

dfki ai next



Wie KI sich in der Welt orientiert

Von der Wahrnehmung zur Handlung



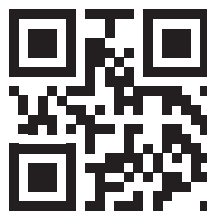
KI für den Menschen – Intelligente Lösungen für die Wissensgesellschaft

Das DFKI forscht seit über 35 Jahren an KI für den Menschen und orientiert sich an gesellschaftlicher Relevanz und wissenschaftlicher Exzellenz in den entscheidenden zukunftsorientierten Forschungs- und Anwendungsgebieten der Künstlichen Intelligenz.

Die Deutsche Forschungszentrum für Künstliche Intelligenz GmbH (DFKI) wurde 1988 als gemeinnützige Public-Private-Partnership gegründet. Das DFKI unterhält Standorte in Kaiserslautern, Saarbrücken, Bremen, Niedersachsen (Osnabrück und Oldenburg), Darmstadt, Labore in Berlin und Lübeck sowie eine Außenstelle in Trier.

In 29 Forschungsbereichen, 13 Kompetenzzentren und acht Living Labs werden ausgehend von anwendungsorientierter Grundlagenforschung Produktfunktionen, Prototypen und patentfähige Lösungen im Bereich der Informations- und Kommunikationstechnologie entwickelt. Die Finanzierung erfolgt über Zuwendungen öffentlicher Fördermittelgeber sowie durch Entwicklungsaufträge aus der Industrie.

Projektergebnisse und Meilensteine werden periodisch institutionell und durch ein international besetztes Expertengremium (Wissenschaftlicher Beirat) begutachtet. Neben dem Bundesministerium für Forschung, Technologie und Raumfahrt und den Bundesländern Rheinland-Pfalz, Saarland, Bremen, Niedersachsen und Hessen sind im DFKI-Aufsichtsrat zahlreiche namhafte deutsche und internationale Hochtechnologie-Unternehmen aus einem breiten Branchenspektrum vertreten.



www.dfki.de

KI föderal und europäisch denken

Internationale Kooperation, wissenschaftliche Sichtbarkeit und strategische Souveränität

Künstliche Intelligenz wird nicht im Alleingang vorangebracht. Kooperationen mit starken europäischen oder internationalen Partnern wie DFKI und INRIA oder die Arbeit von CERTAIN zeigen, dass wissenschaftliche Exzellenz dort entsteht, wo Forschung über Instituts- und Ländergrenzen hinweg zusammengeführt wird.

Das Netzwerk der deutschen KI-Kompetenzzentren stärkt die internationale Strahlkraft der deutschen KI-Forschung und etabliert KI aus Europa als Markenzeichen für eine exzellente und vertrauenswürdige KI. Verbünde wie das Robotics Institute Germany (RIG) bündeln Kompetenzen, schaffen Anschlussfähigkeit zwischen Disziplinen und machen sichtbar, dass KI-Forschung heute oft dort am stärksten ist, wo Wahrnehmung, Orientierung und Handeln nicht getrennt, sondern gemeinsam gedacht werden.

Auf den großen internationalen Bühnen – von CVPR, ICML und anderen IEEE-Konferenzen bis zur IJCAI – wird genau diese Forschung verhandelt, die den nächsten Schritt der KI prägt.

Die Hightech Agenda Deutschland setzt dafür den politischen Rahmen. Ihr Ziel ist klar: die technologische Souveränität stärken, die Verfügbarkeit von KI-Kapazitäten verbessern und Deutschland in der nächsten Generation von KI international anschlussfähig halten.

Diese Ausgabe nimmt das als Ausgangspunkt. Sie zeigt, wie aus exzellenter Forschung, internationaler Kooperation und domänen-spezifischer Vernetzung eine herausragende KI-Landschaft entstehen kann.

Wie KI sich in der Welt orientiert

Von der Wahrnehmung zur Handlung

Künstliche Intelligenz steht an einem Punkt, an dem bloße Leistungssteigerung nicht mehr genügt. Die entscheidende Frage ist nicht, wie Systeme noch größer, schneller oder allgegenwärtiger werden. Die entscheidende Frage ist, welche Art von KI Europa entwickeln will, wenn technologische Handlungsfähigkeit, wissenschaftliche Eigenständigkeit und gesellschaftliche Verlässlichkeit zusammenkommen sollen.

Genau hier beginnt heute der eigentliche Anspruch von Forschung: KI nicht nur leistungsfähiger zu machen, sondern sie zur Orientierung in einer komplexen, offenen und von Menschen geprägten Welt zu befähigen.

Deutschland und Europa bringen dafür viel mit: wissenschaftliche Exzellenz, wachsende politische Aufmerksamkeit und ein geschärftes öffentliches Bewusstsein dafür, dass KI längst keine ferne Zukunftstechnologie mehr ist, sondern zur erfolgskritischen Strukturfrage von Wirtschaft, Wissenschaft und Gesellschaft geworden ist. Entscheidend ist nun, aus diesen Einsichten auch Zielvorgaben zu machen. Wer KI nur als globalen Wettlauf um Skalierung begreift, wird am Ende fremden Entwicklungslogiken unterworfen sein. Wer KI als Forschungsfrage ernst nimmt, kann eigene Maßstäbe setzen – bei Transparenz, Energieeffizienz, Kontrollierbarkeit und realer Einbettung. Technologische Souveränität beginnt deshalb nicht erst beim Markterfolg, sondern bei der Fähigkeit, die Grundlagen kommender Systeme selbst zu bestimmen.

Wenn Europa seine Abhängigkeit von geschlossenen außereuropäischen Systemen verringern will, muss es gerade dort investieren, wo die nächste KI-Generation entsteht. Dazu gehören neue Rechnerarchitekturen, energieeffiziente Verarbeitungsparadigmen und Infra-

strukturen, die nicht allein auf maximale Größe, sondern auf Belastbarkeit, Sicherheit, Effizienz und wissenschaftlich-industrielle Steuerbarkeit ausgelegt sind.

Noch grundlegender ist jedoch die Frage, welches Modell von KI wir überhaupt weiterentwickeln wollen. Die Erfolge großer generativer Systeme sind unübersehbar. Ebenso unübersehbar sind ihre Grenzen: enormer Datenbedarf, intransparente Entscheidungsprozesse, hoher Energieverbrauch und ein Grad an Unzuverlässigkeit, der in sicherheitskritischen oder hochspezialisierten Kontexten nicht ausreicht. Für viele Anwendungen in Industrie, Medizin, Energiewirtschaft oder Mobilität genügt es nicht, dass ein System häufig überzeugend wirkt. Es muss nachvollziehbar arbeiten, überprüfbar bleiben und seine Entscheidungen in belastbare Zusammenhänge stellen können. Gerade deshalb gewinnt eine Forschungsrichtung an Gewicht, die für Europa weit mehr ist als eine technische Alternative: hybride, neuro-explizite KI. Sie verbindet lernende Verfahren mit symbolischen, wissensbasierten und analytischen Komponenten. Dazu passt die Abkehr vom schlichten Größer-ist-besser-Denken. Kleine, spezialisierte Sprachmodelle und frugale KI-Architekturen zeigen, dass Effizienz selbst ein Fortschritt sein kann. Sie benötigen weniger Energie, weniger Spezialhardware und lassen sich präziser an konkrete Domänen anpassen – an in-

dustrielle Prozesse ebenso wie an wissenschaftliche Datenräume oder sensible Wissensumgebungen. Gerade für eine Forschungs- und Wirtschaftslandschaft, die stark von hochspezialisierten, regulierten und mittelständisch geprägten Strukturen lebt, ist das kein Nebenaspekt, sondern ein strategischer Vorteil. Verlässlichkeit unter realen Bedingungen ist hier wichtiger als maximale Allgemeinheit um jeden Preis.

Hinzu kommt ein weiterer Wandel: KI erzeugt nicht mehr nur textuelle oder visuelle Inhalte, sie beginnt zu planen, zu koordinieren und in längeren Aufgabenketten zu handeln. Mit der agentischen Wende verschiebt sich der Fokus erneut. Sobald Systeme in Handlungszusammenhänge eintreten, werden Transparenz, Eingriffsmöglichkeiten und menschliche Aufsicht zentral.

Der nächste Schritt der KI-Forschung besteht deshalb nicht darin, Autonomie zu entgrenzen, sondern sie so zu gestalten, dass sie kontrollierbar, nachvollziehbar und sozial wie technisch eingebettet bleibt. Erst dann wird aus Rechenleistung verlässliche Handlungsfähigkeit. Genau an diesem Punkt berührt sich die technologische

mit der wissenschaftlichen Herausforderung in dieser Ausgabe der ai next. Die eigentliche Zukunftsfrage lautet nicht nur, was KI erkennen kann, sondern wie sie sich in realen, dynamischen und mehrdeutigen Umgebungen orientiert. Wie wird aus Wahrnehmung belastbares Wissen? Wie aus Wissen eine robuste Einordnung? Und wie daraus Handeln, das nicht bloß effizient, sondern verantwortlich ist? In dieser Bewegung von Wahrnehmen über Orientieren zu Handeln liegt nicht nur eine plausible Heftdramaturgie, sondern ein präziser Forschungshorizont für die kommenden Jahre.

