

# Vertrauenswürdige KI für die Prozessindustrie

## Anwendung in datenarmen Szenarien

#Automatisierung #J-Kurve #Machine Learning #Datenqualität  
#Vertrauenswürdigkeit

*Die Prozessindustrie steht vor einem tiefgreifenden Wandel, da Künstliche Intelligenz (KI) über klassische Automatisierung hinaus datenbasierte und adaptive Entscheidungen ermöglicht. Die Einführung von KI als Basistechnologie folgt der Logik einer J-Kurve: Anfangs entstehen hohe Investitionen, langfristig jedoch deutliche Produktivitätsgewinne. An einem Anwendungsfall zeigen wir, dass KI-Algorithmen auch in datenarmen Szenarien robuste Vorhersagen liefern. Vertrauenswürdige KI ist dabei die zentrale Voraussetzung, um Effizienz, Sicherheit und regulatorische Anforderungen zu vereinen.*

*The process industry is undergoing a profound transformation as artificial intelligence (AI) extends classical automation with data-driven and adaptive decision-making. The introduction of AI as a general-purpose technology follows the logic of a J-curve: initial investments are high, but long-term productivity gains are substantial. Using a case study, we demonstrate that AI algorithms can deliver robust predictions even in data-scarce scenarios. Trustworthy AI is a central prerequisite to combine efficiency, safety, and compliance with regulatory requirements.*

Penelope Mück,  
Justus Viga,  
KROHNE Messtechnik;  
Alexander Löser,  
Berliner Hochschule  
für Technik;  
Dagmar Dirzus,  
KROHNE Messtechnik

► PEER-REVIEW:  
28.08.2025

### 1. Einleitung

Die Prozessindustrie steht vor einem tiefgreifenden Wandel: Digitale Technologien, steigende Komplexität und Effizienzdruck lassen das Thema Autonomie mehr Bedeutung gewinnen. Dabei geht es nicht nur um klassische Automatisierung, sondern auch um KI-Systeme, die datenbasiert, adaptiv und vorausschauend selbstständig agieren. Ein Beispiel liefert die Logistik: Amazon steuert mit über einer Million Robotern und 1,56 Millionen Mitarbeitenden komplexe Abläufe bereits KI-gestützt und weitgehend autonom [2]. Diese Entwicklung erreicht nun auch die Prozessindustrie, mit vergleichbarem Potenzial, aber höheren technischen und sicherheitsrelevanten Anforderungen.

Die zentrale Frage lautet: Wie kann ein höherer Autonomiegrad erreicht werden, ohne Sicherheit, Transparenz und Kontrollierbarkeit zu gefährden?

### 2. Vier Bausteine hin zur Autonomie vs. Aufwände entlang der J-Kurve

**Vier Bausteine:** Abbildung 1 zeigt die vier notwendigen Bausteine für Autonomie und Intelligenz in der Prozessindustrie angelehnt an [3], [4]. Baustein 1, die Grundvoraussetzung, ist die vollständige Digitalisierung und Vernetzung aller Komponenten. Baustein 2 ist „Beyond Diagnostics“, d. h. Plug-&-Produce, Advanced Analytics und Predictive Maintenance: Neue Komponenten, beispielsweise Messgeräte, können ohne aufwendige manuelle Konfiguration in bestehende Anlagen integriert werden. Im Betrieb können durch kontinuierliche Messung und Auswertung von Daten Probleme innerhalb einer Anlage frühzeitig erkannt und so proaktiv behandelt werden. Unter anderem können seltene Ereignisse erkannt und ein Herunterfahren der Anlage vermieden werden. Baustein 3 ist der „Process

## 4 Bausteine der Automatisierung

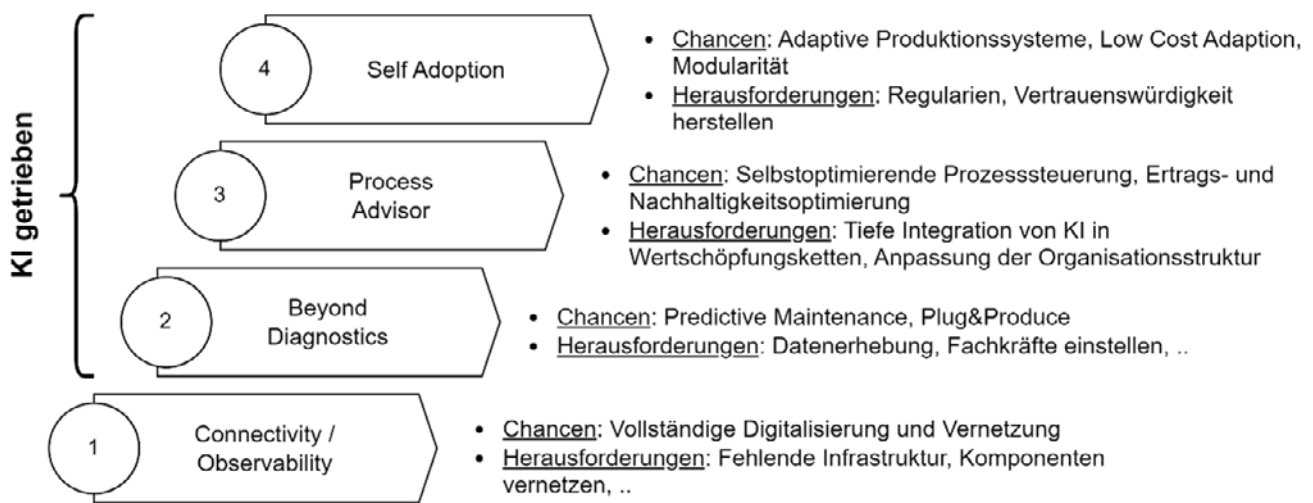


Abbildung 1: Die vier Bausteine der Automatisierung. Baustein 1 schafft die Datengrundlage, während Bausteine 2-4 zunehmend KI-getrieben sind und spezifische Chancen und Herausforderungen für den Weg zur autonomen Prozessführung aufzeigen.

Advisor“, hierunter fällt die selbstoptimierende Prozesssteuerung, insbesondere um die Verwendung von Ressourcen zu optimieren und nachhaltig zu wirtschaften. Baustein 4 nennen wir „Self Adoption“. Komponenten in verteilten modularen Produktionssystemen können sich mit minimalen Kosten in Echtzeit austauschen, um sich an veränderte Bedingungen anzupassen.

**J-Kurve: KI-Transformation senkt zunächst die Produktivität bevor Skalierungseffekte einsetzen.** Die Autoren von [5] bezeichnen KI als eine Basistechnologie (General Purpose Technology, GPT), die zahlreiche Industrien grundlegend verändern kann, ähnlich wie einst die Elektrizität oder das Internet. Die sogenannte J-Kurve beschreibt die Einführung einer Basistechnologie: Sie verursacht zunächst hohe Investitionen, ohne dass kurzfristig signifikante Produktivitätsgewinne sichtbar werden. Beispiele für diese Kosten entstehen auf Ebene jedes Bausteins: Zunächst müssen Daten erhoben, Komponenten vernetzt, neue Infrastrukturen gebaut und entsprechende Fachkräfte eingestellt werden. KI muss tief in Prozesse und

Wertschöpfungsketten eingebunden werden. Zusätzliche Kosten entstehen durch neue Organisationsstrukturen. Die Einhaltung komplexer Regularien erzeugt zusätzliche Aufwände.

Viele Unternehmen scheuen sich daher, umfassende KI-Projekte zu starten. Der initiale Rückgang im ROI kann jedoch ein normaler Bestandteil des Transformationsprozesses sein. Erst nach einer gewissen Zeit, wenn Systeme integriert, Daten aufbereitet und Kompetenzen aufgebaut sind, steigen die Erträge überproportional an. Die Studie von [5] zeigt, dass frühe Investitionen in KI langfristig erheblichen wirtschaftlichen Mehrwert erzeugen können. Vorausgesetzt, Unternehmen sind bereit, organisatorische Anpassungen vorzunehmen und die „Talsohle“ der J-Kurve strategisch zu durchqueren.

