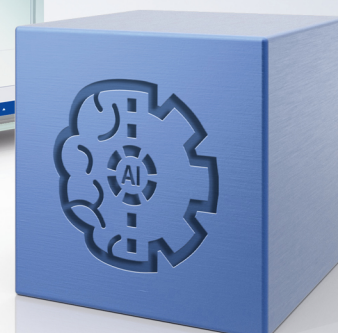
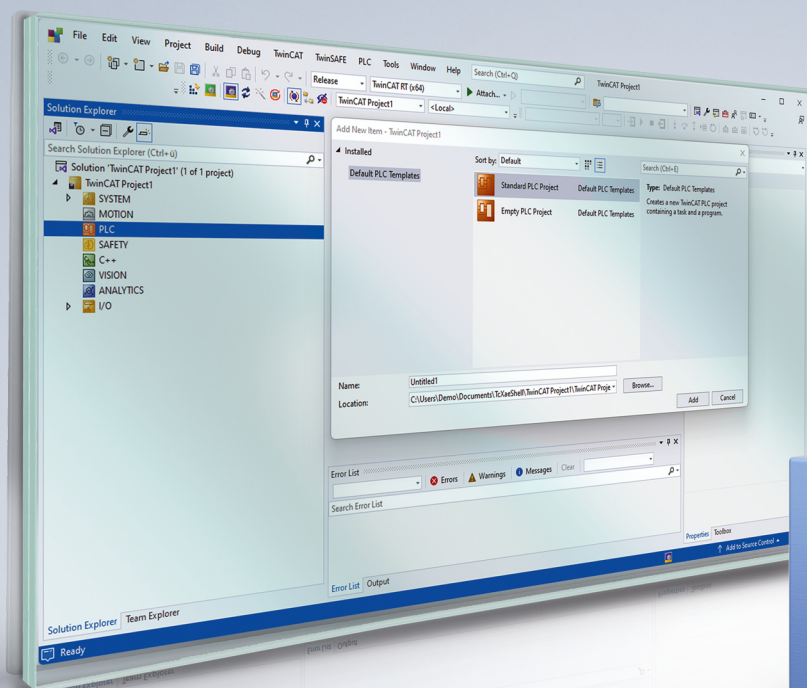


Handbuch | DE

TE3850

TwinCAT 3 | Machine Learning Creator



Inhaltsverzeichnis

1	Vorwort.....	5
1.1	Hinweise zur Dokumentation	5
1.2	Zu Ihrer Sicherheit.....	6
1.3	Hinweise zur Informationssicherheit	7
1.4	Ausgabestände dieser Dokumentation	8
2	Übersicht.....	9
3	Voraussetzungen	10
4	Lizenzierung	11
5	Workflows und UI-Ansichten	13
5.1	Navigation	13
5.2	Workflows für Lizenzadministratoren	13
5.2.1	Mitglieder der Organisation verwalten.....	15
5.2.2	Lizenzen verwalten	16
5.3	Workflows für User	16
5.3.1	Projekte anlegen	17
5.3.2	Datensatz erstellen und verwalten	18
5.3.3	Modelltraining	22
5.3.4	Modell validieren	25
5.3.5	Modell-Download	26
6	Support und Service	29

1 Vorwort

1.1 Hinweise zur Dokumentation

Diese Beschreibung wendet sich ausschließlich an ausgebildetes Fachpersonal der Steuerungs- und Automatisierungstechnik, das mit den geltenden nationalen Normen vertraut ist.

Zur Installation und Inbetriebnahme der Komponenten ist die Beachtung der Dokumentation und der nachfolgenden Hinweise und Erklärungen unbedingt notwendig.

Das Fachpersonal ist verpflichtet, stets die aktuell gültige Dokumentation zu verwenden.

Das Fachpersonal hat sicherzustellen, dass die Anwendung bzw. der Einsatz der beschriebenen Produkte alle Sicherheitsanforderungen, einschließlich sämtlicher anwendbaren Gesetze, Vorschriften, Bestimmungen und Normen erfüllt.

Disclaimer

Diese Dokumentation wurde sorgfältig erstellt. Die beschriebenen Produkte werden jedoch ständig weiterentwickelt.

Wir behalten uns das Recht vor, die Dokumentation jederzeit und ohne Ankündigung zu überarbeiten und zu ändern.

Aus den Angaben, Abbildungen und Beschreibungen in dieser Dokumentation können keine Ansprüche auf Änderung bereits gelieferter Produkte geltend gemacht werden.

Marken

Beckhoff®, ATRO®, EtherCAT®, EtherCAT G®, EtherCAT G10®, EtherCAT P®, MX-System®, Safety over EtherCAT®, TC/BSD®, TwinCAT®, TwinCAT/BSD®, TwinSAFE®, XFC®, XPlanar® und XTS® sind eingetragene und lizenzierte Marken der Beckhoff Automation GmbH.

Die Verwendung anderer in dieser Dokumentation enthaltenen Marken oder Kennzeichen durch Dritte kann zu einer Verletzung von Rechten der Inhaber der entsprechenden Kennzeichnungen führen.



EtherCAT® ist eine eingetragene Marke und patentierte Technologie, lizenziert durch die Beckhoff Automation GmbH, Deutschland.

Copyright

© Beckhoff Automation GmbH & Co. KG, Deutschland.

Weitergabe sowie Vervielfältigung dieses Dokuments, Verwertung und Mitteilung seines Inhalts sind verboten, soweit nicht ausdrücklich gestattet.

Zuwiderhandlungen verpflichten zu Schadenersatz. Alle Rechte für den Fall der Patent-, Gebrauchsmuster- oder Geschmacksmustereintragung vorbehalten.

Fremdmarken

In dieser Dokumentation können Marken Dritter verwendet werden. Die zugehörigen Markenvermerke finden Sie unter: <https://www.beckhoff.com/trademarks>.

1.2 Zu Ihrer Sicherheit

Sicherheitsbestimmungen

Lesen Sie die folgenden Erklärungen zu Ihrer Sicherheit.

Beachten und befolgen Sie stets produktspezifische Sicherheitshinweise, die Sie gegebenenfalls an den entsprechenden Stellen in diesem Dokument vorfinden.

Haftungsausschluss

Die gesamten Komponenten werden je nach Anwendungsbestimmungen in bestimmten Hard- und Software-Konfigurationen ausgeliefert. Änderungen der Hard- oder Software-Konfiguration, die über die dokumentierten Möglichkeiten hinausgehen, sind unzulässig und bewirken den Haftungsausschluss der Beckhoff Automation GmbH & Co. KG.

Qualifikation des Personals

Diese Beschreibung wendet sich ausschließlich an ausgebildetes Fachpersonal der Steuerungs-, Automatisierungs- und Antriebstechnik, das mit den geltenden Normen vertraut ist.

Signalwörter

Im Folgenden werden die Signalwörter eingeordnet, die in der Dokumentation verwendet werden. Um Personen- und Sachschäden zu vermeiden, lesen und befolgen Sie die Sicherheits- und Warnhinweise.

Warnungen vor Personenschäden

GEFAHR

Es besteht eine Gefährdung mit hohem Risikograd, die den Tod oder eine schwere Verletzung zur Folge hat.

WARNUNG

Es besteht eine Gefährdung mit mittlerem Risikograd, die den Tod oder eine schwere Verletzung zur Folge haben kann.

VORSICHT

Es besteht eine Gefährdung mit geringem Risikograd, die eine mittelschwere oder leichte Verletzung zur Folge haben kann.

Warnung vor Umwelt- oder Sachschäden

HINWEIS

Es besteht eine mögliche Schädigung für Umwelt, Geräte oder Daten.

Information zum Umgang mit dem Produkt



Diese Information beinhaltet z. B.:
Handlungsempfehlungen, Hilfestellungen oder weiterführende Informationen zum Produkt.

1.3 Hinweise zur Informationssicherheit

Die Produkte der Beckhoff Automation GmbH & Co. KG (Beckhoff) sind, sofern sie online zu erreichen sind, mit Security-Funktionen ausgestattet, die den sicheren Betrieb von Anlagen, Systemen, Maschinen und Netzwerken unterstützen. Trotz der Security-Funktionen sind die Erstellung, Implementierung und ständige Aktualisierung eines ganzheitlichen Security-Konzepts für den Betrieb notwendig, um die jeweilige Anlage, das System, die Maschine und die Netzwerke gegen Cyber-Bedrohungen zu schützen. Die von Beckhoff verkauften Produkte bilden dabei nur einen Teil des gesamtheitlichen Security-Konzepts. Der Kunde ist dafür verantwortlich, dass unbefugte Zugriffe durch Dritte auf seine Anlagen, Systeme, Maschinen und Netzwerke verhindert werden. Letztere sollten nur mit dem Unternehmensnetzwerk oder dem Internet verbunden werden, wenn entsprechende Schutzmaßnahmen eingerichtet wurden.

