

AI Credits Definition

Updated: September 9, 2026

Please Note: Treasure Data, Inc. dba Treasure AI may update these rates and terms from time to time in order to reflect the addition of new products and services and/or make other changes that are not adverse to Customer or its use of the Services. Any such changes shall be effective immediately upon publishing to this webpage.

AI Credits

Purchased Credits may be redeemed for any of the metered entitlements shown below, at the rate of one (1) Credit for the number of entitlements shown (with proportionate consumption of fractional Credits):

AI Credit Consumption Ratios: 1 Credit				
Meter	Availability			
	AI Agent Foundry	Personalization AI Suite	Engagement AI Suite	Creative AI Suite
600 Conversations	✓	✓	✓	✓
1 Million RT Profiles *	–	✓	✓	–
10 Million Personalization Calls	–	✓	–	–
15 Million Incoming Events	–	✓	✓	–
25 Million Trigger Activations	–	✓	✓	–
1 Million Email Messages	–	–	✓	–
100 Thousand In-App Messages	–	✓	✓	–
10 Million LINE OA Messages	–	–	✓	–
10 Million Mobile Push Messages	–	–	✓	–
10 Thousand SMS Messages	–	–	✓	–
40 Million Predictions - ML Models	–	✓	✓	✓
100 Million Predictions - RFM Model	–	✓	✓	✓
* Measured and incurred monthly.				
Calculation Examples:				
Customer consumes 550 Conversations = .92 Credits				
Customer consumes 300 Conversations + 500,000 Messages = 1.00 Credit				

AI Meter Definitions:

AI Meter	Definition
Conversations	A Conversation is the metric of usage of Treasure AI's Treasure Studio, Treasure Code, and agents created or used from the AI Agent Foundry. A Conversation is the unit used to meter AI usage between two entities—which may be a human and an agent or an agent and another agent—leveraging one of our standard Large Language Models (LLM(s)) as a basis. As an empirical baseline, one (1.0) Conversation corresponds approximately to a typical session of five back-and-forths using Treasure AI's base LLM. This five-back-and-forth example is illustrative and is not a fixed one-to-one conversion rule. Actual Conversation consumption reflects the amount and complexity of LLM usage in each session, including the model selected and the context processed. Shorter sessions may be metered as a fraction of a Conversation, measured to two decimal places, while longer or more complex sessions may be metered as multiple Conversations. Conversations are metered on all Production Instances.
Real Time Profile	A Real Time (RT) Profile is a collection of Core ICDP attributes and RT Attributes related to a single individual held in the Real Time Decision Engine. Each RT Profile may have up to a combined two hundred (200) Core ICDP attributes and RT Attributes, with a maximum size of 500 bytes per profile. Each RT Profile is counted as one (1) RT Profile, regardless of whether it is a Known RT Profile (has Personally Identifiable Information) or an Unknown RT Profile (no PII). RT Profile quantities are measured on a monthly basis to determine Credit consumption. The number of RT Profiles calculated for a given month will be the 4th highest number of RT Profiles held in all Production Instances counted on a daily basis during the month.
Personalization Calls	A Real Time Personalization Call (2.0) is a call to the Personalization API which, using a Real Time (RT) user identifier, returns a payload of up to 20 RT Attributes and Core ICDP attributes which are stored in the RT Decision Engine. A Web Personalization Call (1.0) is a Profile API call that uses native iOS and Android SDK infrastructure to provide Treasure AI Mobile SDK features for React Native apps. The number of Personalization Calls for a given period will be the total number of Personalization Calls recorded in all Production Instances for that period.
Incoming Events	An Incoming Event is an incoming streaming event that is ingested into the Real Time (RT) Decision Engine. The number of Incoming Events is the total number of Incoming Events recorded in the RT Decision Engine(s) in all Production Instances for a given period.
Trigger Activations	A Trigger Activation is an outgoing streaming event that is generated by the Real Time (RT) Decision Engine, which can result in a call to downstream streaming connections for an action to be taken. The number of Trigger Activations is the total number of outgoing events recorded in the RT Decision Engine(s) in all Production Instances for a given period.
Email Messages	Email Messages measures the number of messages sent by a customer using Engage Studio in all Production Instances within a given period.
Mobile Push Messages	Mobile Push Messages measures the number of messages sent by a customer using Engage Studio in all Production Instances within a given period.
In-App and In-Browser Messages	In-App Messages measures the number of personalized, visually rich messages directly inside a mobile app or browser. This feature is available to customers with Engage Studio, and is metered in all Production Instances within a given period.

AI Meter	Definition
LINE OA Messages	LINE OA messages measures the number of messages sent through campaigns through LINE Official Accounts. This feature is available to customers with Engage Studio, and is metered in all Production Instances within a given period.
SMS Messages	SMS Messages measures the number of SMS, MMS, RCS, or other text-based communication sent using the SMS/Messaging Service in Engage Studio in all Production Instances within a given period. Metering is based on the number of individual message segments sent, where a standard message is up to 160 GSM-7 characters (or 70 characters when non-GSM characters are present).
Predictions	Predictions are generated as a result of training and prediction runs of Profiles through a machine learning model. The models include but are not limited to Next Best Product (NBP), Next Best Action (NBA), Recency-Frequency-Monetary (RFM), and Customer Lifetime Value (CLTV). The number of Predictions are recorded as the number of outputs from a prediction run for all Production Instances for a given period.
AutoML Tier 1/2/3	AutoML automates machine learning (ML) process steps such as data processing, feature engineering, model selection and hyper-parameter tuning to reduce time to deploy ML models. It provides users access to the AutoML capabilities from Treasure Workflows. AutoML is available in the following options: AutoML Tier1: Up to 256GB RAM Maximum runtime hours: 500 Unit Hours 4 AutoML tasks can be run concurrently 5 TD users can run AutoML tasks AutoML Tier2: Up to 384GB Maximum runtime hours: 1000 Unit Hours 8 AutoML tasks can be run concurrently 15 TD users can run AutoML tasks AutoML Tier3: Up to 512GB RAM Maximum runtime hours: 2000 Unit Hours 16 AutoML tasks can be run concurrently Unlimited TD users can run AutoML tasks