### 3. KI in der Praxis

#### 3.1 Trainingsdaten für KI in der Prozessindustrie

Industrielle Automatisierung stellt hohe Anforderungen an Künstliche Intelligenz, da Prozesse komplex, kontinuier-

lich und sicherheitskritisch sind. Die folgenden Abschnitte geben einen Überblick über zentrale Konzepte und Verfahren KI-gestützter Automatisierung.

**Überwachte vs. unüberwachte Lernverfahren:** Maschinelles Lernen (ML) lässt sich grob in zwei Hauptkategorien einteilen: überwachtes und unüberwachtes Lernen. Beim überwachten Lernen werden Algorithmen mit Daten trainiert, die bereits ein Ergebnis (Label) enthalten, etwa zur Klassifikation (z. B. automatische Erkennung fehlerhafter Produkte) oder Regression (z. B. Vorhersage des Energieverbrauchs). Unüberwachtes Lernen dagegen nutzt Daten ohne vorgegebene Ergebnisse. Ziel ist es, Muster oder Strukturen zu entdecken, etwa durch Clustering, das ähnliche Zustände oder Verhaltensweisen gruppiert. Das ist z. B. nützlich zur besseren Prozesssteuerung oder dem frühzeitigen Erkennen ungewöhnlicher Anlagenzustände.

**Trainingsdaten und Trainingsprozess:** Der überwiegende Teil des Aufwands für eine Machine-Learning-Lösung entfällt auf die sorgfältige Datenkuratierung.

rung und den Trainingsprozess, da hier die Qualität und Repräsentativität der Daten entscheidend ist. Anschließend wird das trainierte Modell mithilfe eines separaten Testdatensatzes bewertet, um sicherzustellen, dass es auf unbekannte Daten generalisiert. Dieser iterative Prozess „Training, Evaluierung, Anpassung“ ist essenziell, damit Modelle in realen Anwendungen robust und verlässlich funktionieren.

**Verwertbare Trainingsdaten sind einer der Kostentreiber der J-Kurve:** Prozesse in der Industrie, die Trainingsdaten erzeugen sollten, sind oft bisher nicht darauf ausgelegt. Das Kernproblem bisher war, dass es keinen Zwang gab, z. B. über Regulatorik, in Trainingsdaten zu investieren. Es sind daher wenig Daten verfügbar mit vielen fehlenden Werten und ohne Label. Es bestehen oft zu wenig oder ungenaue Kenntnisse über den Prozess, insbesondere für

den Einsatz von Machine Learning. Das liegt zum Teil an schwierigen Prozessbedingungen, aber auch an fehlenden Messgeräten o. Ä., die Aufschluss über Zustände an Teilkomponenten einer Anlage geben könnten. Wenn dann Daten aufgenommen und gesichtet werden stellt sich ein mangelndes Grounding heraus: Die Prozessbedingungen sind anders als angenommen sind und es treten Anomalien auf, wie z. B. Temperaturspitzen oder sich bildende Feststoffe. Hinzu kommt die Individualität verschiedener Prozesse und Anlagen, die es erschwert, ein allgemein gültiges KI-Modell zu erstellen. Meist müssen daher für jede neue Anlage neue Daten aufgenommen werden und ein neues Modell trainiert werden. Auch beim Ressourcenbedarf für das Ausführen von KI, die sogenannte Inferenz, gibt es Einschränkungen. KI soll häufig on-the-edge ausgeführt werden, sei es auf Microcontrollern in Messge-

räten oder einfachen Computern mit GPU-Unterstützung.

### 3.2 Lösungsansätze

**Abhilfe können vortrainierte Foundation Models und In-Context Learning schaffen.** Auf generellen Daten vortrainierte Modelle lassen sich auch ohne erneutes Training anpassen, z. B. für textintensive Prozesse in Service, R&D, Support oder Backoffice. Dabei kommt sogenanntes *In-Context Learning* (oder *Test Time Compute*) zum Einsatz: Das Modell nutzt ein langes Eingabefenster als eine Art Arbeitsspeicher. Maschinelle Verfahren, wie Reasoning oder Agenten, agieren mit Tools wie Websuchen, Suchen in Archiven oder anderen Diensten, um diesen Speicher mit Wissensrepräsentationen für eine Aufgabe zu füllen. Anschließend wendet das Modell logische Verfahren wie das Schritt-für-Schritt-Denken oder die Verifikation darauf an, um Aufgaben zu lösen oder in Teilaufgaben zu zerlegen. Diese Methode funktioniert auch für Modelle hinter der Firewall, benötigt kein Training und kann relativ einfach in Prompts programmiert werden. So sinkt die Einstiegshürde für viele Automatisierungsszenarien deutlich. Für die Prozessindustrie sind nicht nur vortrainierte Modelle im Bereich Sprache relevant, sondern auch im Bereich Bild in der Qualitätskontrolle oder auf strukturierten Daten im Bereich Zeitreihendaten.

**Feedback-Loops und kontinuierliches Lernen.** Der Report [6] identifiziert einen zentralen Grund, warum viele GenAI-Projekte scheitern: Das Fehlen von Lernen und Anpassung im Unternehmenskontext, oft genannt als „Learning Gap“. Dabei geht es nicht nur um technische Modellqualität, sondern vielmehr um die Unfähigkeit der eingesetzten Werkzeuge, sich in bestehende Unternehmensabläufe einzufügen und

**J-Kurve**

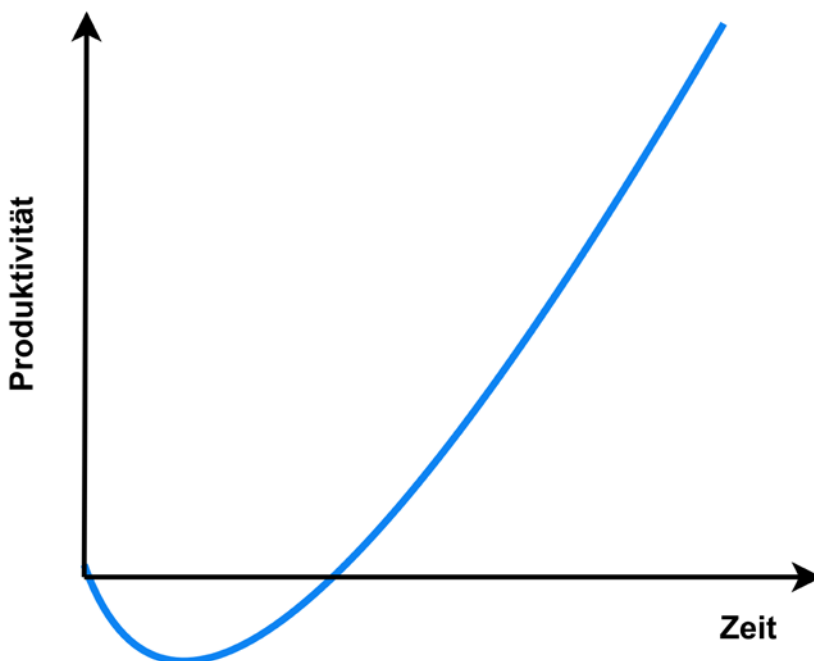


Abbildung 2: Die J-Kurve [5] beschreibt den anfänglichen Produktivitätsrückgang bei der Einführung neuer Technologien (z. B. KI), gefolgt von einem langfristigen Anstieg, sobald Investitionen in Prozesse, Schulungen und Systeme wirksam werden.

daraus im Sinne von Workflow-Adaptation und kontinuierlichem Lernen Lehren zu ziehen. Hier spielen Feedback-Loops eine wichtige Rolle: Systeme, die mit Nutzer- oder Prozessfeedback versorgt werden, können damit ihre Vorhersagen verbessern oder sich an veränderte Rahmenbedingungen anpassen. Solche Schleifen sind insbesondere für automatisierte Systeme mit Echtzeitkomponenten oder adaptivem Verhalten relevant. Diese Feedbackdaten sind exklusiv und helfen, die Kosten in der J-Kurve zu vermindern. Das Unternehmen kann die Qualität der Trainingsdaten kontrollieren, da die datenerzeugenden Prozesse im Unternehmen stattfinden. Aktuell werden Feedbackdaten leider oftmals an kommerzielle Tools wie etwa ChatGPT weitergegeben. Das bedeutet, dass fremde Modelle mit unseren Daten lernen können und wir nicht mehr in der Lage sind die vier Bausteine unabhängig von diesen Anbietern umzusetzen.

