where it matters most.

Logistics Execution Command Center

Welcome!

Ask Oracle

Orders at Risk of Delayed Shipment

Seven orders scheduled to ship today are at risk due to incomplete picking and planning. Most orders are at risk due to incomplete picking and planning, creating a risk of missed on-time shipments without prompt action.

Warehouse Metrics

Pending Orders by Task Type

Order Nbr	Task Type	Priority
WMSOTM_26B_PS_14	LPNUNITS	0
WMSOTM_26B_SY_02	LPNUNITS	0
WMSOTM_26B_SY_03	LPNUNITS	20
WMSOTM_26B_SY_06	LPNUNITS	20
WMSOTM_26B_SY_07	LPNUNITS	0
WMSOTM_26B_SY_08	LPNUNITS	0

Planned Orders

Order Releases: Route and Pickups

Showing top 10 of 12 distinct order releases (descending by risk severity, then lowest packing percentage, then lowest packing percentage).

WMS Order	Order Release
WMSOTM_26B_SY_09	WMSOTM_26B_SY_09
WMSOTM_26B_SY_06	WMSOTM_26B_SY_06

OTM-WMS Workflow Context:

- After order releases are planned into shipments in OTM, the data is sent to the warehouse management system to create corresponding WMS orders aligned to shipment requirements. This is done using the Four Part Key.
- Fulfillment depends on inventory availability, warehouse capacity, and operational efficiency.
- Once WMS orders are created, an allocation wave assigns resources. Orders may be fully or partially allocated.

- Fully allocated orders proceed to packing. Partially allocated orders may miss the planned fulfillment window, which ends at the OTM shipment start time.
- If an order is only partially allocated and cannot be fulfilled in time, it should be replanned to a new shipment scheduled to start the next day.
- For allocated orders, tasks corresponding to these orders are given higher priorities to complete faster.

Order Line Matching Between WMS and OTM via Four Part Key

We're using the WMS-OTM integration to improve accuracy and reliability when identifying and exchanging order data between WMS and OTM.

Four-Part Integration Keys

The WMS-OTM Integration setup requires the following fields to be populated:

- ERP_SOURCE_HDR_REF = Refnum with Qual ORDER NUMBER
- ERP_SOURCE_SHIPMENT_REF = Refnums with Quals ORDER LINE NUMBER, FULFILLMENT LINE NUMBER.EXT, FULFILLMENT LINE NUMBER
- ERP_SOURCE_LINE_REF = Refnum with Qual ORDER LINE NUMBER
- ERP_SOURCE_SYSTEM_REF = FlexField Number

These updates provide more improved fulfillment execution, ensuring the agent can reliably identify the orders that need to ship today, and locate the corresponding OTM orders for those shipments.

Note: For orders to be considered in the agentic app, they must be within 3 days.

Unified View of Planned Orders (WMS + OTM)

A single, consolidated **Planned Orders** pane shows WMS orders alongside their related OTM transportation orders, ranked by risk. This unified context makes it easy to understand whether risk is coming from warehouse execution, transportation planning (or both) without toggling screens.

Once you expand the **Planned Orders** pane, you can view more details, including the pickup window.

Planned Orders

Order Releases: Route and Pickup Windows

Showing top 10 of 12 distinct order releases (deduplicated). Pickup windows are localized to each source location's timezone;

WMS Order	Order Release	Source → Destination	Pickup Window	
WMSOTM_26B_SY_09	WMSOTM_26B_SY_09	SEATTLE MANUFACTURING → SUMMIT AUTO	Fri, Feb 27 2026, 7:58 PM MST – -	P R
WMSOTM_26B_SY_06	WMSOTM_26B_SY_06	SEATTLE MANUFACTURING → SUMMIT AUTO	Tue, Mar 3 2026, 5:44 PM MST – -	P R
WMSOTM_26B_SY_02	WMSOTM_26B_SY_02	SEATTLE MANUFACTURING → SUMMIT AUTO	Fri, Feb 20 2026, 4:03 PM MST – -	P R
WMSOTM_26B_PS_14	WMSOTM_26B_PS_14	SEATTLE MANUFACTURING → SUMMIT AUTO	Fri, Feb 20 2026, 10:38 AM MST – -	P R
WMSOTM_26B_SY_08	WMSOTM_26B_SY_08	SEATTLE MANUFACTURING → SUMMIT AUTO	Thu, Feb 19 2026, 11:47 AM MST – -	P R

Run Wave for At Risk Orders

From Planned Orders view, you will have an option from the **Priority Actions** pane to directly Run a Wave for At Risk Orders.