Zusätzlich sollten die Empfehlungen von Beckhoff zu entsprechenden Schutzmaßnahmen beachtet werden. Weiterführende Informationen über Informationssicherheit und Industrial Security finden Sie in unserem <https://www.beckhoff.de/secguide>.

Die Produkte und Lösungen von Beckhoff werden ständig weiterentwickelt. Dies betrifft auch die Security-Funktionen. Aufgrund der stetigen Weiterentwicklung empfiehlt Beckhoff ausdrücklich, die Produkte ständig auf dem aktuellen Stand zu halten und nach Bereitstellung von Updates diese auf die Produkte aufzuspielen. Die Verwendung veralteter oder nicht mehr unterstützter Produktversionen kann das Risiko von Cyber-Bedrohungen erhöhen.

Um stets über Hinweise zur Informationssicherheit zu Produkten von Beckhoff informiert zu sein, abonnieren Sie den RSS Feed unter <https://www.beckhoff.de/secinfo>.

1.4 Ausgabestände dieser Dokumentation

Version	Änderungen
1.0.0	Erste Veröffentlichung

2 Übersicht

Der TwinCAT **M**achine **L**earning **C**reator (ML-Creator, MLC) ist eine webbasierte Engineering-Umgebung zur Erstellung von KI-Modellen für industrielle Anwendungen. Die Software führt den Anwender durch den standardisierten Workflow aus Datenaufbereitung, Modelltraining, Validierung und Bereitstellung.

Ziel ist die automatisierte Generierung von KI-Modellen ohne tiefgehende Kenntnisse in Data Science oder Machine Learning. Der Fokus liegt auf der Unterstützung von Automatisierungs- und Prozessexperten bei der Umsetzung datenbasierter Anwendungen.

Typische Anwendungen sind die KI-basierte Qualitätsprüfung, Klassifikation und Prozessüberwachung auf Basis von Bilddaten.

Weitere aktuelle Anwendungsbeispiele finden Sie auf der Startseite des MLC unter: www.beckhoff.com/mlc

3 Voraussetzungen

Webbrowser

Für die Nutzung des TwinCAT Machine Learning Creator ist ein aktueller Webbrowser erforderlich. Unterstützt werden gängige Browser wie Google Chrome, Microsoft Edge oder Mozilla Firefox in aktueller Version.

MyBeckhoff Nutzerkonto

Zur Authentifizierung beim TwinCAT Machine Learning Creator ist ein gültiges MyBeckhoff-Nutzerkonto notwendig.

Falls Sie noch kein Nutzerkonto erstellt haben, können Sie dies hier tun:

<https://www.beckhoff.com/en-en/mybeckhoff-registration/index.aspx>

Das MyBeckhoff-Nutzerkonto ist ein zentralisiertes Nutzerkonto für eine Reihe von Beckhoff-Diensten. Sie können es bspw. auch verwenden, um TwinCAT Produkte über den TwinCAT Package Manager herunterzuladen oder sich beim MX-Designer anzumelden.

4 Lizenzierung

Überblick

Der TwinCAT Machine Learning Creator wird als abonnementbasierter Cloud-Service bereitgestellt. Die Nutzung erfolgt über benutzergebundene Lizenzen mit definierter Laufzeit und begrenztem Kontingent an Rechenressourcen (Compute Hours).

Für produktive Anwendungen stehen verschiedene Lizenzvarianten zur Verfügung, die sich hinsichtlich Laufzeit (monatlich oder jährlich), verfügbarer Compute Hours sowie Nutzungsumfang unterscheiden.

Die Lizenzen ermöglichen die Erstellung eigener KI-Modelle auf Basis von Bilddaten. Je nach Lizenztyp können die erzeugten Modelle entweder ausschließlich in TwinCAT ausgeführt oder als ONNX-Dateien für externe Laufzeitumgebungen exportiert werden.

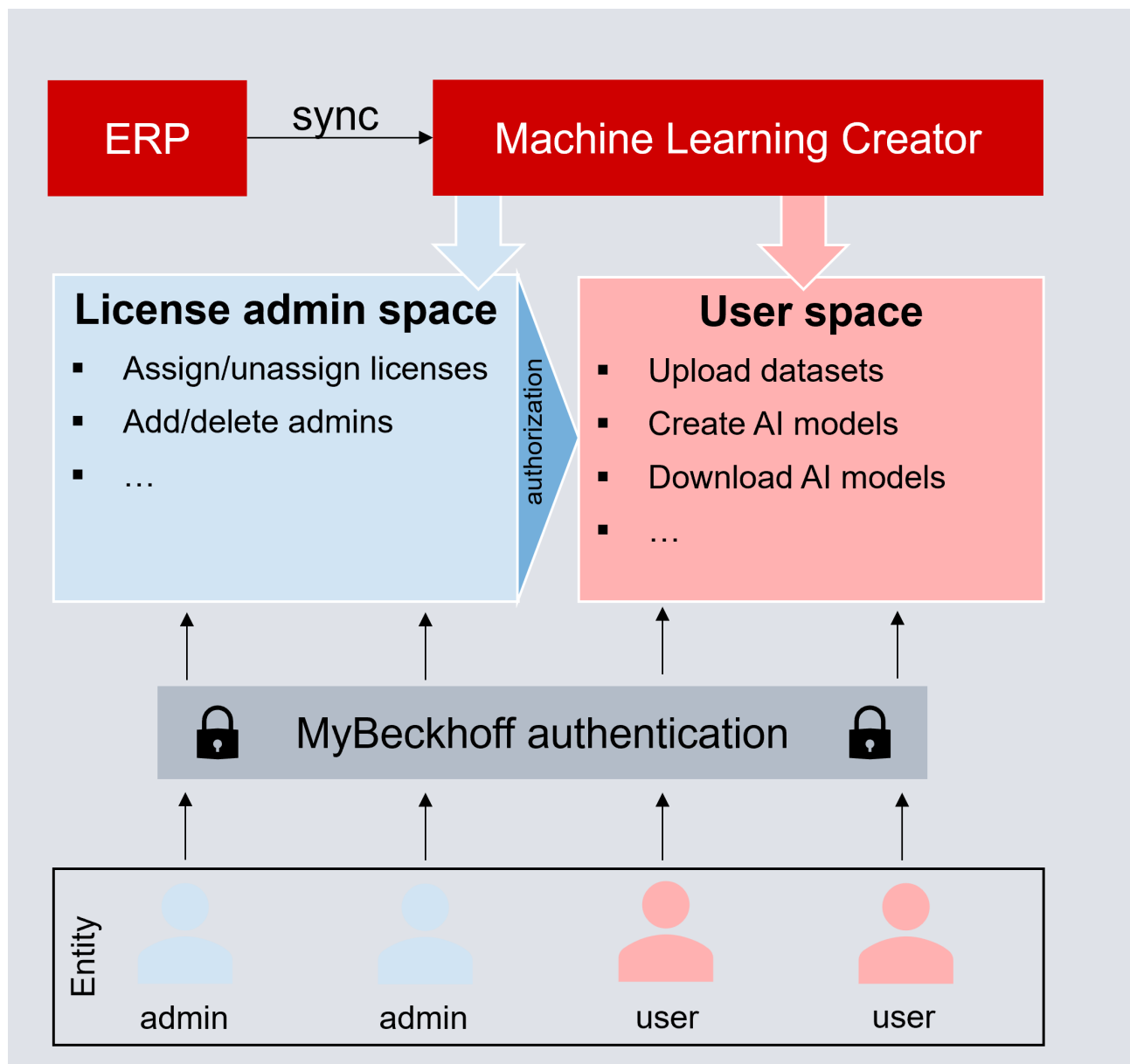
Zusätzliche Rechenressourcen können über Erweiterungspakete bereitgestellt werden. Nicht genutzte Compute Hours verfallen am Ende der jeweiligen Laufzeit.

Weitere Informationen finden Sie auf der Produktseite unter: <https://www.beckhoff.com/te3850>

User Onboarding

Die Lizenzverwaltung erfolgt zentral über das Beckhoff ERP und wird automatisch mit dem TwinCAT Machine Learning Creator synchronisiert. Eine manuelle Lizenzübertragung über License Request- und Response-Dateien, wie in vielen TwinCAT Produkten üblich, ist nicht erforderlich.

Der Zugriff auf die Plattform erfolgt über ein MyBeckhoff-Nutzerkonto. Lizenzen werden innerhalb einer Organisation durch definierte Lizenzadministratoren [► 13] verwaltet. Diese weisen einzelnen Benutzern [► 16] die verfügbaren Lizenzen zu oder entziehen sie ihnen bei Bedarf.