Europa muss für diesen Horizont keine fremden Modelle kopieren. Seine Stärke kann gerade darin liegen, KI als Forschungsgegenstand der Verlässlichkeit, der Orientierungssicherheit und der realweltlichen Einbettung zu entwickeln – offen für internationale Zusammenarbeit, aber klar in den eigenen wissenschaftlichen Maßstäben. Nicht das lauteste System wird den nächsten Schritt bestimmen, sondern dasjenige, das Leistung mit Nachvollziehbarkeit, Effizienz und Verantwortung verbindet. Genau dort beginnt eine KI, die sich nicht nur in Datenräumen bewegt, sondern in der Welt. —●



« Das Versprechen der hybriden, neuro-expliziten KI liegt nicht in der bloßen Vergrößerung bestehender Modelle, sondern in einer potenzierenden Qualitätsidee von maschineller Intelligenz, bei der die Vorteile der probabilistischen und deduktiven KI-Ansätze zusammenkommen. Systeme dieser Art erkennen nicht nur Muster, sondern repräsentieren Wissen explizit, strukturieren ihre Schlüsse und lassen sich unter definierten Bedingungen auditieren. Darin liegt eine wissenschaftliche Perspektive, die Leistungsfähigkeit und Vertrauenswürdigkeit nicht gegeneinander ausspielt, sondern systematisch zusammenführt. »

Prof. Dr. Antonio Krüger, CEO DFKI



Von der Geometrie zur Bedeutung

Drei Zugänge zu 3D-Szenenverständnis, Aufmerksamkeit und Rekonstruktion

Szenenverständnis beginnt dort, wo Geometrie, Beziehung und Perspektive zusammenkommen. Auf diese Dynamik haben DFKI-Forschende mit ihrer Arbeit einen Schwerpunkt gelegt und auf der CVPR 2026 vorgestellt: Sie begreift räumliche Rekonstruktion nicht mehr als Endpunkt, sondern als Ausgangsbasis für semantische Einordnung, menschenzentrierte Orientierung und gezielte Eingriffe in komplexe 360-Grad-Umgebungen.

ReLaGS steht für einen grundlegenden Schritt im 3D-Szenenverständnis. Das Verfahren verbindet Gaussian Splatting mit relationaler Sprache und einer hierarchischen Szenenstruktur, sodass Maschinen Objekte nicht nur erfassen, sondern auch ihre Beziehungen und semantischen Ebenen modellieren können. So wird aus reiner Rekonstruktion eine lesbare Szene.

Aus geometrischen Punkten und Flächen wird eine strukturierte Umgebung, in der sich räumliche Konstel-

lationen, Teil-Ganzes-Beziehungen und sprachlich beschreibbare Zusammenhänge explizit darstellen lassen. Die wissenschaftliche Pointe liegt darin, dass ReLaGS hierarchische und relationale 3D-Repräsentation in einem einheitlichen, szenenagnostischen Ansatz zusammenführt. So entsteht eine offene 3D-Szenengraph-Struktur ohne szenenspezifisches Training; zugleich verbessert das System die geometrische und semantische Stabilität der Rekonstruktion durch Pruning und robuste

ReLaGS

Szenen lesbar machen

Aus reiner Rekonstruktion wird eine strukturierte Umgebung mit Objekten, Ebenen und Beziehungen.

DriverGaze360

Orientierung verstehen

Wohin richten Menschen ihre Aufmerksamkeit – rundum, nicht nur nach vorn?

Inpaint360GS

In Szenen eingreifen

Fehlendes wird ergänzt, ohne die innere Struktur der Szene zu zerstören.



Merkmalsaggregation. In einer Welt, in welcher strukturierte 3D-Modelle unserer Welt zunehmend an internationaler Bedeutung gewinnen, lehren ExpertInnen Maschinen nicht nur, was sie in einer Szene wahrnehmen können – sondern wie ihre Elemente zusammenhängen. DriverGaze360 verschiebt diese Frage in eine andere, ebenso wichtige Richtung: Wie orientieren sich Menschen selbst in komplexen Szenen?

Die DFKI-Forschenden haben dafür eine Datensammlung mit rund einer Million blickannotierter Frames aus einem realistischen 360-Grad-Fahrsimulator aufgebaut und zeigen damit, wie sich Fahreraufmerksamkeit unter kontrollierten, aber realitätsnahen Bedingungen erfassen lässt. Das ist nicht nur für die Analyse menschlicher Wahrnehmung relevant, sondern auch für Assistenzsysteme, die verstehen müssen, worauf FahrerInnen in kritischen Momenten tatsächlich achten. Damit wird ein blinder Fleck bisheriger Forschung sichtbar, denn klassische Datensätze erfassen vor allem frontale Sicht und unterschätzen seitliche oder rückwärtige Aufmerksamkeitsdynamiken, etwa beim Spurwechsel, beim Einfädeln oder beim Blick in den Rückraum.

DriverGaze360 ist somit mehr als ein Verkehrsthema. Die Arbeit zeigt, dass Szenenverständnis nicht nur eine Frage maschineller Segmentierung und Relationen ist, sondern auch eine Frage situierter, menschlicher Orientierung unter realistischen Bedingungen. Die zusätzliche Objektführung des Modells ist dabei zentral: Das System lernt nicht nur Blickverteilung, sondern auch, welche Objekte im jeweiligen Moment aufmerksamkeitstragend sind. Dadurch wird Aufmerksamkeit semantisch anschlussfähig und als strukturierte Beziehung zwischen Mensch und Umgebung beschreibbar.

Mit Inpaint360GS tritt schließlich ein dritter Zugang hinzu: das aktive Bearbeiten von Szenen. Das Verfahren ergänzt fehlende Bild- und Raumanteile in 360-Grad-Umgebungen objektbewusst und adressiert damit ein Kernproblem jeder realitätsnahen 3D-Repräsentation. Szenen sind oft unvollständig, verdeckt oder lückenhaft. Entscheidend ist dabei nicht nur, dass etwas rekonstruiert wird, sondern wie. Inpaint360GS zielt darauf, fehlende Bereiche so zu ergänzen, dass Objektlogik, räumliche

Struktur und semantische Kohärenz erhalten bleiben. Hierin liegt die übergeordnete wissenschaftliche Verbindung zu den Arbeiten ReLaGS und DriverGaze360. Wenn wir mit ReLaGS zeigen, wie Szenen lesbar werden, und mittels DriverGaze360 aufdecken, wie Orientierung in ihnen verteilt ist, dann zeigt Inpaint360GS, wie auf dieser Grundlage in Szenen eingegriffen werden kann, ohne ihre innere Struktur zu zerstören. Aus Wahrnehmung wird damit nicht automatisch Handlung, wohl aber die Voraussetzung für eine Form von Bearbeitbarkeit, die geometrische Präzision, semantisches Verständnis und Kontextsensibilität zusammendenkt.

Damit zeigen die drei Arbeiten mehr als methodische Fortschritte. Sie markieren einen Schritt hin zu KI-Systemen, die räumliche Umgebungen nicht nur rekonstruieren, sondern in ihrer Struktur, ihren Beziehungen und ihren Lücken verstehen und verarbeiten können. Wer Szenen, Aufmerksamkeit und Rekonstruktion gemeinsam denkt, schafft die Voraussetzungen für verlässliche Roboter, robustere Assistenzsysteme, präzisere digitale Zwillinge und eine visuelle KI, die sich nicht mit dem Sichtbaren begnügt, sondern den Zusammenhang mit-
—●

CVPR 2026: Bühne für visuelle KI

Die IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) zählt zu den wichtigsten internationalen Konferenzen der Computer-Vision-Forschung und fand 2026 vom 3. bis 7. Juni in Denver statt. Das DFKI diskutierte dort aktuelle Arbeiten zu 3D-Szenenverstehen, relationalem Denken, Aufmerksamkeitsmodellierung und objektbewusster Rekonstruktion. Im Fokus dieser Ausgabe stehen deshalb drei Beiträge aus dem DFKI-Forschungsbereich Erweiterte Realität: ReLaGS, DriverGaze360 und Inpaint360GS. Gemeinsam zeigen sie, wie sich Geometrie, Aufmerksamkeit und Szene-Editing als zusammenhängende Forschungsbewegung lesen lassen.



Die dargestellte Person ist fiktiv und KI-generiert.
Mikroskopie: Louisa Howard, Public Domain

Der Flaschenhals der Mikroskopie

Wie Foundation Models die Bildanalyse in den Lebenswissenschaften beschleunigen

Moderne Mikroskope erzeugen heute Bilddaten in einer Menge und Auflösung, die biologische Forschung grundlegend verändert haben. Der Engpass liegt deshalb immer seltener in der Aufnahme selbst, sondern in der Auswertung: Zellen, Zellkerne und Organellen müssen präzise abgegrenzt und über große Datensätze hinweg konsistent annotiert werden. Erst diese Segmentierung macht mikroskopische Bilder überhaupt zu belastbaren wissenschaftlichen Daten.

Genau hier setzt die in Nature Methods veröffentlichte Arbeit „Segment Anything for Microscopy“ an. Sie prüft, ob sich ein Foundation Model, das für Alltagsbilder entwickelt wurde, auf mikroskopische Daten übertragen lässt, ohne seine Flexibilität zu verlieren. Mit μ SAM zeigen die Forschenden, dass sich das Segment Anything Model für Licht- und Elektronenmikroskopie so weiterentwickeln lässt, dass biologische Strukturen zuverlässiger erfasst und zugleich interaktive wie automatische Segmentierungsaufgaben unterstützt werden.

Die wissenschaftliche Bedeutung dieser Arbeit liegt nicht allein in besseren Ergebnissen. Entscheidend ist, dass hier ein allgemeines KI-Bildmodell in eine Domäne übersetzt wird, deren Bildlogik sich grundlegend von der Alltagsbildwelt unterscheidet. Kontraste, Formen, Objektgrenzen und räumliche Beziehungen folgen in der Mikroskopie anderen Regeln; gerade deshalb ist die Frage wichtig, wie weit sich die Generalisierungsfähigkeit moderner KI tatsächlich tragen lässt.