### Synthetische Daten auf Basis chemischer oder physikalischer Prozesse.

Ein weiteres wichtiges Element in der ML-Methodik ist die Nutzung synthetischer Daten. In Situationen, in denen reale Daten schwer zu erfassen, teuer oder aus Datenschutzgründen nicht verfügbar sind, bieten künstlich generierte Daten eine wertvolle Alternative. Diese können entweder aus bestehenden Daten simuliert oder mit Hilfe generativer KI erzeugt und im Test-Time-Compute-Verfahren verifiziert werden. Synthetische Daten ermöglichen es, Trainingsdatensätze zu erweitern, seltene Szenarien gezielt zu adressieren und die Robustheit von Modellen unter kontrollierten Bedingungen zu testen. Modelle wie TabPFN [7] und Phi-4 [8] zeigen, dass rein synthetisch generierte Datensätze zu starker Generalisierung führen können, durch kontrollierte Variation und breite Abdeckung möglicher Szenarien. In der Physik erlauben syn-

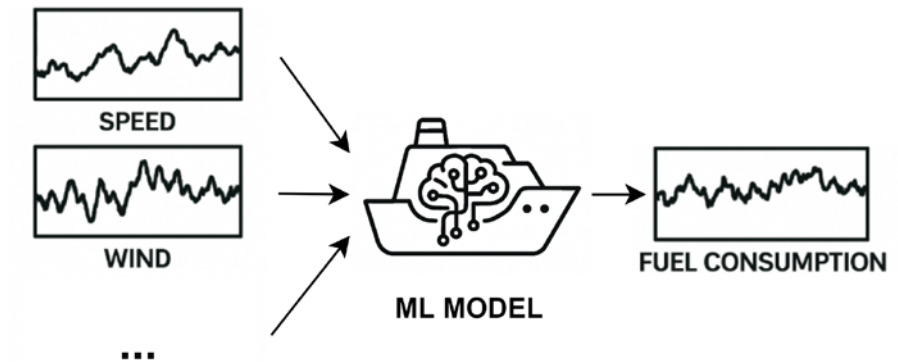


Abbildung 3: Vorhersage des Treibstoffverbrauchs von Schiffen als Hilfsgröße zur Prozessoptimierung.

thetische Daten präzise Simulationen, wo Experimente nicht machbar oder zu teuer sind.

### 3.3 Beispiel: Treibstoffverbrauch auf Schiffen

Als Beispiel für den Einsatz von traditioneller KI und GenAI auf strukturierten Daten und Zeitreihen in Anlagen verwenden wir die Optimierung des operativen Betriebs von Hochseeschiffen in Hinblick auf Treibstoffverbrauch und Treibhausgasemissionen. Solch ein Schiff kann als schwimmende Anlage gesehen werden, in der in Rohrsystemen Treibstoff (meistens Diesel oder Schweröl) zu verschiedenen Verbrauchern gepumpt und dort in Energie umgewandelt wird. Durch geschickte Optimierung von operativen Prozessparametern, wie z. B. Fahr-Geschwindigkeit des Schiffs, ist es möglich, den Treibstoffverbrauch und Emissionsausstoß zu minimieren, ohne den Betrieb des Schiffs einzuschränken.

**Das ist vergleichbar mit der Ausbeuteoptimierung in einer Anlage der Prozessindustrie.** Es handelt sich um Zeitreihendaten aus verschiedenen Sensoren und Modellen. Da jedes Schiff individuell gebaut ist, erfordert es eine eigene Betrachtung. Labels zu guten oder schlechten Zuständen fehlen und müssen indirekt über Treibstoffverbrauch und Kontext erschlossen werden. Die KI muss on-the-edge einsetzbar sein, da auf See keine ständige Internetver-

bindung besteht. Ziel ist ein effizienter, wirtschaftlicher Betrieb sowie die Erkennung anomaler Zustände.

Um die internen Parameter des Prozesses zu messen, wurden von [9] KROHNE OPTIMASS Coriolis Durchflussmessgeräte in die Treibstoffleitungen vor jedem Verbraucher eingebaut. Die gewonnenen Messwerte des Treibstoffverbrauchs wurden mit den GPS-Positionsdaten der Schiffe und Umgebungsdaten (z. B. Windstärke, Wellengang und Wassertiefe) aus öffentlich verfügbaren Quellen zusammengeführt.

### Selbstüberwachtes Lernen ohne Labels.

Da es in unserem Beispiel keine Labels gibt, mit denen wir überwachtes Lernen verwenden können, greifen wir auf einen selbstüberwachten Ansatz zurück. Statt direkt optimale Prozessbereiche oder Anomalien vorherzusagen, trainieren wir das KI-Modell zunächst darauf, eine Hilfsgröße, den Treibstoffverbrauch, vorherzusagen. Dieser ist in unseren Daten vorhanden und lässt sich so durch gewöhnliches überwachtes Training erlernen. Anschließend lässt sich in einem zweiten Schritt eine Optimierung auf dem gelernten Modell mit konventionellen Optimierungstechniken durchführen, über die sich die optimalen Prozessparameter für minimalem Treibstoffverbrauch ermitteln lassen, z. B. mittels Gradientenabstieg oder Monte-Carlo-Simulation.



Aus dem selbstüberwachten Ansatz lassen sich konkrete Aufgaben ableiten.

**Aufgabe „Tabulare Regression“:** Eine klassische tabulare Regression, bei der das Modell die Zeitreihendaten punktuell erhält und keinerlei zeitliche Abhängigkeiten berücksichtigen kann. Es wird angenommen, dass der Treibstoffverbrauch zum Zeitpunkt ausschließlich von den Eingangsvariablen und einem Fehlerterm abhängt:

$$y_t = f(x_t) + \epsilon_t \quad (1)$$

**Aufgabe „Zeitreihen-Regression“:** Eine Zeitreihen-Regression, bei der zusätzlich noch ein kleines Fenster mit zeitlichem Kontext als Eingang für das KI-Modell zur Verfügung stehen. Es wird angenommen, dass der Treibstoffverbrauch zum Zeitpunkt von den Eingangsvariablen der letzten Zeitschritte und einem Fehlerterm abhängig ist:

$$y_t = f(x_t, x_{t-1}, \dots, x_{t-k}) + \epsilon_t \quad (2)$$

Für die Umsetzung vergleichen wir verschiedene Modellklassen:

1. **Physikbasierte Approximation:** Ein polynomiell Modell (3. Ordnung) dient als Baseline, angelehnt an das Admiralitätsgesetz, eine häufig in der Schiffsperformance-Analyse verwendete physikalische Näherung.
2. **Entscheidungsbäume:** CatBoost (Dorogush et. al, 2018) als state-of-the-art Gradient-Boosting-Algorithmus für strukturierte Daten.
3. **Neuronale Netze:** Multilayer Perceptron (MLP) als Standardarchitektur für tabellarische Daten sowie Long Short-Term Memory (LSTM)-Netzwerke für sequenzielle Modellierung.
4. **Foundation Models:** TabPFN, ein Transformer-basiertes Foundation Model für tabellarische Daten, das durch In-Context Learning auch mit

wenigen Beispieldaten leistungsfähige Regressionsmodelle erzeugt, ohne erneutes Training.