Lizenzadministratoren können alle zugewiesenen Lizenzen ihrer Organisation einsehen und steuern den Zugriff auf Funktionen entsprechend dem jeweiligen Lizenzumfang.

Mit der allerersten Bestellung einer MLC-Lizenz durch einen Kunden wird ein erster Lizenzadministrator festgelegt. Nachfolgend können Lizenzadministratoren eigenständig hinzugefügt und entfernt werden.

Die Bestellung der MLC-Lizenzen erfolgt, wie für alle anderen Beckhoff Produkte, über den Beckhoff Vertrieb.

5 Workflows und UI-Ansichten

5.1 Navigation

Navigationselemente

Der TwinCAT Machine Learning Creator verwendet zwei zentrale Navigationselemente:
Die Benutzernavigation sowie das Seitenmenü.

Benutzernavigation

Die Benutzernavigation befindet sich oben rechts in der Kopfzeile der Anwendung.

Über die Benutzernavigation stehen allgemeine Benutzer- und Verwaltungsfunktionen zur Verfügung.
Zusätzlich kann über die Kopfzeile die Sprache der Benutzeroberfläche ausgewählt werden.

Folgende Bereiche können über die Benutzernavigation geöffnet werden (hier am Beispiel der Sprache en-US):

- Projects Overview
- Recent Updates
- My Profile
- License Admin (nur für Lizenzadministratoren)
- Contact
- Sign Out

Seitenmenü

Das Seitenmenü befindet sich auf der linken Seite der Benutzeroberfläche.

Über das Seitenmenü erfolgt die Navigation innerhalb eines Projekts.

Das Seitenmenü kann ein- und ausgeklappt werden.

Je nach Kontext stehen unter anderem folgende Bereiche zur Verfügung:

- [Dataset \[► 18\]](#)
- [Models \[► 22\]](#)
- Performance

Über die einzelnen Menüeinträge kann zwischen den jeweiligen Projektansichten gewechselt werden.

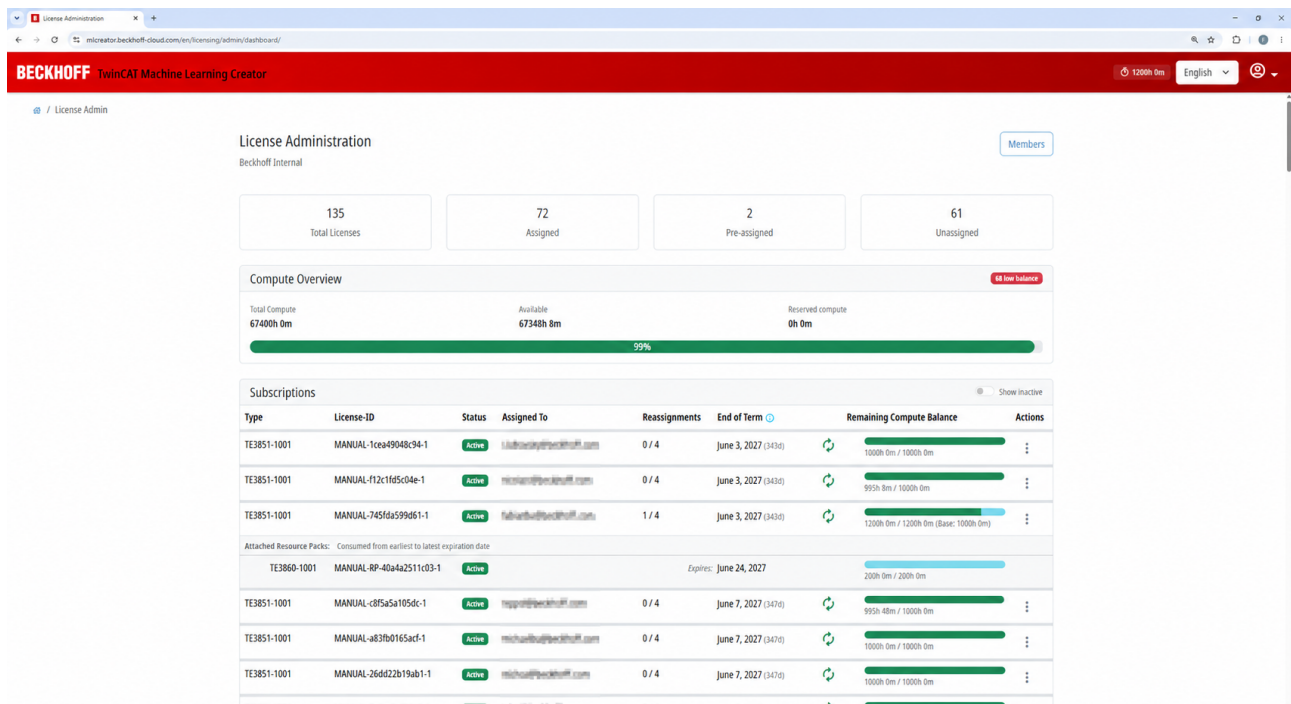
Home Button (Beckhoff Logo)

Durch Klicken auf das Beckhoff Logo wird für Lizenzadministratoren das [Lizenzadministratoren-Panel \[► 13\]](#) und für User der [Project Overview \[► 17\]](#) geöffnet.

5.2 Workflows für Lizenzadministratoren

Lizenzadministration

Das Lizenzadministratoren-Panel dient der zentralen Verwaltung aller dem Unternehmen zugeordneten MLC-Lizenzen und -Ressourcen. Der Zugriff ist ausschließlich für Benutzer mit der Rolle „Lizenzadministrator“ möglich.



Members

Über die Schaltfläche **Members** [► 15] können Benutzer dem jeweiligen Unternehmen zugeordnet werden. Nur Member können zum Lizenzadministrator oder User befähigt werden.

Übersicht

Im oberen Bereich werden aggregierte Kennzahlen zur Lizenzverteilung angezeigt:

- **Total Licenses:** Gesamtanzahl verfügbarer Lizenzen
- **Assigned:** Anzahl aktuell zugewiesener Lizenzen
- **Pre-assigned:** Anzahl bereits vorab zugewiesener Lizenzen (Member hat noch nicht bestätigt)
- **Unassigned:** Anzahl nicht zugewiesener Lizenzen

Diese Werte geben einen schnellen Überblick über die aktuelle Lizenzverteilung.

Compute Overview

Der Bereich „Compute Overview“ zeigt die verfügbaren Rechenressourcen:

- **Total Compute:** Gesamtes Compute-Kontingent aller Lizenzen und Resource Packs
- **Available:** Noch verfügbares Compute-Kontingent
- **Reserved compute:** Bereits reservierte, aber noch nicht verbrauchte Ressourcen

Die Fortschrittsanzeige visualisiert die aktuelle Auslastung. Ein Hinweis auf ein niedriges Restkontingent wird bei Bedarf direkt in diesem Bereich angezeigt.

Subscriptions

In der Tabelle „Subscriptions“ werden alle Lizenzzuweisungen im Detail dargestellt:

- **Type:** Lizenztyp (z. B. TE3851-1000)
- **License-ID:** Eindeutige Lizenzkennung
- **Status:** Aktueller Status (z. B. Active, Grace, Expired)
- **Assigned To:** Zugewiesener Benutzer
- **Reassignments:** Anzahl verbleibender Neuzuweisungen innerhalb der Laufzeit
- **End of Term:** Ablaufdatum des aktuellen Terms

- **Remaining Compute Balance:** Verbleibendes Compute-Kontingent der Lizenz

Aktionen

Über das Drei-Punkte-Menü stehen je nach Lizenzstatus folgende Aktionen [► 16] zur Verfügung:

- **Assign:** Lizenz einem Benutzer zuweisen
- **Reassign:** Übertragung der Lizenz auf einen anderen Benutzer
- **Unassign:** Entfernen der aktuellen Zuweisung
- **Attach Pack:** Zuweisen eines Resource Packs zur Erweiterung des Compute-Kontingents

Attached Resource Packs

Wenn ein Resource Pack einer Lizenz zugewiesen wurde, wird es im Bereich *Attached Resource Packs* direkt unter der zugehörigen Subscription angezeigt. Dabei gilt:

- Das Resource Pack ist an eine konkrete Lizenzinstanz gebunden. Es wird nicht separat einem User, sondern der Lizenz zugeordnet.
- Das zugewiesene Pack erhöht das verfügbare Compute-Kontingent der Lizenz.
- Die Anzeige zeigt das Pack mit eigener License-ID, Status und Ablaufdatum.
- Die Compute-Anzeige der Lizenz berücksichtigt das zusätzliche Kontingent aus dem Pack.