« Mit μ SAM zeigen wir, dass sich Foundation Models auch in hochspezialisierten wissenschaftlichen Bildwelten produktiv einsetzen lassen, wenn man sie methodisch konsequent an die Domäne anpasst. Auf dieser Grundlage arbeiten wir nun an einer Ausgründung für die Analyse von Zellkulturen, die das Potenzial dieser Forschung in eine konkrete Anwendung überführen soll. »

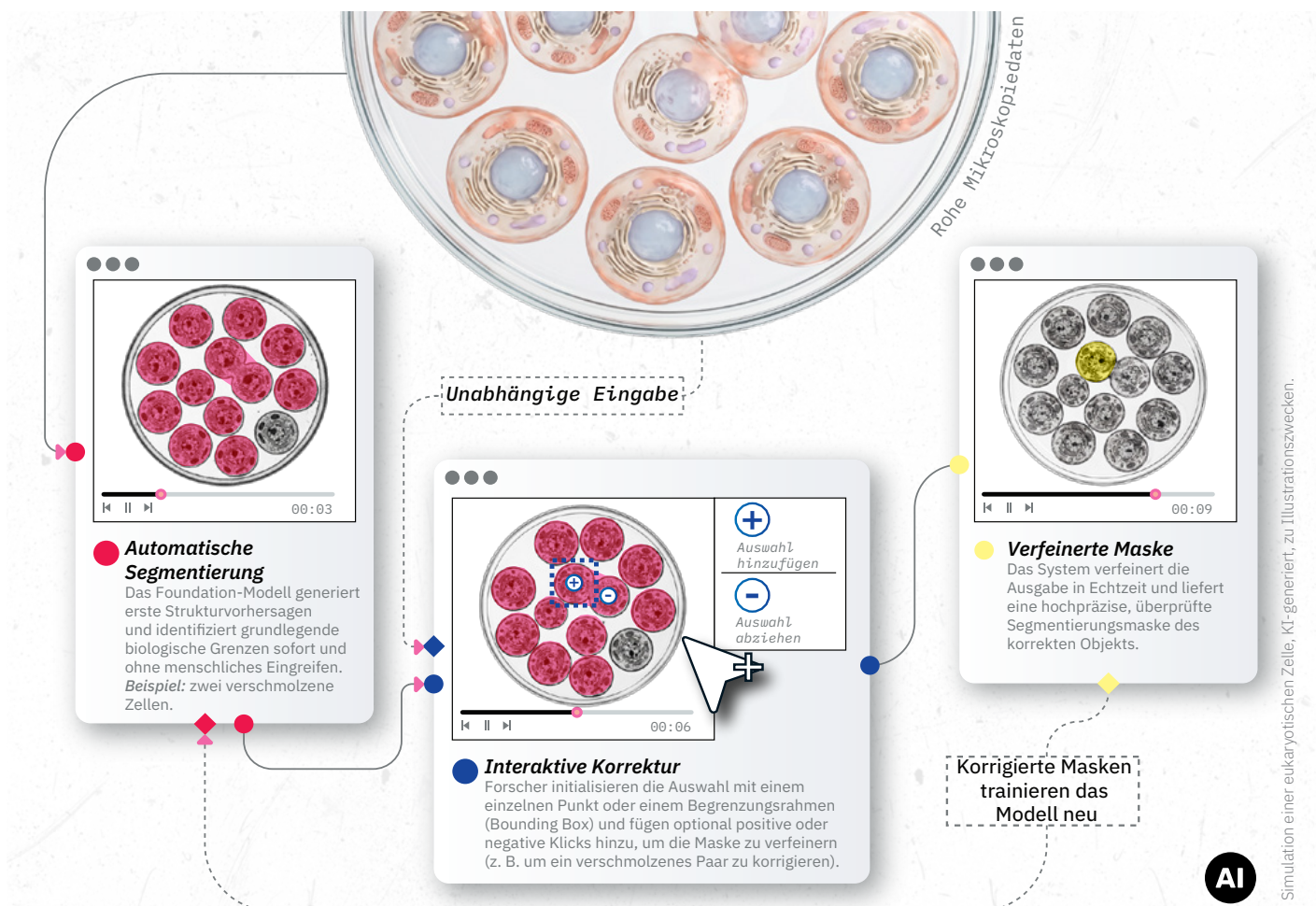
Prof. Dr. Andreas Dengel,
Leiter Forschungsbereich Smarte Daten &
Wissensdienste, Geschäftsführender Direktor
DFKI Kaiserslautern

Arbeitsprozess ein. Aus einer langsamen Routine wird so ein kooperativer Workflow, der Annotation beschleunigt, ohne die wissenschaftliche Kontrolle aufzugeben.

Für die Lebenswissenschaften ist das mehr als eine technische Verbesserung. Zellbiologie, Krebsforschung und Entwicklungsbiologie produzieren Datenmengen, die sich manuell nur mit erheblichem Zeitaufwand erschließen lassen. Die wissenschaftlichen Ergebnisse zeigen, wie sich dieser Flaschenhals entschärfen lässt: nicht durch den Ersatz wissenschaftlicher Arbeit, sondern durch eine präzisere Verbindung von Modell, Interface und Domänenwissen.

Die Arbeit ist damit auch ein Beitrag zur internationalen Debatte über Foundation Models. Sie zeigt, dass die Stärke solcher Systeme nicht in universeller Anwendbarkeit im Abstrakten liegt, sondern in der präzisen Übersetzung in eine konkrete wissenschaftliche Praxis. Genau dort entscheidet sich, ob ein Modell bloß beeindruckt oder in der wissenschaftlichen Praxis wirklich wirkt. —●

Besonders relevant wird der Ansatz dort, wo menschliche Expertise und Modellleistung ineinandergreifen. μ SAM ist als interaktive Arbeitsumgebung angelegt: Forschende setzen wenige Hinweise, das System ergänzt die Konturen und Korrekturen fließen in den weiteren





Wenn KI nur Englisch spricht, bleibt sie exklusiv

Wie künstliche Trainingsdaten helfen, Modelle für ressourcenarme Sprachen zu schaffen

Große Sprachmodelle funktionieren besonders gut in Sprachen, für die viele digitale Daten verfügbar sind. Für Englisch oder Deutsch existieren riesige Mengen an Texten, während zahlreiche andere Sprachen im Training moderner KI-Systeme kaum vertreten sind. Die Folge: Für zahlreiche Sprachgemeinschaften existieren bislang keine leistungsfähigen KI-Anwendungen, wodurch auch das wirtschaftliche und gesellschaftliche Potenzial moderner Sprachtechnologien weitgehend ungenutzt bleibt.

Kann künstlich erzeugtes Trainingsmaterial diese Lücke schließen? Dieser Frage ist ein Forschungsteam des DFKI, der Technischen Universität Brunn und des Kempelen Institute of Intelligent Technologies in Bratislava nachgegangen. Im Mittelpunkt stand dabei eine zentrale Herausforderung: Unter welchen Bedingungen können Large Language Models (LLMs) Trainingsdaten erzeugen, die qualitativ hochwertig genug sind, dass sie für das Training weiterer KI-Modelle tatsächlich geeignet sind.

Für die Studie wurden vier Open-Source-LLMs unterschiedlicher Größe eingesetzt, darunter Modelle der Familien Llama und Gemma. Sie erzeugten Trainingsdaten für elf typologisch unterschiedliche Sprachen, von ressourcenstarken Sprachen wie Englisch und Deutsch bis hin zu Sprachen mit deutlich geringerer digitaler Präsenz wie Walisisch, Rumänisch, Aserbaidschanisch, Slowenisch oder Telugu. Die Forschenden untersuchten

dabei drei zentrale Aufgaben der Sprachverarbeitung: das Erkennen von Nutzerabsichten, die Zuordnung von Texten zu Themen und die Analyse von positiven oder negativen Meinungen.

Um die Qualität der künstlich erzeugten Daten zu überprüfen, trainierte das Team damit neue KI-Modelle. Deren Ergebnisse wurden anschließend mit Modellen verglichen, die mit echten, von Menschen annotierten Goldstandard-Daten trainiert wurden. So ließ sich präzise bestimmen, wie belastbar die künstlich erzeugten Datensätze tatsächlich sind.

Mit dieser Strategie schrumpfte der Abstand zu Modellen, die auf echten Trainingsdaten basieren, deutlich. Besonders beim Erkennen von Nutzerabsichten erreichten die Modelle nahezu das Niveau der Goldstandard-Daten.

