

IDA Classification

Lernen statt Regeln schreiben

IDA Classification ordnet Ihre Dokumente automatisch ohne komplexe Regelwerke und ohne Programmierkenntnisse. Durch fortschrittliches Few-Shot-Learning genügen wenige Beispieldokumente, um neue Klassen zu trainieren. Die Kombination aus textueller und visueller Analyse sorgt für außergewöhnliche Genauigkeit selbst bei minimalen Unterschieden. Das Ergebnis: Drastisch verkürzte Einrichtungszeiten, minimaler Wartungsaufwand und verbesserte Straight-Through-Processing-Raten – auch bei über 200 Dokumentklassen.

FEATURES

Lernen mit wenigen Trainingsdaten

IDA Classification nutzt fortschrittliches Few-Shot-Learning, das sowohl visuelle als auch textuelle Merkmale analysiert. Neue Dokumentklassen sind in Stunden statt Wochen einsatzbereit.

Wenn sich Ihre Dokumententypen ändern oder erweitern, passt sich das System ohne aufwändige Neukonfiguration an. Das minimiert nicht nur die Zeit bis zur Wertschöpfung, sondern reduziert auch den Wartungsaufwand erheblich.

No-Code-Training

Benutzer:innen ohne Programmierkenntnisse können eigenständig Klassifikationsmodelle erstellen, trainieren und anpassen. Die browserbasierte grafische Oberfläche macht Machine Learning einfach zugänglich.

Einfache Bereitstellung und Integration

IDA wird entweder vor Ort (on-premises) oder in der Cloud als Java-Anwendung oder als Containerisierung mit Docker bereitgestellt. Die gRPC-API (alternativ: REST-API) ermöglicht eine schnelle Integration.

KONFIGURATIONEN

Eingabe: JSON (proprietäres „PAI-File“, z. B. aus vorheriger Recognition)

Ausgabe: PDF, PDF/A, JSON

VORAUSSETZUNGEN

Für 64-Bit-Systeme

Linux: Ubuntu 18.04-25.10, Debian 11-13, CentOS 8-10, Red Hat 8.x-10.x, LEAP 15.4-15.6, 16.0; SLES 15 SP 4-7

Windows: 10, 11

Windows Server: 2016, 2019, 2022, 2025

Docker

Mind. **12 GB Festplattenspeicher**

+ unterschiedlicher Speicherbedarf der selbst trainierten Modelle je nach Einstellungen (bis zu ca. 1 GB für vollständiges Training)

Mind. **16 GB Arbeitsspeicher (RAM)**

Das volle Potenzial unübertroffener OCR-Qualität

IDA Classification basiert auf **IDA Recognition**, der OCR-Engine für hervorragende Ergebnisse in den schwierigsten Szenarien. Selbst bei verzerrten Scans, schlechter Bildqualität und schwer lesbarer Handschrift liefert IDA Recognition die hochwertige Textbasis, die für verlässliche Klassifikation entscheidend ist. Warum das wichtig ist: Klassifikationsqualität steht und fällt mit der Qualität der Eingabedaten. IDA Recognition erfasst Maschinenschrift, Handschrift, Kontrollkästchen, Tabellen und historische Schriften als perfekte Grundlage für nachfolgende Klassifikationsprozesse.

Automatische Dokumententrennung

Scanstapel mit über 100 aufeinanderfolgenden Seiten automatisch in Einzeldokumente trennen? Mit IDA Classification ist es möglich, ein neuronales Netz zu trainieren, das automatisch die Grenzen von Dokumenten in großen Dateien erkennt. Die manuelle Trennung während des Scannens entfällt damit komplett, was Prozesse beschleunigt und Fehlerquellen eliminiert.

ZWEI ANSÄTZE FÜR VERSCHIEDENE SZENARIEN

Indem IDA mit einer kleinen Menge Beispieldokumente in Form von Verzeichnissen gefüttert wird, erkennt das System automatisch signifikante Merkmale und lernt weiter. IDA bietet zwei **Klassifizierungsmethoden**, die sowohl für einseitige als auch für mehrseitige Dokumente anwendbar sind.

Dokumentenklassifikation: Den Kontext im Ganzen verstehen

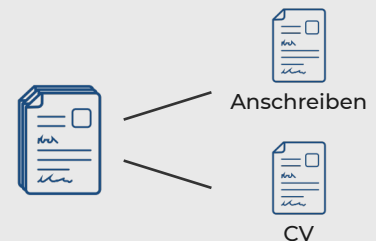
- Analysiert den gesamten Text eines Dokuments für holistische Klassifizierung.
- Ideal für Szenarien wie Posteingang automatisch sortieren



Bewerbung

Seitenklassifikation: Mehrseitige Dokumente intelligent aufteilen

- Analysiert einzelne Seiten unter Berücksichtigung von Vorgänger- und Folgeseiten.
- Ideal für Szenarien wie Bewerbungsunterlagen in Anschreiben, Lebenslauf und Zeugnisse sortieren



Beide Klassifikationsmethoden basieren entweder auf einem trainierbaren neuronalen Netz oder einem Wörterbuch (Bag of Words-Modell). Auf Grundlage der sich ergebenden Klassen können Dokumente an verschiedene nachgelagerte Prozesse weitergeleitet werden, z. B. an eine Datenextraktion.