Die Modelle werden auf einem einheitlichen Datensatz für beide definierten Aufgaben evaluiert. Dafür werden die Daten von einem Kreuzfahrtschiff verwendet, die über ein Jahr aufgenommen wurden. Für die Evaluation wird eine fünffache Kreuzvalidierung verwendet. Für Zeitreihen Regression werden aus den zusammenhängenden Zeitreihendaten zusätzlich überlappende 1-Stunden-Fenster erstellt. Da in den Kontext des TabPFN-Modells nicht alle Trainingsdaten passen, wurde nur eine zufällige Stichprobe der Daten gewählt, bei der in zwei Experimenten einmal 500 und einmal 1.000 Datenpunkte verwendet wurden.

Als Eingangswerte verwenden wir die GPS-basierte Geschwindigkeit des Schiffs zusätzlich zu den Umgebungsdaten, wie der Meerestiefe, der Temperatur sowie Wetterdaten von Wind, Wellengang und Wasserströmung. Als Zielgröße wird der gesamte momentane Treibstoffverbrauch über alle Verbraucher verwendet.

## 3.3.1 Erkenntnisse und Diskussion

Abbildung 4 zeigt die Ergebnisse der KI-Modell Evaluation für beiden Aufgaben.

**Auswertung der Vorhersage.** In der tabularen Regressionsaufgabe zeigte sich, dass einfache physikbasierte Modelle wie das polynomiale Modell bereits grundlegende Zusammenhänge zwischen Geschwindigkeit und Verbrauch abbilden können. Das Vernachlässigen systematisch wichtiger externe Einflussgrößen wie Wind, Wellen oder Strömung führte jedoch zu einer deutlich geringeren Vorhersagegenauigkeit im Vergleich zu anderen datengetriebenen Modellen.

**CatBoost dominiert in datensparsamen Settings.** Im direkten Vergleich übertraf CatBoost alle anderen klassischen Modelle hinsichtlich Genauigkeit und Stabilität, sowohl in der tabularen als auch in der zeitbasierten Aufgabe. Die hohe Dateneffizienz dieses Algorithmus machte sich besonders bei der begrenzten Datenmenge bemerkbar. Im Gegensatz dazu zeigte das MLP, obwohl es komplexere nicht-lineare Zusammenhänge modellieren kann, eine höhere Varianz und etwas schwächere Leistung. Das deutet darauf hin, dass einfache neuronale Netze in datensparsamen Umgebungen schnell an ihre Grenzen stoßen können. Auch das sequenzielle LSTM schnitt schlecht ab, was vermutlich auf zu wenige Trainingsbeispiele und eine unzureichende Modellanpassung zurückzuführen ist.

**Besonders hervorzuheben ist das Foundation Model TabPFN.** Obwohl es ursprünglich für tabellarische Klassifikation entwickelt und nicht mit Verbrauchsdaten von Schiffen trainiert wurde, erzielte es in beiden Aufgabenklassen gute bis sehr gute Regressionsresultate und das bereits bei einem sehr kleinen Satz von Beispieldaten (500 bis 1.000 Stichproben). Durch seine Fähigkeit zum In-Context Learning lieferte TabPFN robuste Vorhersagen ohne erneutes Training auf dem Ziel-Datensatz. Diese Eigenschaft macht es besonders interessant für den Einsatz in der Prozessindustrie, wo große, sauber gelabelte Datensätze meist nicht verfügbar sind.

**Fluktuierende Prozessbedingungen.** Ein weiteres Ergebnis der Untersuchung ist der Leistungsgewinn durch die Berücksichtigung zeitlicher Kontextinformation. Insbesondere bei CatBoost und TabPFN zeigte sich, dass die Integration vergangener Zustände in die Vorhersage den mittleren absoluten Fehler (Mean Absolute Error, MAE) redu-

ziert und die Erfassbarkeit dynamischer Verbrauchsmuster verbessert. Dies gilt insbesondere für Anlagen mit stark fluktuierenden Prozessbedingungen, wie es auch in vielen chemischen oder pharmazeutischen Prozessen der Fall ist.

**On-the-Edge-Eignung und Echtzeitfähigkeit.** Alle untersuchten Methoden eignen sich für den Einsatz On-the-Edge, da sie sowohl beim Training als auch bei der Inferenz nur geringe Rechen- und Speicherressourcen benötigen. CatBoost zeigt sich mit seiner CPU-Effizienz besonders geeignet für eingebettete Systeme. MLP und LSTM bleiben aufgrund ihrer geringen Parameterzahl von nur wenigen tausend Parametern leichtgewichtig und auf Edge-Hardware lauffähig; in unseren Experimenten auf Laptop-CPUs lagen die Inferenzzeiten im Millisekundenbereich und bestätigen damit die Echtzeitfähigkeit. TabPFN ist als Foundation Model größer, lässt sich jedoch bereits auf Consumer-Grade-Grafikkarten innerhalb weniger Sekunden ausführen.

**Übertragbarkeit auf die Prozessindustrie.** Die Erkenntnisse aus der Schifffahrtsdomäne lassen sich direkt auf industrielle Prozesse übertragen.

- **Modellierung ohne Label:** Die Verwendung eines Hilfssignals (z. B. Energieverbrauch, Emissionen) als Zielgröße erlaubt es, Modelle auch ohne manuell annotierte Daten zu trainieren.
- **Wenige Trainingsdaten:** Besonders Foundation Models wie TabPFN ermöglichen KI-Einsatz auch in datenarmen Umgebungen, ein häufiges Hindernis in industriellen Settings.
- **On-the-Edge-Eignung:** Die untersuchten Modelle lassen sich effizient direkt auf Sensorhardware oder lokalen Gateways ausführen, was einen entscheidenden Vorteil in vernetzten, aber oft isolierten Produktionsumgebungen darstellt.