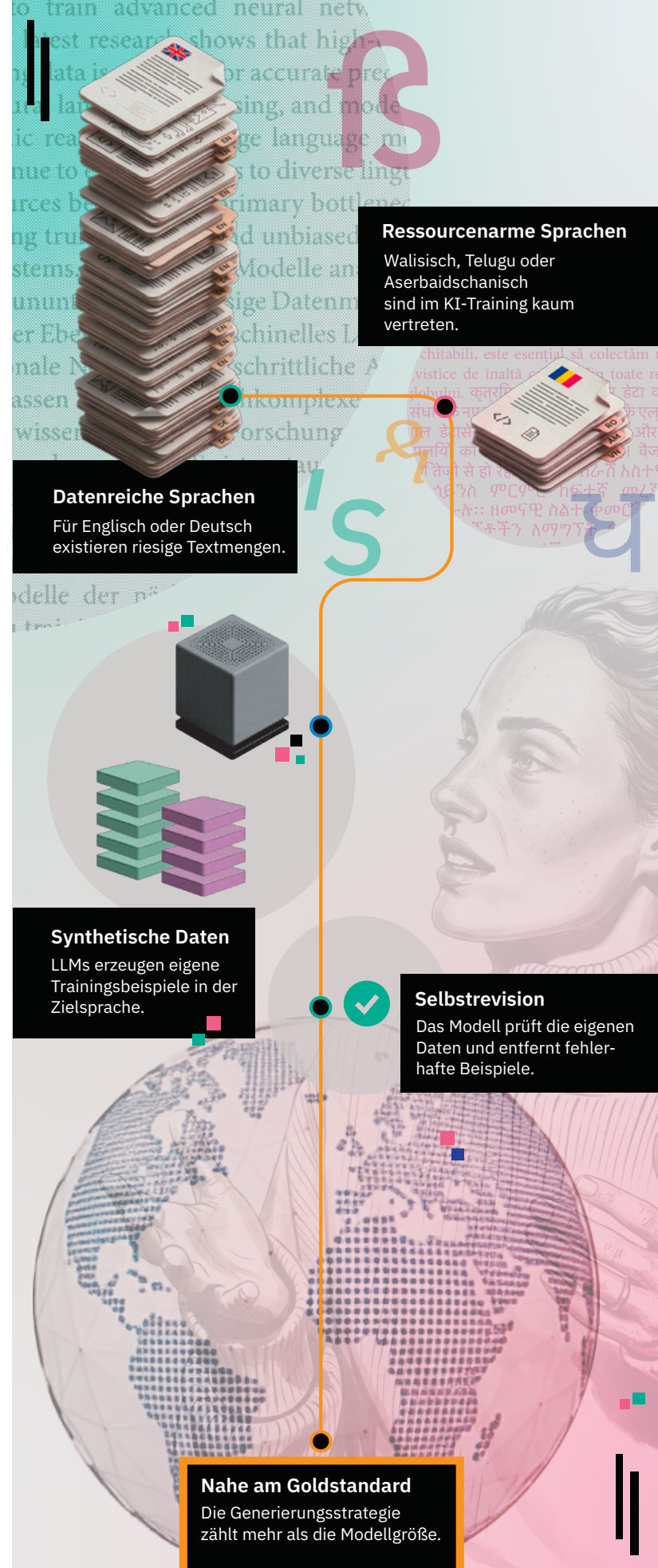
« Oft wird angenommen, dass mehr Parameter automatisch bessere Ergebnisse liefern. Unsere Studie zeigt jedoch, dass die Generierungsstrategie mindestens genauso wichtig ist. Mit den richtigen Prompting- und Revisionsverfahren lassen sich auch für ressourcenarme Sprachen hochwertige Trainingsdaten erzeugen. Das ist eine wichtige Voraussetzung, damit moderne KI nicht nur für wenige dominante Sprachen funktioniert. »

Tatiana Anikina, Erstautorin der Studie und Wissenschaftlerin im Forschungsbereich Sprachtechnologie und Multilingualität am DFKI

Gleichzeitig zeigte sich, dass intelligente Prompting- und Revisionsverfahren einen Teil des Vorteils großer Modelle kompensieren können. Damit eröffnen sich praxisnahe Möglichkeiten, hochwertige Daten auch mit kleineren und kostengünstigeren Open-Source-Modellen zu erzeugen.

Die Ergebnisse zeigen, dass nicht allein die Größe eines Sprachmodells über die Qualität der erzeugten Daten entscheidet. Ausschlaggebend ist vielmehr die Generierungsstrategie. Am erfolgreichsten erwies sich eine Kombination aus Beispielen in der jeweiligen Zielsprache und einer anschließenden automatischen Qualitätskontrolle. In einem zusätzlichen Prüfschritt kontrolliert die KI ihre eigenen Ergebnisse und entfernt fehlerhafte oder unpassende Beispiele.

Trotz der vielversprechenden Ergebnisse erreichen synthetische Daten nicht in allen Sprachen und Aufgaben die gleiche Qualität. Besonders bei stark unterrepräsentierten Sprachen bleiben deutliche Leistungsunterschiede bestehen. Die Studie macht jedoch deutlich, dass diese Lücke kleiner wird, wenn Datengenerierung systematisch evaluiert und mit geeigneten Qualitätsmechanismen kombiniert wird.





Das Ende der „Alleskönner“

Wie KI jetzt maßgeschneiderte High-Speed-Datenbanken baut

Haben Sie sich schon einmal gefragt, warum Software oft so schwerfällig ist? Der Grund liegt meist in einem alten Dogma der Software-Entwicklung: „One size fits all“ (eine Lösung für alle).

Traditionelle Datenbanken wie PostgreSQL oder DuckDB sind wie Schweizer Taschenmesser: sie können nahezu jede Abfrage beantworten. Doch diese Flexibilität kostet Leistung. Für Funktionen, die Sie in Ihrem speziellen Fall gar nicht nutzen, zahlen Sie eine unsichtbare „Flexibilitäts-Steuer“ in Form von Rechenzeit und Speicherauslastung. Allzweck-Analysedatenbanken sind darauf ausgelegt, beliebige Abfragen über beliebige Daten und Schemata auszuführen. Diese Vielseitigkeit ist ihr Vorteil, aber auch ihr Nachteil. Laufzeitinterpretation, Indirektionsschichten und Abstraktionsgrenzen verursachen zusätzlichen Overhead – selbst in optimierten Engines.

Die meisten realen Systeme führen jedoch nicht beliebige Abfragen aus, sondern eine bekannte Menge von SQL-Vorlagen über ein bekanntes Schema, immer und immer wieder. Ein Forschungsteam der TU Darmstadt (Johannes Wehrstein, Timo Eckmann, Matthias Jasny und Carsten Binnig) bricht mit diesem Prinzip. Ihr revolutionärer Ansatz heißt „Bespoke OLAP“ (Maßgeschneiderte Datenanalyse) [2603.02001] – und er zeigt, dass die Zukunft der Software-Entwicklung radikal minimalistisch ist.

Der Trick:

Ein digitaler Maßanzug statt Konfektionsware

Stellen Sie sich vor, Sie müssten nicht mehr eine riesige, universelle Datenbank installieren. Stattdessen analysiert eine künstliche Intelligenz (wie Claude oder GPT) exakt die Daten und die Suchanfragen, die Ihr Unternehmen tatsächlich nutzt.

Anschließend schreibt die KI eine autonome Mini-Datenbank in C++ – und zwar exklusiv für diese eine Aufgabe. Da diese Mini-Datenbank keine Rücksicht auf andere Abfragetypen nehmen muss, fliegt jeglicher Ballast über Bord. Datentypen, Speicherlayouts und Rechenwege werden fest in den Code gegossen.

Bis zu 60-mal schneller als die Industrie-Konkurrenz

Die Ergebnisse der ersten wissenschaftlichen Tests (unter Nutzung des standardisierten TPC-H-Benchmarks) sind ein Paukenschlag in der Tech-Welt: Die von der KI generierten „One-Size-Fits-One“-Datenbanken hängten etablierte und von menschlichen Experten jahrelang optimierte Systeme wie DuckDB, Clickhouse oder Umbra mühelos ab. Die maßgeschneiderten KI-Engines waren 10- bis 60-mal schneller [2603.02001].

Die maßgeschneiderten
KI-Engines waren
10- bis 60-mal schneller

C++



Wo bringt das in der Praxis den Turbo-Boost?

Diese Technologie ist kein reines Theorie-Projekt. Überall dort, wo Unternehmen jede Nacht dieselben riesigen Datenmengen durchrechnen müssen, verändert „Bespoke OLAP“ das Spiel komplett:

E-Commerce & Personalisierung:

Ein Online-Riese wertet jede Nacht das Klick- und Kaufverhalten von Millionen Kundinnen und Kunden aus, um die Startseite für den nächsten Tag zu personalisieren. Statt ein Standard-System stundenlang unter Volllast laufen zu lassen, baut die KI in Minuten eine Mini-Engine, die exakt diese eine nächtliche Analyse in einem Bruchteil der Zeit erledigt.

Finanzsektor & Risikoanalyse:

Banken müssen täglich regulatorische Risiko- und Betrugsberichte erstellen. Die mathematischen Abfragen hierfür verändern sich kaum. Eine maßgeschneiderte KI-Datenbank kann diese standardisierten Berechnungen in Echtzeit durchführen, statt die IT-Infrastruktur über Stunden lahmzulegen.

IoT & Industrie 4.0:

In modernen Fabriken erfassen Tausende Sensoren minütlich Maschinendaten zur vorausschauenden Wartung (Predictive Maintenance). Eine winzige, von der KI generierte Datenbank passt direkt auf die Steuerungsgeräte vor Ort und wertet die Daten verzögerungsfrei aus, ohne sie erst in eine Cloud schicken zu müssen.

Wo steht das Projekt aktuell?

Die Arbeit schlägt bereits hohe Wellen: die wissenschaftlichen Paper zu Bespoke OLAP wurden nach Peer-Review Verfahren im IEEE Data Engineering Bulletin (3/2026), eines der international wichtigsten Journals sowie für die VLDB (9/2026, Boston), eine der weltweit wichtigsten Datenbank-Konferenzen angenommen. Darüber hinaus wurde das Projekt im April 2026 auf dem renommierten Sky Computing Blog der UC Berkeley vorgestellt, zur invite-only HPTS Workshop eingeladen und in Gastvorträgen unter anderem bei Google und SAP diskutiert.