DOKUMENTENTRENNUNG: VON STAPELN ZU EINZELDOKUMENTEN

Der Geschäftsalltag: Gescannte Dokumentenstapel mit 100, 200 oder mehr aufeinanderfolgenden Seiten. Manuelles Trennen ist zeitraubend und fehleranfällig. Mit IDA können Sie ein neuronales Netz trainieren, das Dokumentengrenzen automatisch erkennt. Das System lernt, wo ein Dokument endet und das nächste beginnt.

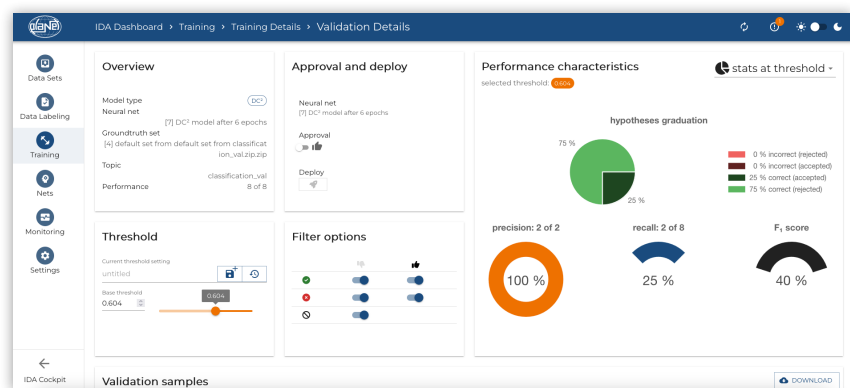


- Keine manuelle Trennung während des Scannens mehr erforderlich
- Wiederverwendung bestehender Seitenklassifikations-Modelle möglich
- Alternative: Regelbasierte Trennung nach fester Seitenzahl für gleichförmige Dokumente
- Für strukturierte Formulare genügt ein leeres Muster pro Klasse

Einschränkung: Aktuell für Dokumente in korrekter Reihenfolge konzipiert.

MODELLTRAINING

IDA bietet eine benutzerfreundliche grafische Oberfläche, die ein einfaches Modelltraining ohne Programmierkenntnisse ermöglicht.



Neuronales Netzwerk: Für die meisten Szenarien die beste Wahl

Für alle komplexen Klassifikationsszenarien empfiehlt sich das Training eines neuronalen Netzes. Es berücksichtigt sowohl visuelle als auch textuelle Merkmale, um ein Modell zu erstellen. Während des Trainings lernen Attention-Mechanismen, sich auf die relevantesten Merkmale zu konzentrieren.

Trainingsanforderungen:

- Minimum: 20 Dokumente pro Klasse + 10 Validierungsdokumente
- Empfohlen: 100+ Dokumente pro Klasse für optimale Genauigkeit
- Sonderfall strukturierte Formulare: 1 leeres Formular pro Klasse genügt

Darüber hinaus steht das **vortrainierte Open-Source-Modell LayoutLM** zur Verfügung, um die Klassifikationsergebnisse für Dokumente mit ähnlichem Layout, aber unterschiedlichen Texten

zu verbessern. LayoutLM ist ein Large Language Model, das speziell für Dokumentenlayouts und Kontextverständnis entwickelt wurde. Für die Nutzung dieses Modells ist GPU-Hardware erforderlich.

Bag of Words (Wörterbuch): Effizient für einfache Fälle

Diese Option nutzt einen **Keyword-Spotting-Ansatz** in Kombination mit der patentierten **PerceptionMatrix**. Das Bag-of-Words-Modell berücksichtigt in erster Linie textuelle Merkmale und ist dadurch für weniger komplexe Klassifikationsaufgaben geeignet. IDA führt eine **Suche innerhalb jedes Dokuments durch, die auf einer anpassbaren Liste von Wörtern** (inkl. Wortgruppen, Sätzen) basiert. Zusätzlich ist es möglich, innerhalb der PerceptionMatrix zu suchen, die alle möglichen Transkriptionen eines gegebenen Textes ohne Informationsverlust speichert.



Erleben Sie IDA Classification selbst und kontaktieren Sie uns für eine individuelle Demo.

[Jetzt Demo buchen](#)