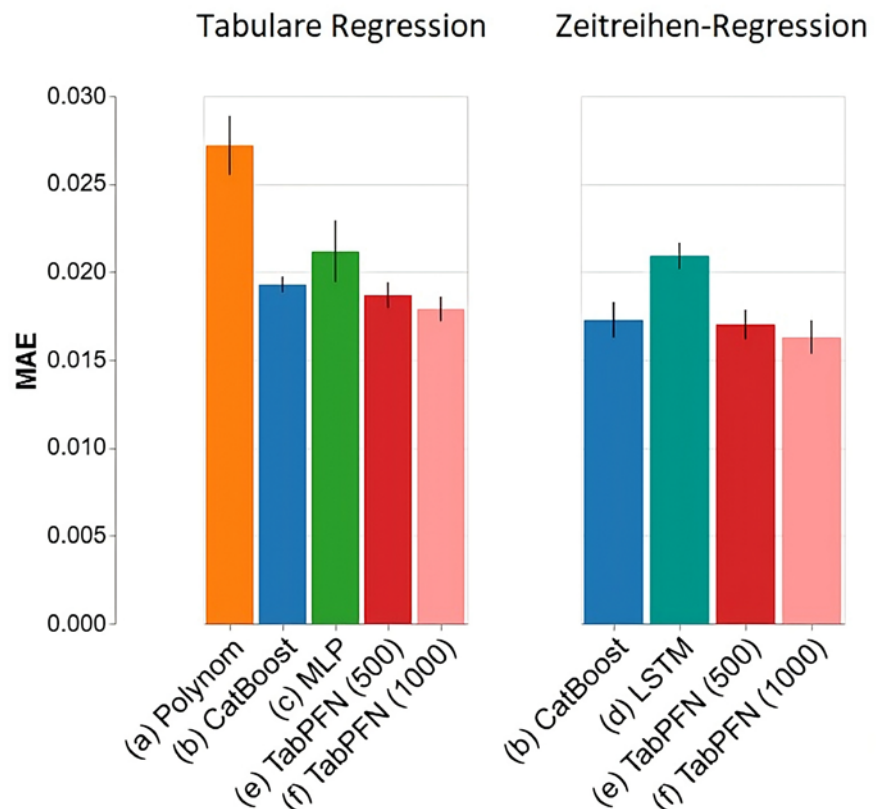


Abbildung 4: Vergleich verschiedener KI-Modelle für die Verbrauchsvorhersage anhand des mittleren absoluten Fehlers (Mean Absolute Error, MAE). Alle Modelle erzielten einen geringeren Fehler als das Polynom (orange), das als Baseline verwendet wurde. Das Hinzunehmen von zeitlichem Kontext (rechts) reduziert den Fehler gegenüber der rein tabellarischen, punktuellen Vorhersage (links). TabPFN schneidet am besten ab, obwohl es die wenigsten Trainingsdaten gesehen hat.

Die Ergebnisse zeigen, dass moderne KI-Methoden, insbesondere solche, die zeitlichen Kontext und externe Einflussgrößen berücksichtigen, in der Lage sind, Prozesse datenbasiert und robust abzubilden. In Kombination mit klassischen Optimierungsverfahren lässt sich eine KI-gestützte Prozessoptimierung praktisch umsetzen. Besonders hervorzuheben ist, dass dies auch mit kleinen Datenmengen und ohne explizites Expertenwissen über den Prozess gelingen kann.

### 3.4 Vertrauenswürdige KI

**Weiterer Kostentreiber in der J-Kurve: Regulatorik.** Mit der zunehmenden Nutzung KI-gestützter Systeme wird Regulatorik ein weiterer Kostentreiber in der J-Kurve. Besonders relevante Regulatorik für KI sind der AI Act, der

Digital Services Act (DSA) [10] und der Digital Markets Act (DMA) [11]. Der AI Act [12] stuft viele industrielle Anwendungen als hochriskant ein, was strenge Anforderungen an Transparenz, Datenqualität und Kontrolle nach sich zieht. Virtuelle Assistenten unterliegen dabei denselben Vorgaben, wenn sie mit Menschen interagieren oder Entscheidungen treffen. Zusätzlich greifen Datenschutzpflichten nach DSGVO, etwa bei der Verarbeitung biometrischer Daten. Unternehmen müssen daher nicht nur technisch, sondern auch rechtlich belastbare KI-Lösungen entwickeln.

**Zuverlässige und quellentreue Modelle.** Die Anwendung künstlicher Intelligenz in zeitkritischen Domänen erfordert ein tiefgehendes Verständnis für Vertrauenswürdigkeit, wie zuverlässig

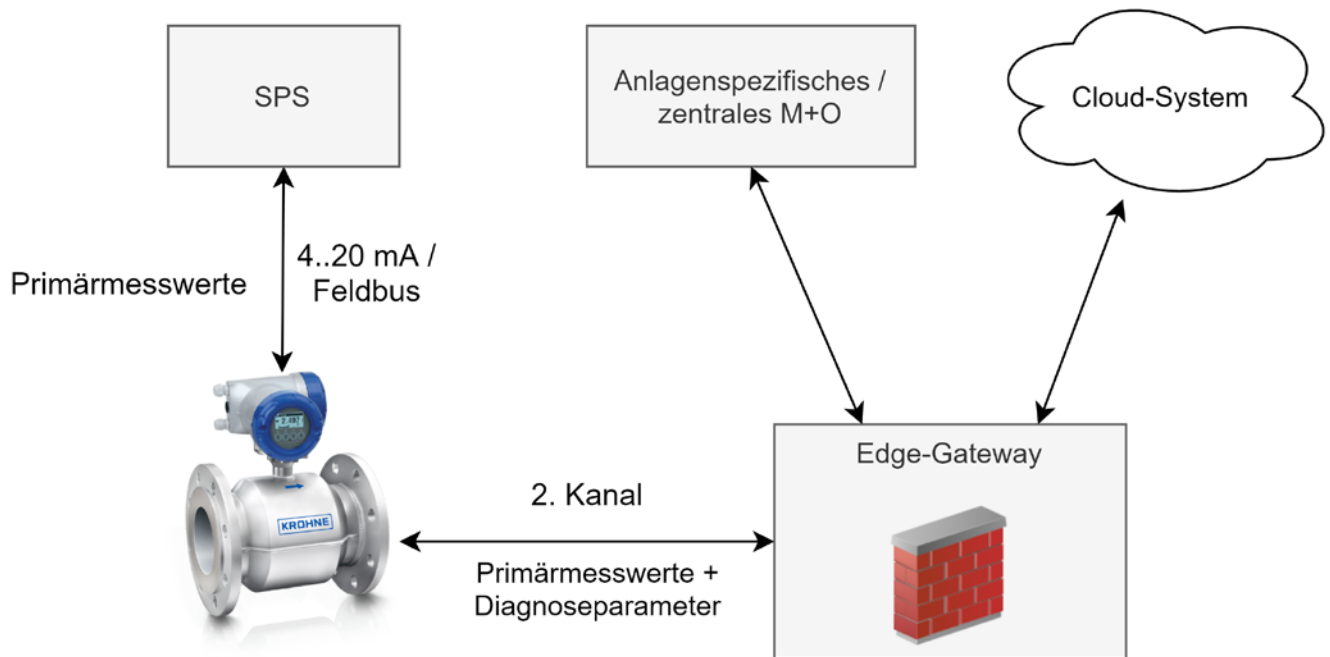


Abbildung 5: Anbindung von Messgeräten in Brownfield-Anlagen über den 2. Kanal mithilfe eines Edge-Gateways.

und verantwortungsvoll ein KI-System handelt, und Faithfulness, also Quellentreue, die Fähigkeit Informationen korrekt wiederzugeben, ohne Inhalte zu erfinden oder verfälschen. Die Ethikleitlinien der Europäischen Kommission für vertrauenswürdige KI [13] formulieren sieben zentrale Anforderungen: (1) Vorrang menschlichen Handelns und Aufsicht, (2) technische Robustheit und Sicherheit, (3) Schutz der Privatsphäre und Datenqualität, (4) Transparenz, (5) Fairness, (6) gesellschaftliches Wohlergehen und (7) Rechenschaftspflicht. Für KI-Systeme, die auf Zeitreihendaten basieren, erweisen sich insbesondere vier dieser Dimensionen als technisch besonders herausfordernd.

**Technische Herausforderungen bei zeitabhängigen Modellen.** Technische Robustheit, insbesondere für Daten, die sich kontinuierlich ändern, ist in der Prozessindustrie besonders relevant. Concept Drift, also die Änderung statistischer Eigenschaften über die Zeit hinweg [14], ist für KI ein besonders schwieriges Problem. Diese Dynamik gefährdet

grundlegende Modellcharakteristika wie Genauigkeit, Reproduzierbarkeit und Stabilität. Anders als bei statischen Daten, bei denen validierte Modelle oft dauerhaft zuverlässig arbeiten, erfordern Zeitreihenmodelle kontinuierliche Überwachung und gegebenenfalls Re-Kalibrierung. Spezielle Evaluationsverfahren wie Walk-Forward-Validation [15], [16] sind notwendig, um Driftprozesse zu erkennen und modellseitig zu kompensieren.