Priority Actions (1)

Critical

OTM-WMS Logistics Orders Agent

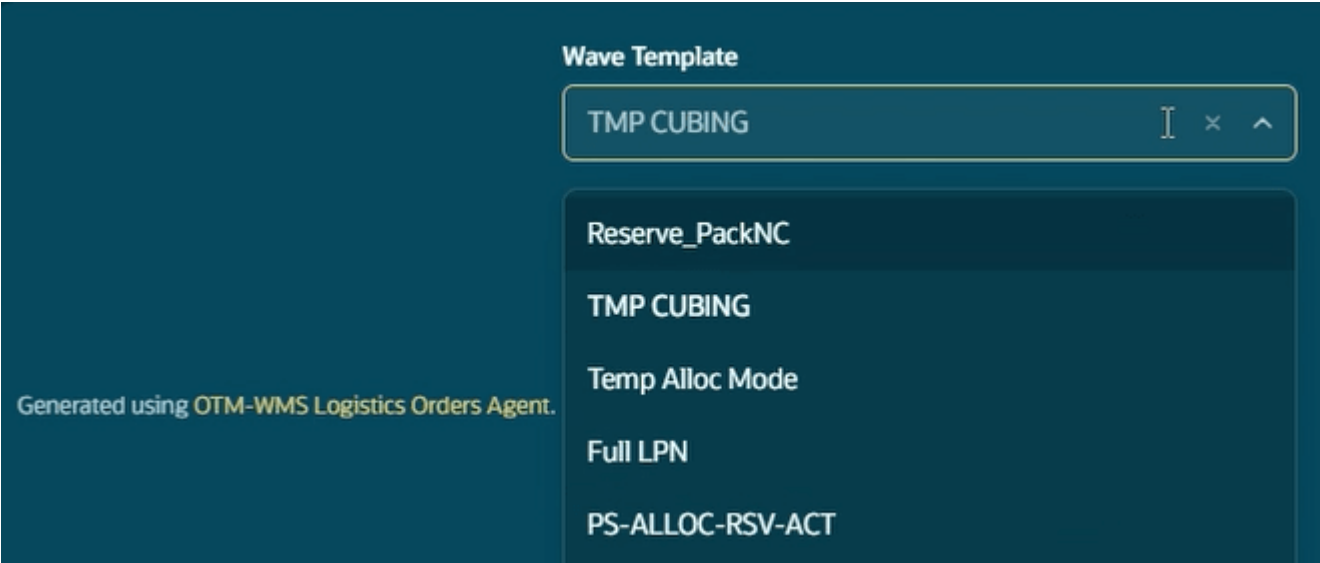
×

At-Risk WMS Orders

Orders WMSOTM_26B_SY_09 are in Created/Partly Allocated status—allocate now.

Run Wave for At Risk Orders

After clicking the Run Wave button, you'll see a **Wave Template** field where you can press up and down arrow keys, or type and search for a template and run the wave.



Query for Order Line Status

From the Planned Orders pane, you also have the option click on each hyperlinked WMS order to query inside order lines and search in active and reserve locations for order summary information. The query will search for orders and their associated items and provide a high-level summary of the available inventory in the warehouse. This will give you an idea of what inventory will be fulfilled.

Order Releases: Route and Pickup Windows		
Showing 10 of 12 de-duplicated order releases. Pickup windows are localized to each source		
WMS Order	Order Release	Source → Destination
WMSOTM 26B SY 09	WMSOTM_26B_SY_09	SEATTLE MANUFACTURING → SUMMIT AUTO
WMSOTM 26B SY 06	WMSOTM_26B_SY_06	SEATTLE MANUFACTURING → SUMMIT AUTO

Communications Pane

The **Communications** pane will prompt you with suggested email communications that you can populate and send. For example, if there are any orders in critical status, you might want to send an email communicating this. Eliminating the need to hunt for data, managers can generate and send communications with a click.

Critical Shipment & Orders Alert

To

Enter email address...

CC

Enter cc list...

Subject

Critical: Shipment & Orders 2026-02-25

Paragraph

Small

B

/

U

≡

1.