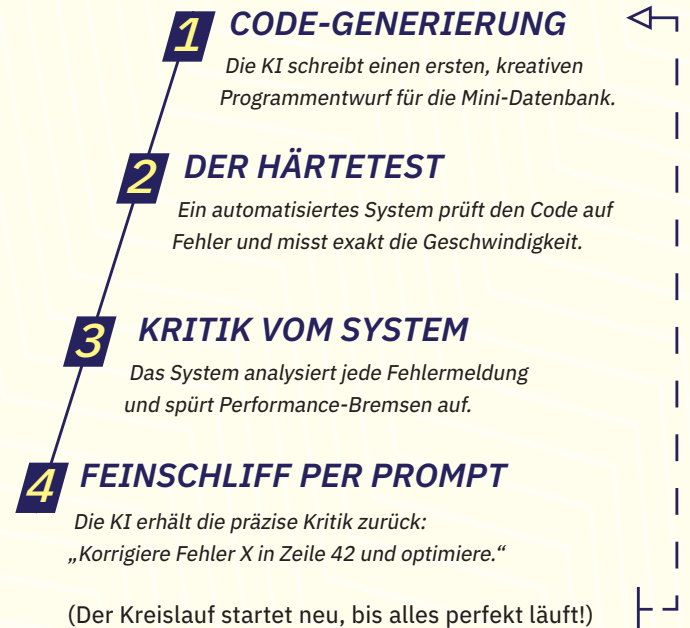
Das Fazit für die Zukunft

„Bespoke OLAP“ zeigt uns eine neue Ära der Software-Architektur. KI wird in Zukunft nicht mehr nur Programmierern beim Tippen assistieren. Sie wird die Rolle des Chef-Architekten übernehmen, der in Sekundenschnelle maßgeschneiderte Digitalsysteme erschafft, von denen menschliche Entwickler bisher nur träumen konnten.

Der Clou:

Die KI im Dauertest (Die Feedback-Schleife)

Normalerweise scheitern KI-Modelle daran, fehlerfreien und extrem schnellen Systemcode für komplexe Datenstrukturen zu schreiben. Die Forscher haben dieses Problem mit einer genialen, geschlossenen Feedback-Schleife gelöst, die völlig ohne menschliches Zutun funktioniert.



Wie geht es weiter?

Mit dem Spin-off SynnoDB (<https://synnodb.com/>) möchte das Team um Prof. Carsten Binnig das Potenzial der Technik ausschöpfen und die innovative Bespoke-OLAP-Technologie auf den Markt bringen. Normale Datenbanken sind Allzweckautos. SynnoDB baut stattdessen für jede Kundin und jeden Kunden ein maßgeschneidertes Rennfahrzeug, das exakt auf die benötigten Daten und Abfragen abgestimmt ist und dadurch deutlich schneller arbeitet. Aktuell gibt es auf der Webseite des Unternehmens eine Try-yourself-Demo, die erlaubt, selbst kleine Minidatenbanken zu synthetisieren, sowie einen Live-Playground, in dem man:

- Abfragen auf einer echten von SynnoDB erzeugten Engine ausführen kann.
- Die Geschwindigkeit mit DuckDB und Umbra vergleichen kann.
- Den erzeugten C++-Quellcode ansehen kann.
- Nachverfolgen kann, wie die Engine aufgebaut wurde. —●



X-Hacking

Wenn Erklärbarkeit zur Täuschung wird

DFKI-Forschende decken mit ihrem Paper auf, wie erklärbare KI zur methodischen Falle werden kann: Wenn AutoML aus vielen ähnlich guten Modellen genau jene Variante auswählt, deren Erklärung ins gewünschte Narrativ passt, wird Transparenz selbst zum Prüfstein.

XAI-Verfahren sollen maschinelle Lernsysteme nachvollziehbarer machen. Doch genau dort, wo Modelle ihre Entscheidungen plausibel erklären, beginnt ein methodisches Problem, das im Paper als X-Hacking beschrieben wird: Erklärungen lassen sich gezielt in eine gewünschte Richtung verschieben, obwohl die Modelle in ihrer Vorhersageleistung nahezu gleich gut sind.

Der Ausgangspunkt ist eine einfache, aber folgenreiche Beobachtung. In vielen Datensätzen existieren mehrere Modelle, die ähnlich präzise arbeiten, ihre Ergebnisse aber sehr unterschiedlich begründen. Wer aus dieser Menge „vertretbarer“ Modelle gezielt auswählt, kann genau jene Erklärung bevorzugen, die ein gewünschtes Narrativ stützt. Das Paper überträgt damit das aus der Statistik bekannte Prinzip des p-Hacking auf erklärbare KI: Nicht nur die Vorhersage wird optimiert, sondern auch die Begründung des Resultats.

Besonders relevant wird das durch AutoML. Automatisierte Such- und Optimierungsverfahren erleichtern die Entwicklung von Modellen enorm, vergrößern aber zugleich den Raum möglicher Entscheidungen. Darin liegt die wissenschaftliche Pointe von X-Hacking: Wenn viele Modelle, Vorverarbeitungsschritte und Parametereinstellungen zur Verfügung stellen, kann die Suche nach

ICML: Wo maschinelles Lernen auf Anwendung trifft

Die ICML gilt als eine der drei weltweit führenden Konferenzen für maschinelles Lernen (neben NeurIPS und ICLR). Ihr Profil ist geprägt von theoretischer Analyse, algorithmischer Innovation und statistischem Lernen – der mit starken Schnittstellen zu Reinforcement Learning, Robotik und Optimierungstheorie. Die Konferenz ist ein Ort des Austauschs für Forschung aus KI, Statistik und Data Science sowie anwendungsnahe Felder wie Computer Vision, Computational Biology, Speech und Robotics. Auf der ICML 2025 präsentierte das DFKI-Team um Prof. Sebastian Vollmer das Paper „X-Hacking: The Threat of Misguided AutoML“. Es überträgt das Prinzip des p-Hacking auf erklärbare KI und zeigt, wie AutoML Pipelines systematisch nach Modellen mit gewünschten, aber potenziell irreführenden Erklärungen durchsucht werden können.



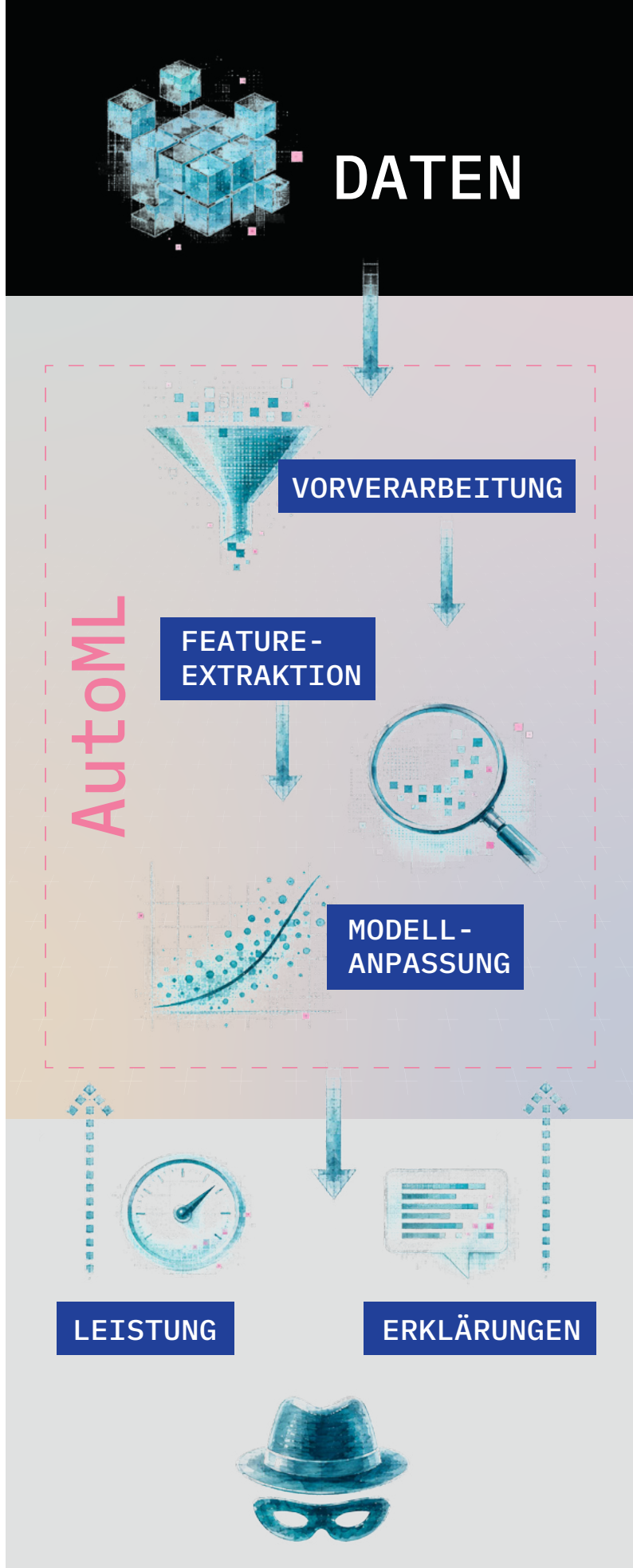
« In einer Zeit, in der KI Entscheidungen erklärt, aber nicht immer versteht, müssen wir als Wissenschaft Verantwortung für die Tiefe dieser Erklärungen übernehmen – und für ihre Grenzen. Ziel ist eine Wissenskultur, die nicht nur auf Genauigkeit, sondern auch auf Ehrlichkeit in der Erklärbarkeit setzt. »

Prof. Dr. Sebastian Vollmer,
Leiter Forschungsbereich Data Science
und Ihre Anwendungen am DFKI

einer „passenden“ Erklärung systematisch Teil des Analyseprozesses werden. Erklärbarkeit erscheint dann nicht mehr automatisch als Korrektiv, sondern selbst als etwas, das methodisch geprüft werden muss.