**Neuere Ansätze wie Foundation Models mit Transformer-basierter Architektur** erlauben durch erweiterte Kontexteingaben ein kontinuierliches parameterfreies Update der aktuell in der Anlage gemessenen Daten zur Inferenzzeit. Das ermöglicht kontinuierliche Vorhersagen zur Inferenzzeit. Beispiel für diese Systeme sind TabPFN, siehe Leaderboard [17]. Dies fördert nicht nur die Robustheit gegen Drift, sondern ermöglicht auch eine höhere funktionale Autonomie, da Modelle Entscheidungen auf Basis längerer historischer Muster treffen und bei kontinuierlichen

Updates treffen können. Diese zwei Aspekte sind besonders wichtig bei kontextsensitiven Prognosen, z. B. im Energiemanagement.

**Die Dimension Transparenz** stellt bei sequenziellen Daten besonders hohe Anforderungen. Während bei klassischen ML-Modellen Feature-Importance ausreichend sein kann, müssen bei Zeitreihen zusätzlich die zeitliche Relevanz einzelner Datenpunkte und kausale Beziehungen zwischen Zeitabschnitten nachvollziehbar gemacht werden. Dies erfordert komplexe Erklärungsmechanismen [18], etwa attention-basierte Visualisierungen oder zeitfensterspezifische Attributionstechniken [19].

**Datenqualitätsmanagement** wird primär durch die zeitliche Struktur der Daten herausgefordert. Temporale Konsistenz muss sichergestellt, fehlende Datenpunkte korrekt behandelt und Sensordrift sowie Kalibrierungsabweichungen über längere Zeiträume erkannt werden. Die Unterscheidung zwischen saisonalen Mustern und ech-

Tabelle 1: Herausforderungen und Lösungsansätze für vertrauenswürdige KI-Systeme auf Zeitreihen werden entlang der Vertrauenswürdigkeitsdimensionen der EU-Leitlinien dargestellt. Die Tabelle zeigt zentrale Probleme und passende technische Lösungen.

| Dimension (nach EU Guidelines)                         | Herausforderungen bei Zeitreihendaten                                                               | Typische technische Anforderungen / Lösungen                                                                                                          |
|--------------------------------------------------------|-----------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------|
| Technische Robustheit & Sicherheit                     | Concept Drift, zeitabhängige Degradierung der Modelleleistung                                       | Walk-Forward-Validation, Drift-Erkennung, kontinuierliche Re-Kalibrierung, Nutzung von Foundation Models für kontextstabile Langzeitprognosen         |
| Transparenz                                            | Zeitliche Abhängigkeiten erschweren Interpretation; Kausalität über Zeitfenster                     | Attention-Visualisierung, zeitfensterbasierte Feature-Attribution, erklärbare Architekturen wie Temporal Fusion Transformer (TFT)                     |
| Datenqualität                                          | Sensordrift, fehlende Werte, Akkumulation systematischer Fehler, Lag-Effekte                        | Echtzeit-Qualitätsmetriken, Imputation temporaler Lücken, Validierung zeitlicher Konsistenz, Anomalieerkennung für Sensordegradation                  |
| Vorrang menschlichen Handelns & Aufsicht               | Schnelle, automatisierte Entscheidungen in dynamischen Umgebungen, erschwerte menschliche Kontrolle | Echtzeit-Feedback, adaptive Schwellenwerte, Interventionsmechanismen, erklärbare Ausgaben für schnelles Verständnis                                   |
| Modularität & Spezialisierung (Zusatzdimension)        | Komplexität steigt mit Modellgröße und Echtzeitanforderungen                                        | Arbeitsteilige Systeme mit spezialisierten Sub-agenten (z. B. Drift-Detektor, Unsicherheitsmodul, Erklärbarkeitsmodul)                                |
| Unsicherheitsquantifizierung & Vertrauenskommunikation | Prognoseunsicherheit nimmt mit dem Zeithorizont zu                                                  | Deep Ensembles, Monte Carlo Dropout, Konfidenzintervalle, Kalibrierung der Unsicherheiten, visuelle Unsicherheitskommunikation (z. B. Forecast-Bands) |

ten Qualitätsproblemen ist oft nicht trivial. Zusätzlich erschweren Lag-Effekte die frühzeitige Erkennung von Qualitätsverlusten, da sich deren Folgen häufig erst mit Verzögerung auf die Modellperformance auswirken.

Tabelle 1 zeigt eine erweiterte Übersicht über für Zeitreihen relevante Dimensionen der Vertrauenswürdigkeit mit ihren Herausforderungen und einigen technischen Lösungsmöglichkeiten.

**Bewertung und Verbesserung der Vertrauenswürdigkeit.** Vertrauenswürdige KI erfordert hier menschenzentrierte Schnittstellen, die auch unter Zeitdruck Interventionen ermöglichen und erklärbare Entscheidungsgrundlagen liefern [20]. Vertrauenswürdige KI-Systeme in der Prozessindustrie sind solche, bei denen die automatisierte Bedienung der Anlage sich genauso oder besser nachvollziehbar verhält,

als wenn ein Fachangestellter dieses Bedienen würde. Im Idealfall können die Modelle mit gleichem oder geringerem Risikopotenzial verwendet werden. Außerdem muss Vertrauen bestehen, dass die KI, wenn anormale Situationen auftreten, eine Erklärung liefern kann. Diesen Idealzustand wollen wir erreichen. Methoden zur Bewertung der Datenqualität und Vertrauenswürdigkeit sind wiederum basiert auf ML. Diese Methoden sind wichtige Indikatoren für ESG Compliance.

Die Bewertung vertrauenswürdiger KI-Systeme mit Zeitreihendaten muss sowohl modellspezifische als auch systemische Komponenten berücksichtigen. Rein technische Fehlermaße wie Mean Average Error (MAE) oder Root Mean Squared Error (RMSE) reichen nicht aus, um Vertrauen zu begründen, da sie weder Drift noch Unsicherheiten angemessen abbilden. Stattdessen

sind zeitabhängige Validierungsstrategien wie Walk-Forward- oder Rolling-Window-Validation notwendig, um Modellverhalten realistisch zu beurteilen.

**Ein zentraler Aspekt ist die Unsicherheitsquantifizierung.** Verfahren wie Monte-Carlo Dropout [21] und Deep Ensembles [22] ermöglichen probabilistische Vorhersagen, die über Punktwerte hinausgehen und Vertrauensintervalle bereitstellen. Ensembles, also Kombinationen mehrerer Basismodelle, verbessern dabei nicht nur die Punktgenauigkeit, sondern reduzieren auch das Risiko von Überanpassungen an Driftphasen und erhöhen die Modellrobustheit bei Zeitreihenvorhersagen [23]. Diese Unsicherheitsabschätzungen sind insbesondere bei längerfristigen Prognosen essenziell, da Unsicherheit mit steigendem Vorhersagehorizont zunimmt. Allerdings bedarf



es zusätzlich geeigneter Verfahren zur Kalibrierung dieser Unsicherheiten, da falsch eingeschätzte Konfidenzbereiche zu Fehlentscheidungen führen können.

Auch die Erklärbarkeit sequenzieller Modelle ist entscheidend. Während klassische Attributionsmethoden auf statische Daten zugeschnitten sind, liefern neue architekturelle Ansätze, z. B. der Temporal Fusion Transformer (TFT) [24], kombinierte Lösungen aus Attention-Mechanismen und Unsicherheitsausgaben, die sowohl interpretierbare als auch robuste Vorhersagen ermöglichen. Ergänzt werden diese durch zeitenfensterspezifische Relevanzanalysen, wie sie beispielsweise in den Arbeiten von [19] dargestellt werden.