Unattached Resource Packs

Dieser Bereich zeigt verfügbare, noch nicht zugewiesene Resource Packs:

- **Type:** Typ des Resource Packs, z. B. TE3860-1001
- **License-ID:** Eindeutige Kennung des Resource Packs
- **Compute:** Enthaltene Compute Hours
- **Expiry:** Ablaufdatum
- **Status:** Aktueller Status
- **Compatible Licenses:** Mit welchen Lizenztypen das Resource Pack verwendet werden kann

Über den Schalter **Show inactive** können inaktive Resource Packs ein- oder ausgeblendet werden. Resource Packs werden erst nach der Zuweisung zu einer kompatiblen Lizenz im Bereich Attached Resource Packs angezeigt.

5.2.1 Mitglieder der Organisation verwalten

Über den Button „Member“ gelangen Sie zur Mitgliederverwaltungsansicht.

Neue Benutzer werden durch einen Lizenzadministrator zur Organisation eingeladen. Dazu wird über **Invite Member** die E-Mail-Adresse des neuen Benutzers angegeben. Optional kann hier bereits eine Lizenz vorgegeben werden. Der User bekommt dann diese Lizenz automatisch zugeordnet, sobald dieser die Einladung zur Organisation akzeptiert hat.



MyBeckhoff-Konto erfordert

Im Regelfall sollte die verwendete E-Mail-Adresse bereits einer bestehenden MyBeckhoff-ID zugeordnet sein. Existiert noch kein MyBeckhoff-Konto, muss der Benutzer zunächst ein Konto mit dieser E-Mail-Adresse erstellen.

Nach dem Versand der Einladung erhält der Benutzer eine E-Mail zur Verifikation. Die Einladung muss über den enthaltenen Link bestätigt werden.

Bis zur Bestätigung befindet sich der Benutzer im Status **Pending**.

Nach erfolgreicher Bestätigung wechselt der Status auf **Member**. Danach können dem Benutzer MLC-Lizenzen zugewiesen werden (pre-assigned licenses werden automatisch zugewiesen, sobald der Status Member erreicht wird).

Lizenzadministratoren können (unter **Actions**) Mitgliederrollen verwalten:

- **Promote:** Erhöht ein Mitglied zum Lizenzadministrator.
- **Demote:** Entfernt Administratorrechte und setzt den Benutzer auf den Status *Member* zurück.
- **Remove:** Entfernt den Member aus der Organisation.

Ein Lizenzadministrator kann die eigenen Administratorrechte nicht selbst entfernen. Dadurch wird sichergestellt, dass jederzeit mindestens ein Lizenzadministrator für die Organisation vorhanden ist.

5.2.2 Lizenzen verwalten

Die Verwaltung von Subscription-Lizenzen und Resource Packs erfolgt zentral über das Lizenzadministrationen-Panel.

Subscription-Lizenzen zuweisen

Subscription-Lizenzen können ausschließlich Benutzern mit dem Status **Member** zugewiesen werden. Die Zuweisung erfolgt durch einen Lizenzadministrator.

Jede Subscription-Lizenz kann immer nur einem Benutzer gleichzeitig zugeordnet sein. Die aktuelle Zuweisung wird in der Spalte **Assigned To** angezeigt.

Über **Reassign** kann eine Lizenz einem anderen Benutzer zugewiesen werden. Die Anzahl möglicher Neuzuweisungen innerhalb der Laufzeit ist begrenzt und wird in der Spalte **Reassignments** dargestellt.

Die zulässige Anzahl der Reassignments ist abhängig vom Lizenztyp:

- Jährliche Subscription-Lizenzen: bis zu vier Zuweisungen pro Laufzeit
- Monatliche Subscription-Lizenzen: eine Zuweisung pro Laufzeit

Nicht verbrauchte Reassignments verfallen am Ende der jeweiligen Subscription-Periode.

Lizenzzuweisung entfernen

Über **Unassign** kann die aktuelle Benutzerzuweisung einer Subscription-Lizenz entfernt werden. Die Lizenz steht anschließend wieder zur erneuten Zuweisung innerhalb der erlaubten Reassignment-Grenzen zur Verfügung.

Resource Packs zuweisen

Resource Packs erweitern das verfügbare Compute-Kontingent einer Subscription-Lizenz. Resource Packs werden nicht direkt Benutzern, sondern immer einer kompatiblen Basislizenz zugewiesen.

Nicht zugewiesene Resource Packs werden im Bereich **Unattached Resource Packs** angezeigt.

Über **Attach Pack** kann ein Resource Pack mit einer Subscription-Lizenz verknüpft werden. Nach der Zuordnung bleibt das Resource Pack während der gesamten Laufzeit an diese Basislizenz gebunden und kann nicht auf andere Lizenzen übertragen werden.

Compute Hours werden automatisch anhand des frühesten Ablaufdatums verbraucht. Dadurch werden zuerst Ressourcen genutzt, deren Verfallsdatum am nächsten liegt.

5.3 Workflows für User

Workflowübersicht

Der TwinCAT Machine Learning Creator bildet den vollständigen Workflow zur Erstellung und Bereitstellung von KI-Modellen in einer web-basierten Benutzeroberfläche ab.

Der Workflow besteht aus den folgenden Schritten:

Projekt erstellen

Zunächst wird ein Projekt auf der Plattform angelegt. Das Projekt definiert den Aufgabentypen (z. B. Image Classification) und dient als Rahmen für Datensätze und Modelle.

Datensatz erstellen

Zunächst werden Bilddaten in die Plattform hochgeladen. Die Bilder werden anschließend (oder direkt beim Hochladen) mit Labels versehen, die die gewünschte Zielklasse definieren. Der Datensatz bildet die Grundlage für das spätere Modelltraining.

Modelltraining konfigurieren

Für das Training wird ein neues Modell angelegt und mit einem Datensatz verknüpft. Zusätzlich können trainingsrelevante Parameter sowie Zielhardware, Zielsoftware und Laufzeitanforderungen definiert werden.

Modelltraining ausführen

Nach dem Start des Trainings erzeugt die Plattform automatisiert ein KI-Modell auf Basis der bereitgestellten Trainingsdaten.

Modell validieren und analysieren

Nach Abschluss des Trainings stehen verschiedene Analyse- und Explainability-Funktionen zur Verfügung. Dazu gehören unter anderem statistische Bewertungsverfahren, Konfidenzen, Confusion-Matrizen sowie visuelle Darstellungen zur Nachvollziehbarkeit der Modellentscheidung.

Modell exportieren

Trainierte Modelle können anschließend heruntergeladen werden. Zusätzlich kann PLCopen-XML exportiert werden, um die Integration in TwinCAT-Projekte zu vereinfachen.

5.3.1 Projekte anlegen

Projektübersicht

Nach dem Login wird die Projektübersicht angezeigt. Hier können bestehende Projekte verwaltet sowie neue Projekte angelegt werden.

Ein Projekt dient als organisatorische Einheit für Datensätze, Trainingsläufe und KI-Modelle.

Neues Projekt anlegen

Über **Create Project** wird ein neues Projekt erstellt. Dabei werden Projektname und optionale Metadaten definiert.

Projekt öffnen

Über **Open** wird ein Projekt geöffnet. Anschließend können Datensätze erstellt, Modelle trainiert und Ergebnisse analysiert werden.

Projekt bearbeiten

Über das Symbol  können die Projektmetadaten bearbeitet werden:

- Projektname
- Projektbeschreibung

Projekt zu Favoriten hinzufügen

Über das Symbol  kann ein Projekt als Favorit markiert werden. Favorisierte Projekte werden im Bereich **Favorites** angezeigt.

Projekt löschen

Über das Symbol  kann ein Projekt gelöscht werden.

Tags verwalten

Über **Add tag** können Projekte mit Tags versehen werden. Tags dienen der organisatorischen Strukturierung und erleichtern das Filtern und Auffinden von Projekten.

Suche und Filter

Über das Suchfeld können Projekte anhand ihres Namens oder ihrer Metadaten durchsucht werden. Zusätzlich stehen Filterfunktionen zur Eingrenzung der angezeigten Projekte zur Verfügung.

Sortierung

Die Projektliste wird standardmäßig chronologisch nach der letzten Änderung sortiert. Projekte mit der neuesten Aktivität werden zuerst angezeigt.

5.3.2 Datensatz erstellen und verwalten

Datensätze bilden die Grundlage für das Training von KI-Modellen. Ein Datensatz enthält die Bilddaten sowie die zugehörigen Labels für die gewünschten Zielklassen.

Datensatz erstellen

Mit **Create Dataset** wird ein neuer Datensatz erstellt. Ein Datensatzname muss vergeben werden, weitere Metainformationen können optional eingetragen werden.