⚠

Critical Alert: Shipment & Orders Mismatch

The following orders have OTM releases in **Final/Failed** status while WMS remains in **Created/Allocated/Partial** status. **Allocated.**

Order Nbr	OTM Order Release GID	OTM Status	WMS Status	Ship Date	OTM Early Pickup
24022606	WMSOTM_26B_SY_06	Planned - Final	Allocated	2025-01-30	
24022603	WMSOTM_26B_SY_03	Planned - Final	Allocated	2025-01-30	

Warehouse Metrics Pane for Today’s Shipping Execution

The WMS metrics pane shows what must ship today and the related warehouse task types. Users can drill into order details and then use the **Priority Actions** pane to view the tasks for the orders.

Note: priority task updates will only be available if the task has not been started.

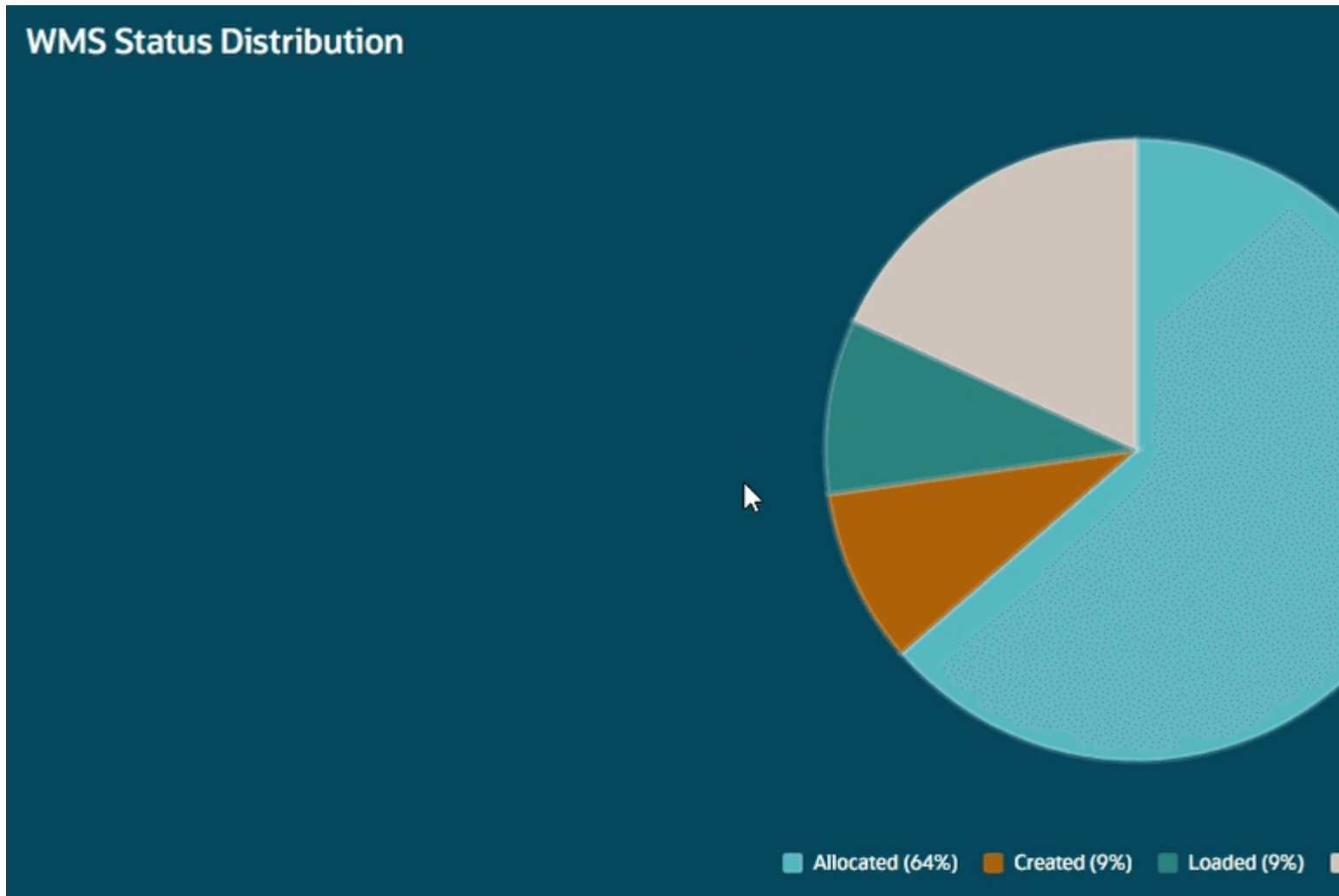
Warehouse Metrics

Pending Orders by Task Type

Here is a grouped summary of pending orders by task type with expanded details available

Order Nbr	Task Type	Priority
WMSOTM_26B_PS_14	LPNUNITS	0
WMSOTM_26B_SY_02	LPNUNITS	20
WMSOTM_26B_SY_03	LPNUNITS	20
WMSOTM_26B_SY_06	LPNUNITS	20

The **WMS Status Distribution** metric displays a distribution of the order status.