Erklärbarkeit für den Endnutzer ist bisher oft schwierig in hochdimensionalen Datenräumen (Bild, Text, etc.) und bisher vor allem für statische und sequenzielle Modelle möglich. In letzterem Falle insbesondere über das Highlighting von Attention-Mechanismen und die Möglichkeit mit dem System reden zu können. Diese Verfahren sind aber bisher eher für Entwickler von Modellen hilfreich.

**Vertrauen ein interaktives Phänomen, das im Zusammenspiel von Mensch und System entsteht.** Studien zeigen, dass Nutzer:innen dann Vertrauen entwickeln, wenn Systeme erklärbar, vorhersehbar und kontrollierbar bleiben [20]. Dies erfordert Visualisierungen für Unsicherheit, Schwellenwertanzeigen und adaptive Warnmechanismen, die auch unter Zeitdruck verständlich und handhabbar bleiben. Zunehmend werden auch selbstüberwachende Modelle entwickelt, die automatisch auf Drift oder Anomalien reagieren und proaktiv Hinweise zur Gültigkeit ihrer eigenen Vorhersagen geben.

### 3.5 Beispiel: Vertrauenswürdige Datengrundlage für Expertendiagnose

Ein erfolgreiches Beispiel für die nachträgliche Erschließung vertrauenswürdiger und qualitativ hochwertiger Daten liefert die Implementierung von Edge-Gateways in Brownfield-Anlagen, die bestehende Feldgeräte über ungenutzte Feldbusschnittstellen wie HART oder Modbus sicher und rückwirkungsfrei anbinden, ohne Eingriff in die bestehende Prozessführung. Industrielle Lösungen, wie das von KROHNE entwickelte Edge-Gateway [1], dass eine solche Anbindung auch unter Praxisbedingungen wirtschaftlich und robust realisierbar ist. Abbildung 4 illustriert den schematischen Aufbau dieser Lösung.

Diese Systeme orientieren sich an den Prinzipien der NAMUR Open Architecture (NOA), die eine sichere und rückwirkungsfreie Erschließung von Geräteinformationen über einen zweiten Kommunikationskanal ermöglichen. Zugleich adressieren sie zentrale Anforderungen der Ethikleitlinien für vertrauenswürdige KI der Europäischen Kommission, etwa in den Dimensionen Transparenz, technischer Robustheit und Vorrang menschlichen Handelns. Durch integrierte Gerätesemantik und IT-Sicherheitsmechanismen wird eine gezielte Filterung und Auswertung der übertragenen Daten ermöglicht. Nur lesende Zugriffe werden zugelassen, schreibende Zugriffe auf prozesskritische Parameter zuverlässig unterbunden. Die so erhobenen Laufzeitdaten erlauben eine prozessnahe Diagnostik und bilden eine belastbare Grundlage für erklärbare, kontextbasierte KI-Analysen. Durch das Einspielen der Kommunikation in ein Cloud-System entsteht ein bidirektionaler Feedback-Loop. Die in der Cloud aggregierten Daten aus verschiedenen Anlagen können zur

Verbesserung von Diagnosealgorithmen und Vorhersagemodellen genutzt werden. Diese Modelle lassen sich gezielt auf die Edge-Systeme zurückspielen, sodass ein kontinuierlicher Lernprozess entsteht, der sowohl die Modellgüte als auch das Vertrauen in KI-gestützte Assistenzsysteme im Feld stärkt.

**Vertrauen durch erweiterte Diagnosewerte.** Ein besonderer Mehrwert entsteht durch den Zugriff auf zusätzliche Roh- und Diagnosewerte über den zweiten Kommunikationskanal. Diese Daten bieten wertvolle Informationen über die Zuverlässigkeit und Vertrauenswürdigkeit der Primärmesswerte, gehen aber in Umfang und Komplexität oft über das hinaus, was im Leitsystem dargestellt oder manuell bewertet werden kann. KI-Algorithmen schließen hier die Lücke. Sie extrahieren aus dem kontinuierlichen Strom von Rohdaten automatisch wesentliche Zustandsinformationen und Kennzahlen, beispielsweise zur Drifterkennung oder Gerätestabilität.

Besonders Hersteller von Messgeräten können durch ihr tiefes Systemverständnis KI-Modelle entwickeln, die interne Diagnosedaten präzise interpretieren und auch subtilen Abweichungen Bedeutung zuweisen. Die auf dem Edge-Gateway ausgeführten Modelle machen diese Einschätzungen vor Ort verfügbar, angereichert mit Visualisierungen, die dem Betriebspersonal eine Einordnung und Validierung der KI-Ergebnisse ermöglichen. Dies stärkt sowohl die Transparenz als auch den Vorrang menschlichen Handelns im Sinne eines vertrauenswürdigen und verantwortungsvollen Einsatzes von KI in der Prozessindustrie.

### 4. Was sind unsere Chancen?

Die Prozessindustrie kann KI unabhängig von großen Plattformen entwickeln, da die entscheidenden Daten aus den

eigenen Abläufen stammen. Das eröffnet die Chance, souverän zu handeln und maßgeschneiderte Lösungen aufzubauen. Besonders wertvoll ist, dass Unternehmen direkt auf hochwertige Prozessdaten zugreifen und diese durch Feedback oder synthetische Daten ergänzen können. Mit Methoden, die auch kleine Datensätze effizient nutzen, lässt sich KI ressourcenschonend einsetzen.

Besonders interessant ist die Möglichkeit, kleinere und spezialisierte Modelle zu modularen Systemen zu kombinieren, die Prozesse adaptiv steuern und kontinuierlich verbessern. Diese bilden eine Grundlage, die in diesem Artikel beschriebenen vier Bausteine für autonome Prozessführung Schritt für Schritt umzusetzen von der Vernetzung über fortgeschrittene Diagnostik und prozessoptimierende Assistenzsysteme bis hin zu adaptiven Produktionssystemen.

## 5. Fazit

Im Ergebnis zeigt sich, dass auch die Nutzung von KI in der Prozessindustrie oft der Logik einer J-Kurve folgt: Anfangs sind Investitionen erforderlich, insbesondere für die Erzeugung und Aufbereitung von Trainingsdaten, für die Weiterentwicklung geeigneter Lernverfahren und für die Sicherung hoher Datenqualität im Hinblick auf regulatorische Vorgaben. Langfristig überwiegen jedoch die Chancen. Souveräne und vertrauenswürdige KI-Lösungen können Effizienz und Nachhaltigkeit deutlich steigern und zugleich

regulatorische Sicherheit bieten. Wer jetzt beginnt, diese Grundlagen zu schaffen, legt den Grundstein für eine autonome, resiliente und zukunftsfähige Prozessindustrie.