Für sensible Anwendungsfelder ist das von erheblicher Bedeutung. In Medizin, Sozialforschung oder Justiz gelten erklärbare Modelle oft als Voraussetzung für Vertrauen, Transparenz und verantwortliche Entscheidungen. Wenn sich aber auch plausible Erklärungen aus einem Spektrum ähnlich guter Modelle selektiv herausgreifen lassen, reicht der Verweis auf Interpretierbarkeit allein nicht mehr aus. An diesem Spannungsverhältnis setzt der Beitrag zur Debatte um Trustworthy AI an: Nicht nur die Genauigkeit eines Modells zählt, sondern die Integrität des gesamten Weges, auf dem Erklärung überhaupt zustande kommt.

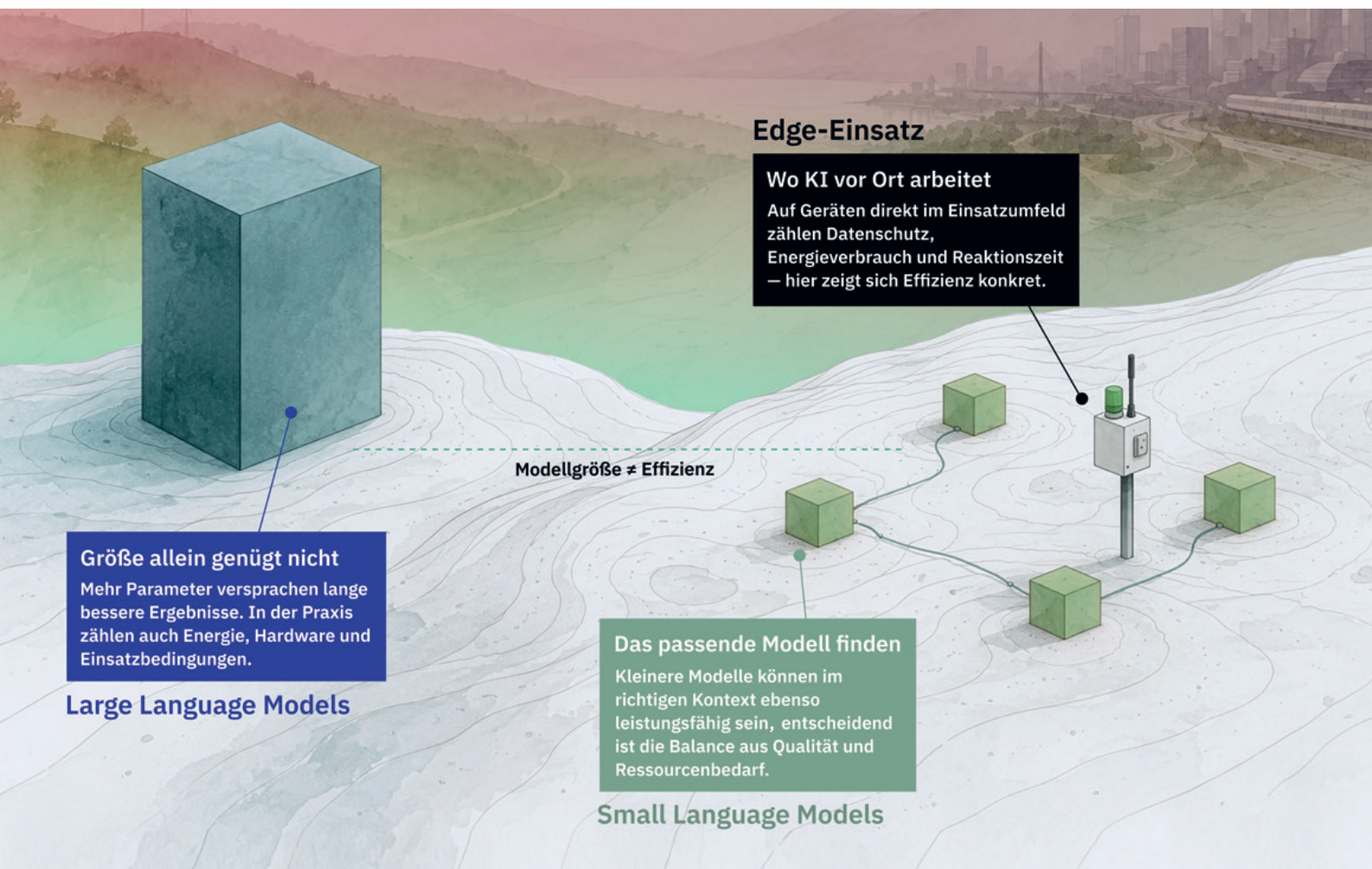
Das Paper belässt es deshalb nicht bei der Problembeschreibung. Es plädiert für eine reflektiertere wissenschaftliche Praxis: für Erklärungshistogramme, die die Spannweite möglicher Erklärungen sichtbar machen, für eine vollständige Dokumentation der Pipeline und für ein stärkeres Bewusstsein dafür, dass automatisierte Modellauswahl nicht nur Effizienz schafft, sondern auch neue Verzerrungen hervorbringen kann. X-Hacking markiert damit keinen Randaspekt erklärbarer KI, sondern eine grundlegende Frage wissenschaftlicher Redlichkeit im Umgang mit automatisierten Modellsuchen. —●



Effizienz als neuer Maßstab für Sprachmodelle

Warum die Größe eines Sprachmodells allein nicht entscheidend ist

Die Entwicklung großer Sprachmodelle folgte lange einer scheinbar einfachen Logik: Mehr Parameter, mehr Rechenleistung und größere Modelle versprechen bessere Ergebnisse. Dieses Skalierungsparadigma prägte die KI-Forschung der vergangenen Jahre. Doch je stärker Sprachmodelle ihren Weg in Unternehmen und gesellschaftliche Anwendungen finden, desto deutlicher zeigt sich: Entscheidend ist nicht mehr allein die Größe, sondern die Praxistauglichkeit.



« Wir schaffen wissenschaftliche Grundlagen für den effizienten Einsatz von Sprachmodellen. Die Ergebnisse stellen wir auf der IJCAI 2026 im Beitrag „Mapping the Efficiency Landscape of Small Language Models“ vor. Sie zeigen, dass nicht die Größe eines Modells entscheidend ist, sondern wie effizient es die Anforderungen eines konkreten Anwendungskontexts erfüllt. »

Fabian Reichwald, Forschungsbereich
Intelligente Märkte am DFKI

Warum größer nicht automatisch besser ist

Ob in industriellen Prozessen, auf Edge-Geräten oder in datenschutzsensiblen Umgebungen: Sprachmodelle müssen unter realen Bedingungen und Anforderungen funktionieren. Für ein mittelständisches Unternehmen, das beispielsweise technische Dokumente zusammenfasst oder Kundenanfragen automatisiert bearbeitet, ist daher nicht zwangsläufig das größte verfügbare Modell die beste Wahl. Entscheidend ist vielmehr, welches Modell die Aufgabe auf der vorhandenen Hardware mit der besten Balance aus Ergebnisqualität, Energieverbrauch und Reaktionszeit erfüllt.

Damit verschiebt sich auch der wissenschaftliche Fokus. Statt Modelle nur nach Größe oder Leistungsfähigkeit zu bewerten, rückt ihre praktische Einsetzbarkeit in den Mittelpunkt. Wie viel Rechenleistung benötigt eine Aufgabe? Wann rechtfertigt ein höherer Energieeinsatz einen messbaren Qualitätsgewinn und wann steigen lediglich Kosten und Ressourcenverbrauch?

Mehr als nur Leistung

Effizienz ist damit keine bloße Nebenbedingung mehr, sondern eine eigenständige wissenschaftliche Größe. Soll ein Unternehmen ein Sprachmodell lokal betreiben, reicht der Blick auf die Modellgröße nicht aus. Entscheidend ist, welches Modell eine konkrete Aufgabe mit möglichst geringem Ressourcenaufwand löst. Erst das Zusammenspiel von Aufgabenleistung, Energieverbrauch und Anwendungskontext macht Effizienz messbar.

Viele verbreitete Annahmen halten dieser Betrachtung nicht stand. Kleinere Modelle arbeiten nicht automatisch effizienter als größere. Zusätzliche Parameter oder höherer Energieeinsatz führen ebenfalls nicht zwangsläufig zu besseren Ergebnissen. Erst die systematische Betrachtung aller Einflussgrößen ermöglicht fundierte Entscheidungen.

Transparenz für fundierte Entscheidungen

Die nächste Entwicklungsstufe der Sprachmodelle besteht deshalb nicht allein darin, leistungsfähigere Modelle zu entwickeln. Ebenso wichtig sind wissenschaftlich fundierte Kriterien, mit denen sich für unterschiedliche Aufgaben das jeweils geeignete Modell auswählen lässt. Nicht die Größe eines Modells allein entscheidet über seinen praktischen Nutzen, sondern seine Eignung für den jeweiligen Anwendungskontext. Effizienz wird damit zu einer wichtigen Orientierung bei der Auswahl von Sprachmodellen, insbesondere dort, wo technische, wirtschaftliche und ökologische Anforderungen zusammenkommen. —●

Zwischen generativer KI und Robotik

Vom Aufgabenverständnis zum verlässlichen Roboterhandeln

Während Sprachmodelle komplexe Zusammenhänge verstehen und plausible Antworten erzeugen können, müssen Robotersysteme zusätzlich mit Kräften, Dynamiken und Unsicherheiten umgehen. Die neue Architekturebene macht KI-generierte Aufgabenrepräsentationen für die Robotik nutzbar und trägt so dazu bei, die Ausführung von Aufgaben in realen Szenarien verlässlicher zu machen. Gleichzeitig soll der Ansatz hochdynamischen Robotersystemen eine flexiblere Anpassung an wechselnde Umgebungen ermöglichen.