Datensatzinformationen

Im linken Bereich werden die wichtigsten Informationen des Datensatzes angezeigt:

- Name des Datensatzes
- Beschreibung
- Zeitpunkt der letzten gespeicherten Version
- Anzahl enthaltener Dateien
- Trainings- und Testsplitt
- Ergebnis einer ersten ad-hoc Datensatzvalidierung

Standardmäßig wird der Datensatz automatisch in Trainings- und Testdaten aufgeteilt:

- 80 % Trainingsdaten
- 20 % Testdaten

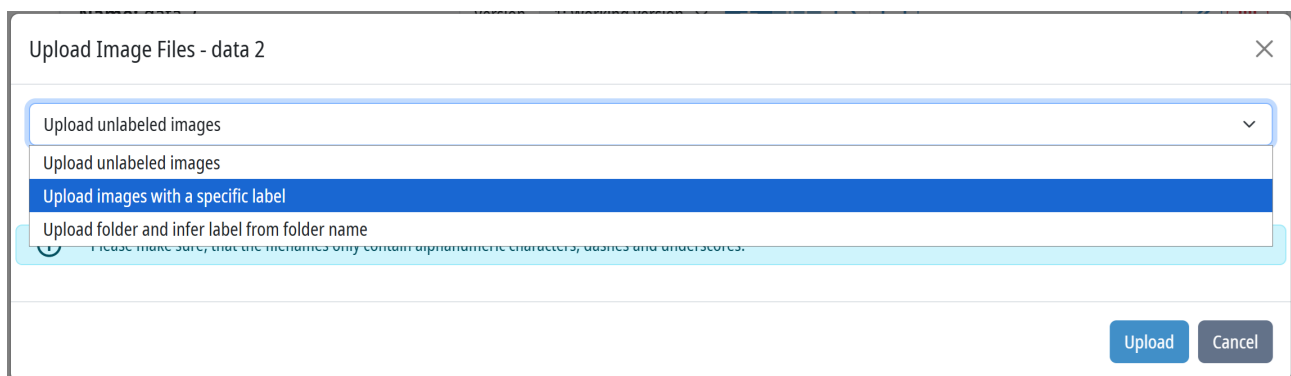
Die Testdaten werden ausschließlich zur späteren Modellvalidierung verwendet.

Aktionen

Bilder hochladen

Über **Upload Image Files** können Bilddateien in den Datensatz importiert werden.

Unterstützt werden gängige Bildformate.



Labels importieren

Über **Import Labels** können bestehende Labelinformationen importiert werden. Dadurch können bereits vorhandene Annotationen aus externen Werkzeugen übernommen werden.

Import Labels - data 2

Example for valid CSV-File

Example for valid JSON-File

File field*

Choose File No file chosen

```
{
  "categories": [
    {"id": 1, "name": "label1"},
    {"id": 2, "name": "label2"}
  ],
  "images": [
    {"id": 1, "file_name": "image1.jpg"},
    {"id": 2, "file_name": "image2.jpg"}
  ],
  "annotations": [
    {"image_id": 1, "category_id": 1},
    {"image_id": 2, "category_id": 2}
  ]
}
```

Label import will only work as expected if the uploaded files have unique names in the dataset.

Submit

Cancel

Daten können alternativ auch auf der Plattform gelabelt werden [► 20].

Datensatzversionen

Anzeige und Auswahl der Datensatzversion. Datensatzversionen stellen einen festen und reproduzierbaren Stand des Datensatzes dar und dienen als Grundlage für das Modelltraining.

Eine noch nicht versionierte Version wird als „Working version“ angezeigt.

Neue Datensatzversion erstellen

Über das **Symbol +** wird eine neue Datensatzversion erzeugt.

Dabei wird der aktuelle Zustand des Datensatzes gespeichert. Nur gespeicherte Datensatzversionen können für das Training von KI-Modellen verwendet werden.

Datensatz klonen

Über das **Symbol**  kann ein Datensatz dupliziert werden.

Der geklonte Datensatz enthält alle Bilder, Labels und Metadaten des ursprünglichen Datensatzes und kann unabhängig weiterbearbeitet werden.

Label Manager

Über das Symbol  wird der Label Manager geöffnet. Hier können Klassenlabels erstellt, bearbeitet und verwaltet werden.

File Explorer

Über das Symbol  wird der File Explorer geöffnet.

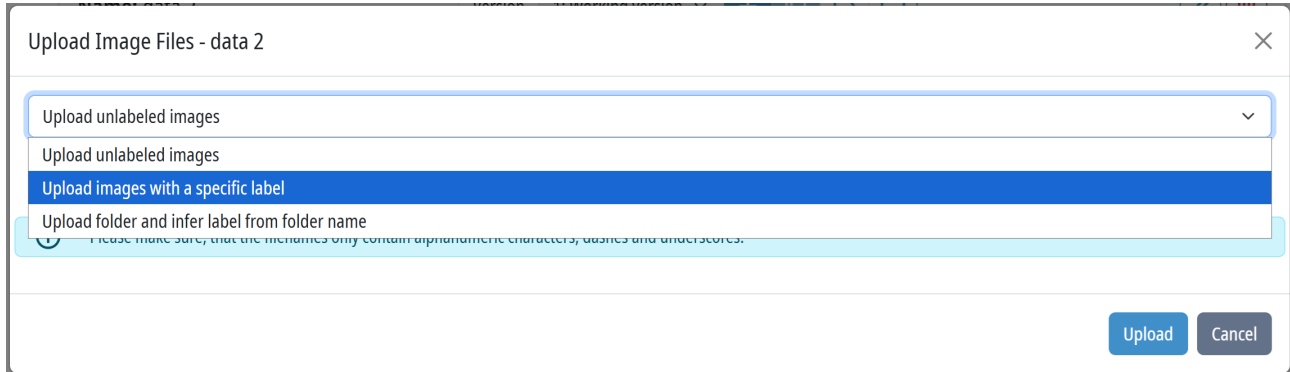
Der File Explorer ermöglicht die vollständige Ansicht und Navigation aller Dateien innerhalb des Datensatzes. Dadurch können Bilder kontrolliert, gefiltert und einzeln überprüft werden. Darüber hinaus können im File Explorer auch Labels erstellt und verändert werden.

5.3.2.1 Labeling Optionen für Bildklassifikation

Für Image-Classification-Datensätze stehen verschiedene Möglichkeiten zur Erstellung und Verwaltung von Labels zur Verfügung.

Labeling beim Datei-Upload

Labels können bereits während des Uploads der Bilddaten vergeben werden.



Labelzuweisung für einen Upload-Batch

Beim Upload kann festgelegt werden, dass alle ausgewählten Bilder einem bestimmten Label zugeordnet werden.

Durch mehrfaches batchweises Hochladen mit unterschiedlichen Labels kann ein Datensatz mit mehreren Klassen erstellt werden.

Beispiel:

- Upload Batch 1 → Label OK
- Upload Batch 2 → Label NOK

Labeling über Ordnerstruktur

Alternativ können Bilder bereits lokal in Unterordnern organisiert werden.

Beim Upload eines Verzeichnisbaums wird der Name des jeweiligen Unterordners automatisch als Label verwendet.

Beispiel:

```
dataset/
├── OK/
│   ├── img_001.png
│   └── img_002.png
└── NOK/
    ├── img_003.png
    └── img_004.png
```

In diesem Beispiel werden automatisch die Labels OK und NOK erzeugt und den jeweiligen Bildern zugeordnet.

Import externer Labels

Labels können aus externen Annotationstools importiert werden.

Import Labels - data 2

Example for valid CSV-File

Example for valid JSON-File

File field*

Choose File No file chosen

```
{
  "categories": [
    {"id": 1, "name": "label1"},
    {"id": 2, "name": "label2"}
  ],
  "images": [
    {"id": 1, "file_name": "image1.jpg"},
    {"id": 2, "file_name": "image2.jpg"}
  ],
  "annotations": [
    {"image_id": 1, "category_id": 1},
    {"image_id": 2, "category_id": 2}
  ]
}
```

Label import will only work as expected if the uploaded files have unique names in the dataset.

Submit

Cancel

Der Import erfolgt über **Import labels**.