Planned Shipments Pane – an Early-Warning System

Delayed shipments or shipments missing confirmed dock appointments are flagged early, enabling managers to reschedule labor and warehouse activity to avoid idle time.

The **Status** column will give you an idea of the percentage of the Load that has been loaded and packed. The Load Statuses will include:

- Appointment missing – load doesn't have an appointment assigned to it.
- Appointment missed - the scheduled load appointment time has passed.
- Appointment at risk – the appointment will be happening in the next 30 minute window.

Logistics Managers can also trigger **Priority Actions** to schedule or request appointments, keeping warehouse flow synchronized with carrier arrival. The appointment will be scheduled based on available appointment times in OTM.

Note: the priority action for scheduling the appointments will only display if the start time is within 24 hours from the current time.

Planned Shipments

Planned Shipments Overview

Here is an overview of 8 shipments excluding those marked On-Time. The shipments cover routes primarily from SEATTLE M including Appointment Missed and Appointment Missing.

Shipment	Route & Duration	Summary	WMS Lo
02181 — Feb 19, 2026 3:40 PM MST	SEATTLE MANUFACTURING → SUMMIT AUTO • 1d 23h 24m	Distance 2497.51 MI • 10 pkgs	OSPPM1
02188 — Feb 18, 2026 6:16 PM MST	SEATTLE MANUFACTURING → SUMMIT AUTO • 2d 23h 24m	Distance 2497.51 MI • 10 pkgs	OSPPM1
02179 — Feb 24, 2026 10:00 AM MST	SEATTLE MANUFACTURING → BALA CYNWYD • 2d 13h 28m	Distance 3271.23 MI • 20 pkgs	OSPPM1

OTM Metrics Pane for Outbound Transportation Health

The OTM **Metrics Pane** provides an at-a-glance view of outbound flow (order counts, packages, weight, volume) and dynamically highlights the most relevant exceptions and KPIs for that day’s operations to support faster, better decisions.

Steps to Enable and Configure

Using the following steps, you can configure access to the **Logistics Execution Command Center** agentic app in WMS so that you can launch it directly from the WMS Redwood menu.

1.

From **Redwood Desktop**, configure the URL via the **External Site** WMS Module -> **external_url** screen parameter. The URL will open in a new browser tab, independent of WMS. No authentication, tokens, or WMS context are passed to the destination site.
- Note:

Normally you would enter the full URL path for the external_url parameter. However Oracle non-WMS applications configured in “SAAS app configuration” have a shortcut available. For example to configure an entry for the new “Logistics Execution Command Center” agentic application, the parameter should be set to (AI-AGENT-STUDIO)/hcmUI/redwoodAI/ora_logistics_execution_command_center.

You can refer to the domain part using the AI-AGENT-STUDIO entry that you would have configured in the **SAAS App Configuration** screen in redwood desktop

2.

Grant access to configure external URLs using the “**access external sites**” permission. This permission is granted by default to ADMIN. EMPLOYEE roles must be explicitly granted this permission.

6 Index

Index

Classification

Classification is the process of predicting which category (or “class”) a data point belongs to, based on patterns in the data. In machine learning, models are trained to recognize similarities among data points so they can accurately classify new, unseen inputs.

Example: Predicting whether an order is "on time" or "delayed."

Regression

Regression is a statistical technique used to model the relationship between a **dependent variable** and one or more **independent variables**. In machine learning, regression is used to predict continuous values.

Example: Estimating how many hours it will take to fulfill an order based on its size, contents, and location.

Tuple

A tuple is an **ordered collection** of elements, often used to represent a single, structured record of related values. Tuples are commonly used in programming and machine learning to organize data inputs.

Example: A tuple might look like this: (OrderID, ItemCount, ShipDate).