Im Zentrum des Ansatzes steht eine neue Form der Aufgabenrepräsentation. Anstatt eine Anweisung wie „Bewege ein Objekt an einen bestimmten Ort“ direkt in Steuerbefehle umzusetzen, analysiert das KI-Modell zunächst die Aufgabe. Dabei erfasst es die Zielsetzung, relevante Bedingungen sowie die Eigenschaften des Roboters und seiner Umgebung und erzeugt daraus eine semantische Aufgabenrepräsentation.

Diese Repräsentation wird anschließend über eine speziell entwickelte domänenspezifische Sprache (Domain-Specific Language, DSL) in eine maschinenlesbare Form überführt.

Die DSL bildet die Schnittstelle zwischen generativer KI und mathematischen Optimierungsverfahren der Robotik. Sie beschreibt die Anforderungen an eine Be-



Der Unitree G1 ist eine der Forschungsplattformen, die im ActGPT-Projekt eingesetzt werden.

« Die eigentliche Herausforderung besteht nicht darin, Robotern beizubringen, Anweisungen zu verstehen, sondern sicherzustellen, dass sie diese unter physikalischen Randbedingungen zuverlässig und robust ausführen können. Generative KI hilft dabei, den Kontext natürlicher Sprache zu verstehen und das Problem zu formulieren, während die Physik entscheidet, ob die Lösung tatsächlich funktioniert. »

Shubham Vyas,
Robotics Innovation Center am DFKI

wegung, aus denen sich ein Optimierungsproblem (Optimal-Control-Problem, OCP) ableiten lässt. Dieses kann mit etablierten Steuerungsverfahren gelöst werden, um Bewegungen zu erzeugen, die zur Aufgabe passen und gleichzeitig physikalisch ausführbar sind.

Die Arbeiten der Forschenden stehen exemplarisch für eine breitere Entwicklung in der KI-Forschung: Entscheidend ist nicht mehr allein, wie Maschinen ihre Umwelt verstehen und modellieren können, sondern wie zuverlässig sie dieses Wissen unter realen Bedingungen in Handlungen überführen. —●

Forschungsprojekt ActGPT

Im vom Bundesministerium für Forschung, Technologie und Raumfahrt (BMFTR) geförderten Projekt „ActGPT“ entwickelt das DFKI Robotics Innovation Center Methoden, um generative KI mit robusten Robotikverfahren zu verknüpfen und KI-generierte Aufgabenrepräsentationen für reale Anwendungsszenarien nutzbar zu machen. Ziel ist es, die verlässliche Ausführung komplexer Roboteraufgaben zu unterstützen und die Entwicklung von Steuerungsstrategien zu vereinfachen.

Der ActGPT Workflow

1 Input



Sprache



Bilder



Sensordaten

2 Generative KI

Versteht die Aufgabe

Neue Architekturschicht

Domänenspezifische Sprache (DSL)

Beschreibt:



Ziele



Kosten-
funktionen



Rand-
bedingungen

Optimalsteuerungsproblem (OCP)

3 Optimalsteuerung

Löst das OCP

Berechnet optimale Trajektorien

4 Roboter

Führt die Bewegung aus

ACT GPT

Die neue Architekturschicht. KI-Modelle übersetzen Aufgaben in strukturierte physikalische Probleme und verbinden so semantisches Wissen direkt mit robuster Robotik.



Die Zukunft der KI-basierten Robotik entsteht im Zusammenspiel

Wie das DFKI wissenschaftliche Exzellenz, Systemkompetenz und reale Erprobung verbindet

Die nächste Entwicklungsstufe der Künstlichen Intelligenz entscheidet sich in der realen Welt: in Robotern, die ihre Umgebung verstehen, aus Erfahrungen lernen und sicher handeln können. Diese Verbindung von KI und physischem Handeln wird unter dem Begriff Embodied AI (verkörperte Künstliche Intelligenz) erforscht und zählt zu den zentralen wissenschaftlichen Herausforderungen der KI-Forschung.

Mit seiner Arbeit an der Schnittstelle von Künstlicher Intelligenz, autonomen Systemen und Robotik leistet das DFKI einen wichtigen Beitrag zur Weiterentwicklung KI-basierter Robotik in Deutschland. Seine besondere Stärke liegt in der Verbindung von wissenschaftlicher Grundlagenforschung, technologischer Systemkompetenz und realitätsnaher Erprobung.

Von KI-Methoden zu intelligenten Robotersystemen

Über verschiedene Forschungsbereiche und Standorte hinweg untersucht das DFKI zentrale Fragestellungen intelligenter Robotik – von Wahrnehmung und Lernen über Planung und autonome Entscheidungsprozesse bis hin zur Mensch-Roboter-Interaktion und praktischen Anwendungsszenarien.

Die Arbeiten umfassen sowohl einzelne KI-Methoden und technologische Komponenten als auch die Entwicklung und Erprobung vollständiger Robotersysteme. Ent-

scheidend ist das koordinierte Zusammenwirken unterschiedlicher technologischer Ebenen: Software, Sensorik, Steuerung und physische Umsetzung müssen zu einem funktionierenden Gesamtsystem integriert werden. Erst dadurch entstehen autonome Roboter, die in dynamischen Umgebungen sicher und flexibel agieren können.

Ganzheitlicher Ansatz und Systemkompetenz

Besonders sichtbar wird dieser Ansatz am Robotics Innovation Center in Bremen. Hier werden komplexe Robotersysteme entwickelt und unter realitätsnahen Bedingungen erprobt. Die entwickelten Forschungsplattformen reichen von autonomen Unterwasserrobotern und Weltraumrovern über Laufroboter und humanoide Systeme bis hin zu robotischen Exoskeletten sowie einzelnen Komponenten wie Manipulatoren und Greifern. Sie ermöglichen die Untersuchung zentraler Herausforderungen intelligenter Robotik – von autonomer Navigation bis zur sicheren Interaktion mit Menschen.

« Robotik ist ein Feld, in dem viele Disziplinen zusammenkommen müssen. Wir betrachten alle Ebenen eines Robotersystems: von Hardware- und Elektronikentwicklung über Software bis hin zu neuen KI-Methoden. Diese Kombination findet man in dieser Breite nicht an vielen Einrichtungen. »

Dr. Sirko Straube, stellvertretender Leiter des Robotics Innovation Center am DFKI

Die Bedeutung von Forschungsplattformen und Testinfrastrukturen

Forschungsplattformen sind ein wichtiger Bestandteil der Robotikforschung. Eigene Systeme ermöglichen es, neue Konzepte und Technologien auf Systemebene zu untersuchen und das Zusammenspiel verschiedener Komponenten zu erforschen. Kommerzielle Plattformen ergänzen diesen Ansatz, indem sie reproduzierbare Experimente unter standardisierten Bedingungen ermöglichen und so die Vergleichbarkeit von Forschungsergebnissen verbessern.

Ebenso wichtig sind spezialisierte Testumgebungen, in denen neue Methoden und Systeme unter realitätsnahen Bedingungen erprobt werden können. Mit der Maritimen Explorationshalle und der Multifunktionshalle in Bremen, die unter anderem eine künstliche Mondkraterlandschaft umfasst, sowie der SmartFactory in Kaisers-

lautern verfügt das DFKI über einzigartige Infrastrukturen für unterschiedliche Bereiche der Robotikforschung. Diese ermöglichen die systematische Validierung neuer Ansätze und die Bewertung ihrer Leistungsfähigkeit in realistischen Testszenarien.

Gemeinsam die Robotikforschung voranbringen

Die Weiterentwicklung intelligenter Robotik erfordert das Zusammenspiel unterschiedlicher wissenschaftlicher Disziplinen und technologischer Kompetenzen, auch über Fach- und Institutionsgrenzen hinweg. Deshalb engagiert sich das DFKI in nationalen und internationalen Forschungsnetzwerken.

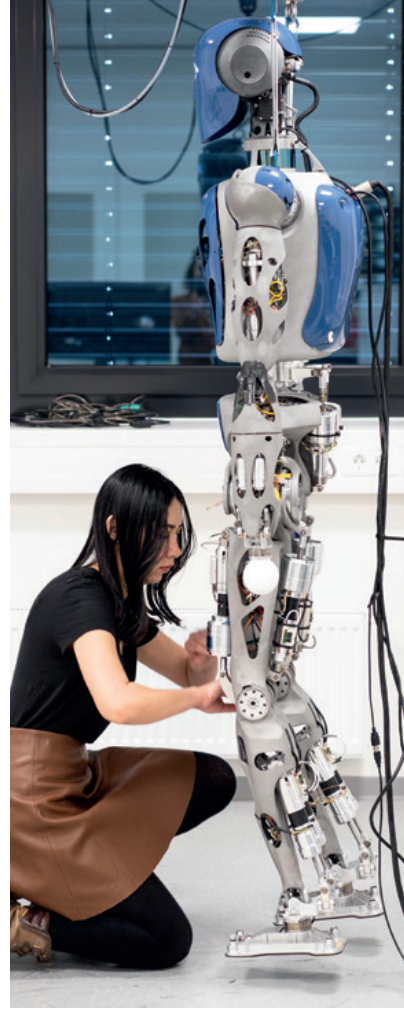
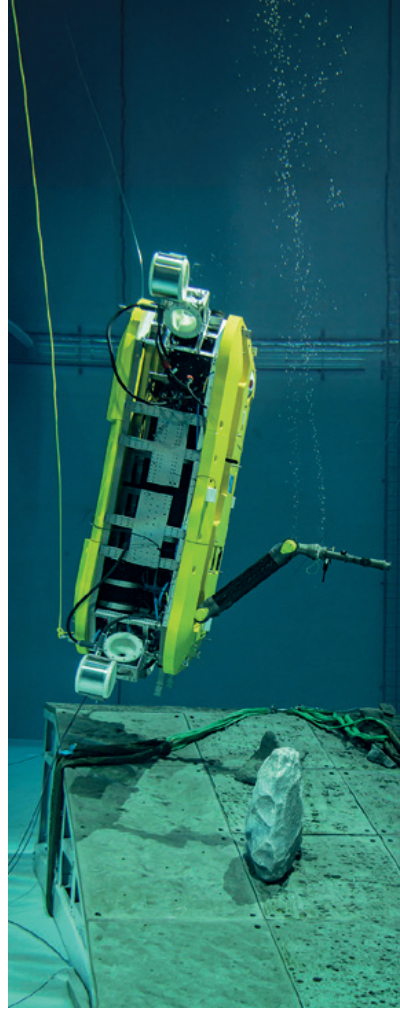
Ein zentrales Beispiel für die Zusammenarbeit ist das Robotics Institute Germany (RIG). Hier bringt das DFKI seine Expertise in autonomer Robotik und Künstlicher Intelligenz sowie seine Erfahrungen aus der Entwicklung und Erprobung komplexer Systeme ein. Die Vernetzung mit führenden Robotikstandorten in Deutschland bündelt unterschiedliche Kompetenzen und schafft einen Rahmen für gemeinsame Forschungsaktivitäten.