## Referenzen

- [1] D. Kuschnerus, „Digitale Kommunikation im Brownfield nachrüsten“, atp magazin, 2025.
- [2] Bünte, O. (2025). Amazon "employs" over one million robots", heise online. Verfügbar unter: <https://www.heise.de/en/news/Amazon-employs-over-one-million-robots-10465702.html>
- [3] Hermann, M., Pentek, T., Otto, B. (2016). Design principles for industrie 4.0 scenarios. In 2016 49th Hawaii international conference on system sciences (HICSS) (pp. 3928-3937). IEEE.
- [4] Diemer, J., Elmer, S., Gaertler, M., Gamer, T., Görg, C., Grotepass, J., ... Winter, J. (2020). KI in der Industrie 4.0: Orientierung, Anwendungsbeispiele, Handlungsempfehlungen. Wegweiser.
- [5] Brynjolfsson, E., Rock, D., Syverson, C. (2019). The productivity j-curve: How intangibles complement general purpose technologies. Verfügbar unter: [https://ide.mit.edu/sites/default/files/publications/2019-04/JCurvebrief.final2\\_.pdf](https://ide.mit.edu/sites/default/files/publications/2019-04/JCurvebrief.final2_.pdf)
- [6] Challapally, A., Pease, C., Raskar, R., Chari, P. (2025). State of AI in business 2025 report (Version 0.1). MIT NANDA. Verfügbar unter: [https://mlq.ai/media/quarterly\\_decks/v0.1\\_State\\_of\\_AI\\_in\\_Business\\_2025\\_Report.pdf](https://mlq.ai/media/quarterly_decks/v0.1_State_of_AI_in_Business_2025_Report.pdf)
- [7] Hollmann, N., Müller, S., Eggensperger, K., & Hutter, F. (2022). TabPFN: A transformer that solves small tabular classification problems in a second. arXiv preprint arXiv:2207.01848.[8] A b - din, M. I., Jyoti, A., Behl, H., Bubeck, S., Eldan, R., Gunasekar, S., Harrison, M., ... Zhang, Y. (2024). Phi-4 technical report. Verfügbar unter: <https://www.microsoft.com/en-us/research/publication/phi-4-technical-report/>
- [8] Viga, J., Mueck, P., Löser, A., Weis, T. (2025). Fuel-Cast: Benchmarking Tabular and Temporal Models for Ship Fuel Consumption. arXiv preprint arXiv:2510.08217.
- [9] Europäische Kommission. (2022). Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market for Digital Services (Digital Services Act) and amending Directive 2000/31/EC. Verfügbar unter: <https://eur-lex.europa.eu/eli/reg/2022/2065/oj>
- [10] Europäische Kommission. (2022). Regulation (EU) 2022/1925 of the European Parliament and of the Council of 14 September 2022 on contestable and fair markets in the digital sector (Digital Markets Act) and amending Directives (EU) 2019/1937 and (EU) 2020/1828". Verfügbar unter: <https://eur-lex.europa.eu/eli/reg/2022/1925/oj>
- [11] Europäische Kommission. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts". Verfügbar unter: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- [12] Europäische Kommission. (2019). Ethics guidelines for trustworthy AI. Verfügbar unter: <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>
- [13] Gama, J., Žliobaitė, I., Bifet, A., Pechenizkiy, M., Bouchachia, A. (2014). A survey on concept drift adaptation. ACM computing surveys (CSUR), 46(4), 1-37.
- [14] Hyndman, R. J., Athanasopoulos, G. (2018). Forecasting: Principles and practice (2nd ed)". Verfügbar unter: <https://otexts.org/fpp2/>
- [15] Cerqueira, V., Torgo, L., Mozetič, I. (2020). Evaluating time series forecasting models: An empirical study on performance estimation methods. Machine Learning, 109(11), 1997-2028.
- [16] Erickson, N., Purucker, L., Tschalzev, A., Holzmüller, D., Desai, P. M., Salinas, D., Hutter, F. (2025). Tabarena: A living benchmark for machine learning on tabular data. arXiv preprint arXiv:2506.16791.
- [17] Schlegel, U., Arnout, H., El-Assady, M., Oelke, D., & Keim, D. A. (2019, October). Towards a rigorous evaluation of XAI methods on time series. In 2019 IEEE/CVF International Conference on Computer Vision Workshop (ICCVW) (pp. 4197-4201). IEEE.
- [18] Tonekaboni, S., Joshi, S., Mccradden, M. D., Goldenberg, A. (2020). Explaining Time Series Predictions through Selective Relevance", Proceedings of the AAAI Conference on Artificial Intelligence, Bd. 34, S. 5434-5441.
- [19] Amershi, S., Weld, D., Vorvoreanu, M., Fournery, A., Nushi, B., Collisson, P., ... Horvitz, E. (2019). Guidelines for human-AI interaction. In Proceedings of the 2019 chi conference on human factors in computing systems (pp. 1-13).
- [20] Gal, Y., Ghahramani, Z. (2016). Dropout as a bayesian approximation: Representing model uncertainty in deep learning. In international conference on machine learning (pp. 1050-1059). PMLR.
- [21] Lakshminarayanan, B., Pritzel, A., Blundell, C. (2017). Simple and scalable predictive uncertainty estimation using deep ensembles. Advances in neural information processing systems, 30.
- [22] Zhang, Y., Chen, X., Li, W., Wang, J. (2025). Deep Ensemble Learning for Time Series Forecasting: A Survey". Verfügbar unter: <https://arxiv.org/abs/2207.03599>
- [23] Lim, B., Arık, S. Ö., Loeff, N., Pfister, T. (2021). Temporal fusion transformers for interpretable multi-horizon time series forecasting. International journal of forecasting, 37(4), 1748-1764.

**Penelope Mück** ist Senior Entwicklungsingenieur bei Krohne Messtechnik GmbH und arbeitet im Bereich digitale Produkte und Künstliche Intelligenz. Ihre Schwerpunkte liegen in den Bereichen Machine Learning auf Zeitreihendaten, vertrauenswürdige KI und Sprachmodelle. Sie promoviert in der Informatik an der Universität Duisburg-Essen bei Prof. Dr.-Ing. Torben Weis. Neben ihrer akademischen Tätigkeit verfügt sie über mehrjährige Industrienerfahrung im Bereich Softwareentwicklung und Data Science, sowohl als Entwickler als auch in Leitungsfunktionen.

### **Penelope Mück**

Krohne Messtechnik GmbH  
Ludwig-Krohne-Str. 5  
47058 Duisburg  
p.mueck@krohne.com

**Justus Viga** ist Senior Development Engineer und Data Scientist bei KROHNE Digital. In seiner Arbeit verbindet er Softwareentwicklung, datengetriebene Analytik und Projektleitung, um industrielle Prozesse durch KI-basierte Methoden zu optimieren. Im Rahmen seiner Promotion erforscht er den Transfer und die Individualisierung von KI-Algorithmen, mit dem Ziel, lernfähige Systeme effizient auf unterschiedliche Anlagen und Prozesse zu übertragen.

### **Justus Viga**

Krohne Messtechnik GmbH  
Ludwig-Krohne-Str. 5  
47058 Duisburg  
j.viga@krohne.com

**Alexander Löser** lehrt und erforscht seit September 2013 Data Science an der Berliner Hochschule für Technik und leitet dort das Forschungszentrum Data Science. Seit 2018 berät Alexander auch das BMBF im nationalen KI-Flaggschiffprogramm Plattform Lernende Systeme und ist Mentor für KI-Startups und klinische Ärzte an der Charité Berlin bei der Erforschung zukünftiger medizinischer Datenprodukte. Seit 2011 hat er IBM, Zalando, eBay, Munich Re, Fresenius, Krohne Messtechnik, Babbel und zahlreichen anderen Unternehmen geholfen, mehr als 52 Datenprodukte mit Geschäftsmodellen auf sechs Plattformen zu entwickeln. Außerdem unterstützte er dabei, ganze KI-Bereiche aufzubauen und deren KI-Strategie zu definieren.

### **Prof. Dr.-Ing habil. Alexander Löser**

Berliner Hochschule für Technik (BHT)  
Luxemburger Straße 10  
13353 Berlin  
aloeser@bht-berlin.de

**Dagmar Dirzus** ist seit Oktober 2021 als Vice President AI & Platform Business bei Krohne tätig. Gemeinsam mit ihrem Team entwickelt sie Daten- und KI-basierte Services. Zuvor war sie mehr als sieben Jahre lang Geschäftsführerin der Gesellschaft Mess- und Automatisierungstechnik im VDI/VDE. Sie promovierte im Fachbereich Fertigungstechnik der RWTH Aachen und ist seit vielen Jahren im Beirat des atp magazins aktiv.

### **Dr. Dagmar Dirzus**

KROHNE Messtechnik GmbH  
Ludwig-Krohne-Straße 5  
47058 Duisburg  
d.dirzus@krohne.com