Unterstützte Formate:

- CSV
- COCO

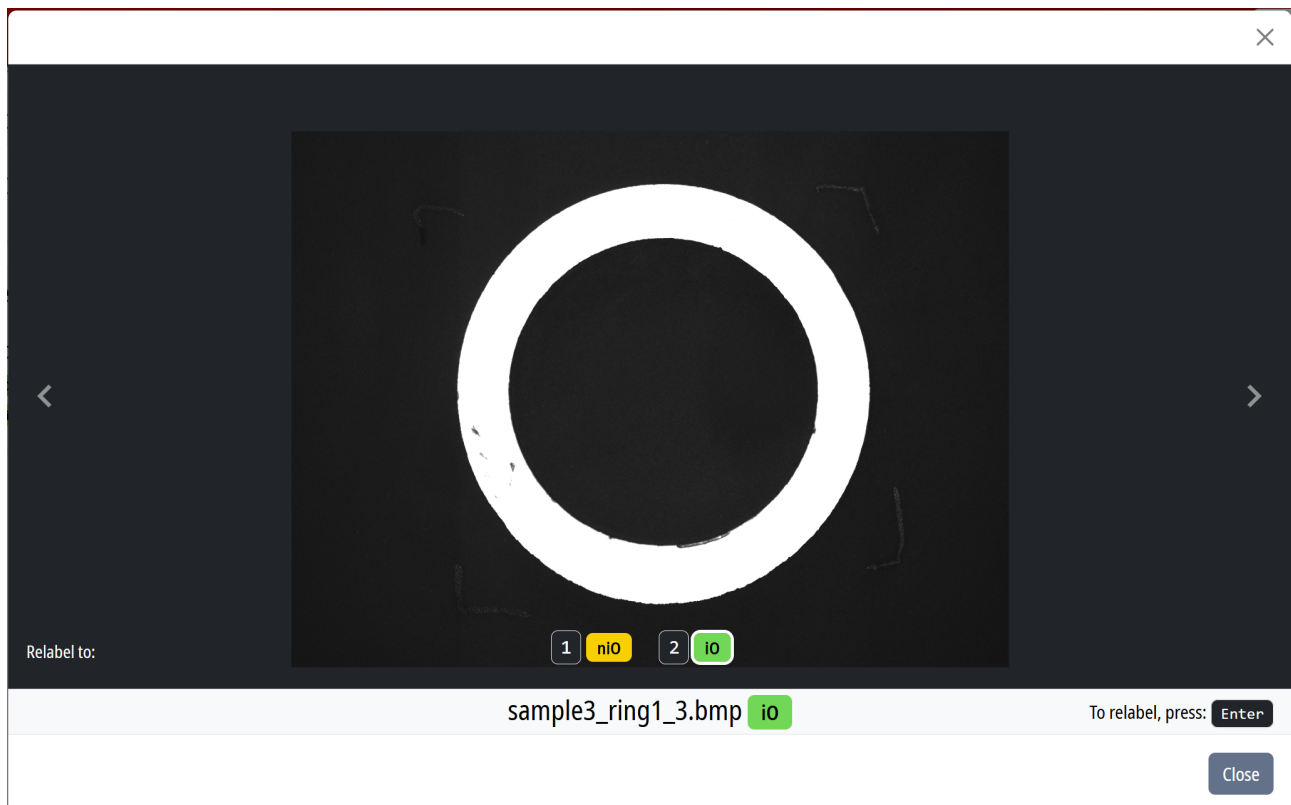
Dadurch können bestehende Annotationen aus externen Workflows übernommen werden.

Labeling im TwinCAT Machine Learning Creator

Labels können direkt innerhalb des TwinCAT Machine Learning Creator erstellt und bearbeitet werden.

Möglichkeiten über den **Label Manager**:

- Neue Labels anlegen
- Labels umbenennen
- Bestehende Labels verwalten



Im **File Explorer** können Bilder einzeln gelabelt werden. Dazu wird ein Bild ausgewählt und anschließend mit der Taste **Enter** in den Label-Modus gewechselt. Im Label-Modus können Bilder bestehenden Klassen zugeordnet oder neu gelabelt werden. Über die Pfeiltasten links und rechts können Sie schnell durch den Datensatz navigieren.

5.3.3 Modelltraining

Modell erstellen

Über **Create Model** wird ein neues KI-Modell angelegt und für das Training konfiguriert.

Voraussetzung für die Modellerstellung ist, dass bereits ein Datensatz hochgeladen und mindestens eine Datensatzversion gespeichert wurde. Nur versionierte Datensätze können für das Training verwendet werden.

Allgemeine Modellinformationen

- **Name:** Eindeutiger Name des Modells.
- **Description:** Optionale Beschreibung des Modells und des vorgesehenen Einsatzzwecks.

Datensatzkonfiguration

- **Dataset:** Auswahl des Datensatzes, der für das Training verwendet werden soll. Es werden nur Datensätze angezeigt, für die bereits eine Version angelegt wurde.
- **Dataset version:** Auswahl der zu verwendenden Datensatzversion. Die Version stellt einen festen und reproduzierbaren Stand des Datensatzes dar. Änderungen am Datensatz nach der Versionierung beeinflussen bestehende Trainingskonfigurationen nicht.

Laufzeit- und Zielsystemkonfiguration

- **Latency threshold in ms:** Definiert die maximal zulässige Ausführungszeit des KI-Modells auf dem Zielsystem [► 24]. Dieser Wert wird während der Modellerstellung berücksichtigt, um ein geeignetes Modell hinsichtlich Genauigkeit und Laufzeitverhalten auszuwählen.
- **Target device:** Auswahl der Zielhardware, auf der das Modell später ausgeführt werden soll.
- **CPU:** Auswahl der verwendeten CPU-Konfiguration des Zielsystems.

- **GPU option:** Optionale Auswahl einer unterstützten GPU-Konfiguration.
- **Target software:** Definiert die TwinCAT-Zielumgebung für die spätere Inferenz, beispielsweise:
 - [TF3820 | TwinCAT 3 Machine Learning Server](#)
 - [TF7810 | TwinCAT 3 Vision Neural Network](#)
- **Execution provider (EP):** Definiert die Ausführungsumgebung der Inferenzengine Machine Learning Server, beispielsweise CPU oder CUDA (NVIDIA GPU).
- **Number of cores used for inference:** Legt die Anzahl der CPU-Kerne fest, die für die spätere Modellausführung verwendet werden. Bei Verwendung von TF7810 bezieht sich dieser Wert auf isolierte CPU-Kerne innerhalb des Vision Job Pools. Bei Verwendung von TF3820 mit Execution Provider = CPU entspricht der Wert der Anzahl verfügbarer CPU-Kerne des User-Mode-Prozesses des TwinCAT Machine Learning Servers. Dabei werden ausschließlich User-Mode-Kerne verwendet. Auf Systemen mit Hybridarchitektur werden ausschließlich Performance-Kerne (P-Cores) berücksichtigt.

Modell erstellen

- Über **Submit** wird die Trainingskonfiguration gespeichert. Das Training muss dann auf der Modellkarte gestartet werden.
- Über **Cancel** wird der Dialog ohne Änderungen geschlossen.

Modellkarte

Die Modellkarte zeigt den aktuellen Status sowie die wichtigsten Konfigurations- und Leistungsdaten eines KI-Modells.

Der Titel der Karte entspricht dem Modellnamen.

Modellaktionen

Im oberen Bereich stehen verschiedene Aktionen zur Verfügung:

- **Analyse öffnen:** Öffnet die Analyse- und Explainability-Ansicht des Modells.
- **Modell herunterladen:** Lädt das trainierte Modell herunter.
- **Modell bearbeiten:** Öffnet die Modellkonfiguration zur Bearbeitung der Metadaten
- **Modell löschen:** Entfernt das Modell aus dem Projekt.
- **Start Training:** Startet einen Trainingslauf für das Modell. Beachten Sie die Regelungen zur [Compute-Hours-Verrechnung](#) [► 24].

Modellinformationen

- **Status:** Zeigt den aktuellen Zustand des Modells an. Mögliche Zustände sind beispielsweise:
 - Submitted
 - Waiting
 - Running
 - Finished (inkl. Zeitstempel, wann das Modell diesen Zustand erreicht hat)
 - Error
- **Description:** Optionale Beschreibung des Modells
- **Dataset:** Zeigt den verwendeten Datensatz einschließlich der verwendeten Datensatzversion. Über das Suchsymbol kann der zugehörige Datensatz geöffnet werden.
- **Labels:** Zeigt die im Datensatz enthaltenen Klassen beziehungsweise Zielkategorien.
- **Target device:** Zeigt die konfigurierte Zielhardware einschließlich CPU-, GPU- und Execution-Provider-Konfiguration.
- **Target Software:** Zeigt die konfigurierte TwinCAT-Zielumgebung für die spätere Inferenz.
- **Target latency threshold:** Zeigt die konfigurierte maximale Ziellatenz für die Modellausführung.

Trainings- und Performanceinformationen

- **Performance:** Zeigt die erreichte Modellgüte als Prozentwert. Die dargestellte Kennzahl basiert auf den während der Validierung berechneten Testergebnissen.