Mit der Verbindung von wissenschaftlicher Exzellenz, technologischer Systemkompetenz und interdisziplinärer Zusammenarbeit leistet das DFKI wichtige Beiträge zur Erforschung der zentralen Fragestellungen der KI-basierten Robotik und trägt damit zur Gestaltung der nächsten Generation intelligenter Robotersysteme bei. —●



Das DFKI als Partner im Robotics Institute Germany (RIG)

Das Robotics Institute Germany (RIG) ist eine vom Bundesministerium für Forschung, Technologie und Raumfahrt (BMFTR) geförderte Initiative zur Vernetzung führender Robotikstandorte in Deutschland. Ziel ist es, wissenschaftliche Exzellenz, internationale Sichtbarkeit, Talentförderung und Technologietransfer in der Robotik gemeinsam zu stärken. Das DFKI beteiligt sich mit drei Forschungsbereichen am RIG: dem Robotics Innovation Center in Bremen sowie den Bereichen Intelligente Netze und Innovative Fabriksysteme in Kaiserslautern. Die Beiträge des DFKI umfassen die Bereitstellung und Nutzung von Forschungsinfrastrukturen, Angebote zur Aus- und Weiterbildung von Fach- und Führungskräften, die Beteiligung an internationalen Robotik-Wettbewerben und Challenges sowie die Vernetzung mit Industriepartnern. —●



DFKI @ IJCAI-ECAI 2026

Perspektiven auf die KI von morgen

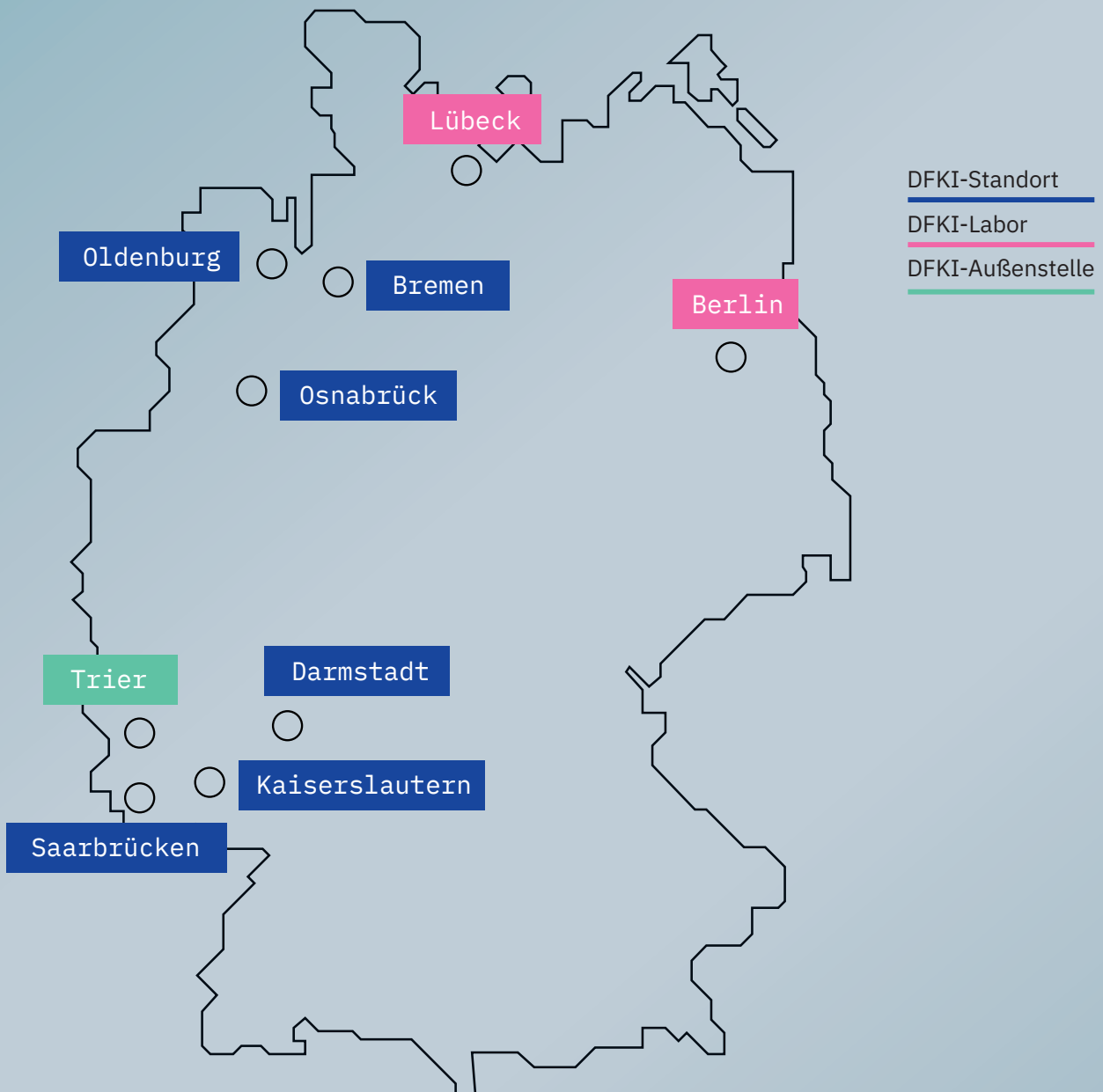
Mit zwölf angenommenen Paper-Beiträgen auf der International Joint Conference on Artificial Intelligence (IJCAI) 2026 ist das DFKI auf einer der weltweit renommiertesten Konferenzen der Künstlichen Intelligenz vertreten. Die Arbeiten spiegeln die wissenschaftliche Vielfalt des Forschungszentrums wider – von Foundation Models und Computer Vision über Robotik, semantische Technologien und erklärbare KI bis hin zu medizinischen Anwendungen, Umweltmonitoring und Software Engineering.

Gemeinsam machen die Beiträge sichtbar, welche Fragestellungen die internationale KI-Forschung derzeit prägen: Wie lassen sich intelligente Systeme effizienter, nachvollziehbarer und interaktiver gestalten? Wie können sie komplexe Zusammenhänge besser verstehen und Menschen in unterschiedlichsten Anwendungsfeldern zuverlässig unterstützen? Damit geben die DFKI-Paper einen Einblick in die nächste Generation intelligenter KI-Systeme.

Für mehr Informationen QR-Code scannen

www.dfki.de





Impressum

dfki ai next; Ausgabe August | 2026; **Herausgeber:** Deutsches Forschungszentrum für Künstliche Intelligenz GmbH (DFKI); Trippstadter Str. 122, 67663 Kaiserslautern; **Tel.:** +49 631 20575 0; **E-Mail:** news@dfki.de; **Redaktion:** Jeremy Gob (verantwortlich), Heike Leonhard, Udo Urban, Andrea Fink, Maria Santos; **Layout:** Lando Lehmann, Juan Mejia; **Satz:** One Vision Design, Saarbrücken; **Titelbild:** KI-generiert, gestalterische Bearbeitung DFKI.

[KI] Abbildung enthält KI-generierte Bildelemente. Konzept, Composing und Gestaltung: **DFKI**. Bei fotorealistischen Abbildungen erfolgt die Kennzeichnung zusätzlich direkt am Bild.

Credits

Titelseite: DFKI, Lando M. Lehmann [KI]; **Seite 5:** DFKI, Oliver Dietze; **Seite 6:** DFKI, Lando M. Lehmann [KI]; **Seite 8:** DFKI, Juan D. Mejía [KI] (Person); Louisa Howard, Public Domain (Mikroskopie); **Seite 9:** DFKI, Juan D. Mejía [KI]; **Seite 11:** DFKI, Lando M. Lehmann [KI]; **Seite 12:** DFKI, Juan D. Mejía [KI]; **Seite 14:** DFKI [KI]; **Seite 15:** DFKI, Lando M. Lehmann [KI]; **Seite 16:** DFKI, Lando M. Lehmann [KI]; **Seite 18:** DFKI, Juan D. Mejía [KI]; **Seite 19:** DFKI, Juan D. Mejía [KI]; **Seite 21:** DFKI, Juan D. Mejía [KI]; **Seite 22 oben:** DFKI, Thomas Frank; DFKI, Annemarie Popp; **Seite 22 unten:** DFKI [KI].