- **Training Time** Gesamtdauer des Trainingsprozesses.
- **Estimated latency:** Geschätzte Laufzeit des Modells auf der konfigurierten Zielhardware.
- **Training History:** Über Loss History kann der Verlauf des Trainingsprozesses visualisiert werden. Dabei werden die Trainingsmetriken einzelner Epochen dargestellt.

5.3.3.1 Abrechnung von Compute Hours

Compute Hours stellen die verfügbaren Rechenressourcen für das Training von KI-Modellen dar.

Das verfügbare Compute-Kontingent hängt vom verwendeten Lizenztyp sowie von optional zugewiesenen Ressource Packs ab. Mit Beginn einer neuen Subscription-Periode wird das in der Lizenz enthaltene Compute-Kontingent erneut bereitgestellt. Nicht genutzte Compute Hours verfallen am Ende der jeweiligen Laufzeit. Zusätzliche Compute Hours können über Ressource Packs erweitert werden.

Wenn kein ausreichendes Compute-Kontingent mehr verfügbar ist, können keine weiteren Trainings gestartet werden. In diesem Fall muss entweder auf die nächste Subscription-Periode gewartet oder ein zusätzliches Ressource Pack zugewiesen werden.

Den aktuellen Wert ihres Compute-Kontingents finden Sie oben rechts in der Kopfzeile; angegeben in Stunden und Minuten.

Reservierung von Compute Hours

Beim Start eines Modelltrainings wird zunächst die voraussichtliche Trainingszeit auf Basis der Trainingskonfiguration und des verwendeten Datensatzes geschätzt.

Ist die geschätzte Trainingszeit kleiner oder gleich dem verfügbaren Compute-Kontingent, wird das Training gestartet. Dabei wird die geschätzte Compute-Zeit zunächst reserviert. Dieser Mechanismus entspricht der temporären Reservierung eines Budgets, beispielsweise bei einer Kreditkartenzahlung.

Nach Abschluss des Trainings wird ausschließlich die tatsächlich benötigte Trainingszeit abgerechnet. Nicht verwendete reservierte Compute Hours werden automatisch wieder freigegeben.

Unzureichendes Compute-Kontingent

Ist die geschätzte Trainingszeit größer als das aktuell verfügbare Compute-Kontingent, kann das Training nicht gestartet werden.

In diesem Fall muss zusätzliches Compute-Kontingent bereitgestellt werden.

Abrechnungszeitpunkt

Die Abrechnung eines Trainings erfolgt immer auf Basis des Compute-Kontingents, das zum Zeitpunkt des Trainingsstarts verfügbar war. Die reservierten Compute Hours bleiben während der gesamten Trainingsdauer an die ursprüngliche Subscription-Periode gebunden.

Beispiel:

Wird ein Training am letzten Tag einer Subscription-Periode gestartet und endet erst in der folgenden Subscription-Periode, wird die gesamte Trainingszeit dennoch dem ursprünglichen Subscription-Term belastet. Grund dafür ist, dass das Compute-Kontingent bereits beim Trainingsstart reserviert wurde.

5.3.3.2 Latency Thershold

Über den Parameter **Latency Threshold** kann die maximal zulässige Inferenzlatenz des KI-Modells definiert werden.

Dadurch kann das Modelltraining gezielt an die Anforderungen der späteren Maschinen- oder Prozesslaufzeit angepasst werden.

Der Benutzer definiert dabei die vorgesehene Zielhardware, die Zielsoftware sowie die maximal erlaubte Laufzeit für die Modellausführung. Der TwinCAT Machine Learning Creator berücksichtigt diese Randbedingungen während der Modellerstellung und optimiert das Modell entsprechend hinsichtlich Laufzeit und Modellkomplexität.

Dadurch können KI-Modelle erzeugt werden, die sowohl die gewünschte Modellgüte als auch die erforderlichen Zeitbedingungen erfüllen.

Beispiel:

In einer Qualitätskontrolle wird ein Produkt per Kamera erfasst und anschließend über eine Ausschleusstation aussortiert. Zwischen Bildaufnahme und Erreichen der Ausschleusposition stehen lediglich 10 ms zur Verfügung. In diesem Fall muss die Inferenz einschließlich Vorverarbeitung und Nachverarbeitung innerhalb dieses Zeitfensters abgeschlossen sein.

Der konfigurierte Latency Threshold ermöglicht die Anpassung des KI-Modells an diese Prozessanforderung.

Balance zwischen Modellgüte und Latenz

Der konfigurierte Latency Threshold beeinflusst die Auswahl und Optimierung der Modellarchitektur.

Wird kein Latenzschwellwert angegeben, erfolgt keine Einschränkung hinsichtlich der späteren Inferenzlaufzeit. Dadurch kann der Optimierungsprozess komplexere Modellarchitekturen auswählen, was in der Regel zu besseren Modellmetriken, wie Accuracy, Precision oder Recall führt.

Im Engineering-Prozess wird daher empfohlen, zunächst ohne Latency Threshold zu trainieren. Dadurch kann der Datensatz iterativ verbessert und bewertet werden, bis ein fachlich überzeugendes Modell erreicht wird.

Erst anschließend sollte der Latency Threshold schrittweise reduziert werden. Auf diese Weise kann überprüft werden, ob weiterhin eine ausreichende Modellgüte bei geringerer Modellkomplexität und reduzierter Inferenzlatenz erreicht wird.

Dieses Vorgehen ermöglicht eine gezielte Balance zwischen Modellperformance und Echtzeitfähigkeit der späteren Anwendung.

5.3.4 Modell validieren

Explainability-Ansicht

Die Explainability-Ansicht dient zur Analyse und Validierung eines trainierten KI-Modells.

Die dargestellten Ergebnisse basieren auf dem Testdatensplit des verwendeten Datensatzes. Standardmäßig werden 20 % des Datensatzes automatisch als Testdaten reserviert. Die verbleibenden 80 % werden für das Modelltraining verwendet.

Die Testdaten werden während des Trainings nicht verwendet und dienen ausschließlich zur Bewertung der Generalisierungsfähigkeit des Modells.

Modellstatistiken

Im oberen linken Bereich werden die wichtigsten Bewertungsmetriken des Modells dargestellt.

- **Test Accuracy:** Beschreibt den Anteil aller korrekt klassifizierten Testdaten. Die Accuracy gibt einen allgemeinen Überblick über die Gesamtleistung des Modells.
- **Test Precision:** Beschreibt die Genauigkeit der positiven Vorhersagen. Die Precision bewertet, wie viele als positiv klassifizierte Samples tatsächlich korrekt sind.
- **Test Recall:** Beschreibt die Fähigkeit des Modells, relevante Samples korrekt zu erkennen. Ein hoher Recall bedeutet, dass nur wenige relevante Samples übersehen werden.
- **Test F1:** Der F1-Score kombiniert Precision und Recall zu einer gemeinsamen Kennzahl. Die Metrik ist insbesondere bei unausgeglichenen Datensätzen hilfreich.

Confidence

Der Bereich **Confidence** zeigt statistische Kennwerte der Modellkonfidenzen:

- Mittelwert (Mean)
- Minimalwert (Min)

- Maximalwert (Max)

Die Konfidenz beschreibt die Sicherheit des Modells bei einer Vorhersage.

Die Confidence Distribution visualisiert die Verteilung aller Vorhersagekonfidenzen als Histogramm. Jeder Balken repräsentiert einen Konfidenzbereich (Bin) und zeigt die Anzahl der Vorhersagen innerhalb dieses Bereichs. Dadurch lässt sich beurteilen, ob das Modell überwiegend sichere oder unsichere Entscheidungen trifft.

Interaktive Filterung

Durch Anklicken eines Histogramm-Bereichs kann auf einen bestimmten Konfidenzbereich gefiltert werden. Die gefilterten Ergebnisse werden anschließend im Bereich **Prediction Behavior** angezeigt. Dort erscheinen nur noch Samples innerhalb der ausgewählten Confidence-Range. Der aktive Filter wird im oberen Bereich von Prediction Behavior angezeigt und kann auch dort gelöscht werden.

Confusion Matrix

Die Confusion Matrix zeigt die Verteilung von korrekten und fehlerhaften Vorhersagen pro Klasse.

- Zeilen entsprechen den tatsächlichen Klassen (True Labels, wie vom Menschen gelabelt).
- Spalten entsprechen den vorhergesagten Klassen (Predicted Labels).

Die Diagonale der Matrix enthält korrekt klassifizierte Samples. Werte außerhalb der Diagonale stellen Fehlklassifikationen dar.

Die Größe der Matrix hängt von der Anzahl der Klassen im Datensatz ab.

Interaktive Filterung

Durch Anklicken einzelner Matrixelemente kann auf bestimmte Vorhersagekombinationen gefiltert werden. Dadurch lassen sich gezielt Fehlklassifikationen oder korrekt klassifizierte Samples analysieren. Die gefilterten Ergebnisse werden im Bereich **Prediction Behavior** dargestellt.

Prediction Behavior

Der Bereich Prediction Behavior zeigt alle Samples aus dem Testdatensatz einschließlich Vorhersageergebnis, Modellkonfidenz und attention map.

Korrekte und fehlerhafte Vorhersagen können separat ein- oder ausgeblendet werden. Zusätzlich können die Ergebnisse nach der Konfidenz sortiert werden.

Attention Maps

Für jedes angezeigte Sample kann optional eine Attention Map dargestellt werden. Die Visualisierung hebt Bildbereiche hervor, die besonders stark zur Entscheidung des KI-Modells beigetragen haben. Dadurch wird nachvollziehbar, welche Merkmale für die Klassifikation relevant waren. Über den Schalter in jedem Bild, kann die Attention Map abgeschaltet werden.

5.3.5 Modell-Download

Über den Modell-Download können trainierte KI-Modelle einschließlich Integrationsartefakten exportiert werden.

Zielsoftware auswählen

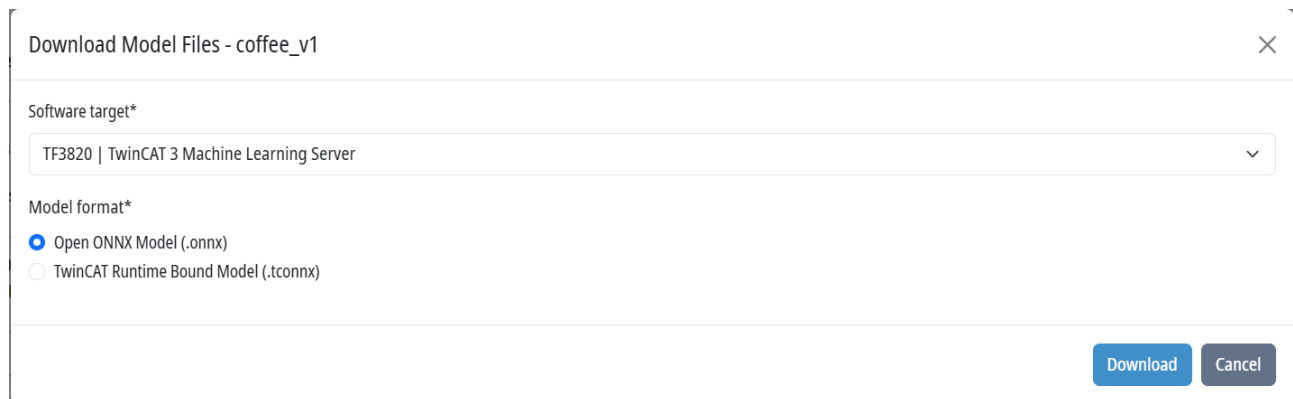
Im Feld **Software target** wird die Zielsoftware für die Generierung des PLC-Codes ausgewählt.

Standardmäßig wird die Zielsoftware verwendet, die bereits während der Modellerstellung definiert wurde.

Die Zielsoftware kann beim Download erneut geändert werden. In diesem Fall wird der PLC-Code für die neu ausgewählte Laufzeitumgebung generiert. Dabei ist zu beachten, dass die während des Trainings berechnete Latenzschätzung nur für die ursprünglich konfigurierte Zielsoftware gültig ist.

Modellformat auswählen

Je nach verwendetem Lizenztyp stehen unterschiedliche Modellformate zur Verfügung.



Download Model Files - coffee_v1

Software target*

TF3820 | TwinCAT 3 Machine Learning Server

Model format*

☒ Open ONNX Model (.onnx)

☐ TwinCAT Runtime Bound Model (.tconnx)

Download Cancel

Open ONNX Model (.onnx)

Das Modell wird als standardisierte ONNX-Datei exportiert.

ONNX ist ein offenes und frameworkunabhängiges Format und kann in beliebigen kompatiblen Laufzeitumgebungen oder KI-Frameworks verwendet werden.

TwinCAT Runtime Bound Model (.tconnx)

Das Modell wird als TwinCAT-gebundenes tcONNX-Modell exportiert.

Dieses Format ist ausschließlich für die Verwendung innerhalb von TwinCAT 3 vorgesehen.

Inhalt des Download-Pakets

Der Download erfolgt als ZIP-Archiv mit mehreren Dateien zur Integration und Dokumentation des KI-Modells.

KI-Modell

Enthält das trainierte Modell entweder als ONNX oder tcONNX abhängig vom gewählten Exportformat.

model_card.json

- Beschreibt das erzeugte KI-Modell einschließlich:
- verwendeter Datensatzsplits
- Eingabe- und Ausgabeformaten
- Klassenlabels
- Zeitstempel
- Trainings- und Laufzeitinformationen
- Latenzschätzungen

PLCopen-XML

Enthält vollständigen PLC-Code zur Integration in TwinCAT 3.

Der generierte Code umfasst:

- Bildaufnahme
- Bildvorverarbeitung
- KI-Inferenz
- Nachverarbeitung der Ergebnisse

Die Datei kann direkt in TwinCAT 3 importiert werden.

Modelname.json

Je nach ausgewähltem Software-Target wird auch eine spezifische json generiert, welche von der Inferenz-Software benötigt wird.

readme.md

Enthält eine detaillierte Beschreibung zur Integration der generierten PLCopen-XML in ein TwinCAT-3-Projekt. Zusätzlich werden erforderliche Konfigurationsschritte und projektspezifische Anpassungen beschrieben.

6 Support und Service

Beckhoff und seine weltweiten Partnerfirmen bieten einen umfassenden Support und Service, der eine schnelle und kompetente Unterstützung bei allen Fragen zu Beckhoff Produkten und Systemlösungen zur Verfügung stellt.

Downloadfinder

Unser Downloadfinder beinhaltet alle Dateien, die wir Ihnen zum Herunterladen anbieten. Sie finden dort Applikationsberichte, technische Dokumentationen, technische Zeichnungen, Konfigurationsdateien und vieles mehr.

Die Downloads sind in verschiedenen Formaten erhältlich.

Beckhoff Niederlassungen und Vertretungen

Wenden Sie sich bitte an Ihre Beckhoff Niederlassung oder Ihre Vertretung für den lokalen Support und Service zu Beckhoff Produkten!

Die Adressen der weltweiten Beckhoff Niederlassungen und Vertretungen entnehmen Sie bitte unserer Internetseite: www.beckhoff.com

Dort finden Sie auch weitere Dokumentationen zu Beckhoff Komponenten.

Beckhoff Support

Der Support bietet Ihnen einen umfangreichen technischen Support, der Sie nicht nur bei dem Einsatz einzelner Beckhoff Produkte, sondern auch bei weiteren umfassenden Dienstleistungen unterstützt:

- Support
- Planung, Programmierung und Inbetriebnahme komplexer Automatisierungssysteme
- umfangreiches Schulungsprogramm für Beckhoff Systemkomponenten

Hotline: +49 5246 963-157

E-Mail: support@beckhoff.com

Beckhoff Service

Das Beckhoff Service-Center unterstützt Sie rund um den After-Sales-Service:

- Vor-Ort-Service
- Reparaturservice
- Ersatzteilservice
- Hotline-Service

Hotline: +49 5246 963-460

E-Mail: service@beckhoff.com

Beckhoff Unternehmenszentrale

Beckhoff Automation GmbH & Co. KG

Hülshorstweg 20
33415 Verl
Deutschland

Telefon: +49 5246 963-0

E-Mail: info@beckhoff.com

Internet: www.beckhoff.com

Trademark statements

Beckhoff®, ATRO®, EtherCAT®, EtherCAT G®, EtherCAT G10®, EtherCAT P®, MX-System®, Safety over EtherCAT®, TC/BSD®, TwinCAT®, TwinCAT/BSD®, TwinSAFE®, XFC®, XPlanar® and XTS® are registered and licensed trademarks of Beckhoff Automation GmbH.

Third-party trademark statements

Chrome, Chromium and Google are trademarks of Google LLC.

Excel, IntelliSense, Microsoft, Microsoft Azure, Microsoft Edge, PowerShell, Visual Studio, Windows and Xbox are trademarks of the Microsoft group of companies.

Mozilla and Firefox are trademarks of the Mozilla Foundation in the U.S. and other countries.

NVIDIA, NVIDIA RTX and CUDA are trademarks of NVIDIA Corporation.

Mehr Informationen:
www.beckhoff.com/te3850

Beckhoff Automation GmbH & Co. KG
Hülshorstweg 20
33415 Verl
Deutschland
Telefon: +49 5246 9630
info@beckhoff.com
www.beckhoff.com

