

Glossar

Deep Learning	Teilbereich des maschinellen Lernens, welche viele (hierarchische) Schichten von neuronalen Netzen nutzt. Deep Learning kann effizient komplexe Informationen verarbeiten und so präzise Vorhersagen treffen (Universität Tübingen o.J.).
Deep-Learning-basierte Pipelines	Workflows, die verschiedene Phasen umfassen, um Modelle für maschinelles Lernen, die auf Deep Learning basieren, zu entwickeln, zu trainieren, bereitzustellen und zu verwalten (IBM o.J.).
Vision-Language-Modell	Vision Language Models (VLMs) sind Modelle der künstlichen Intelligenz (KI), die Funktionen der Computer Vision und der Verarbeitung natürlicher Sprache (NLP) miteinander verbinden. VLMs lernen, die Beziehungen zwischen Textdaten und visuellen Daten wie Bildern oder Videos abzubilden, sodass diese Modelle Text aus visuellen Eingaben generieren oder Prompts in natürlicher Sprache im Kontext visueller Informationen verstehen können (IBM o.J.).
Critical Point Thinking	Critical Point Thinking (CPT) ist ein Denkansatz, der darauf abzielt, die wichtigsten oder einflussreichsten Aspekte eines Systems, Problems oder einer Situation zu identifizieren, um fundiertere Entscheidungen zu treffen und effektivere Lösungen zu entwickeln. Es betont, dass nicht alle Faktoren gleich wichtig sind und die Konzentration auf die entscheidenden Punkte zu besseren Ergebnissen führt (Zhu et al. 2025).
Halluzinationen	Bezeichnet bei KI-Anwendungen Antworten bzw. Ausgaben, welche zwar plausibel erscheinen, aber nicht auf Fakten basieren und falsch sind. Eine KI-Anwendung generiert den wahrscheinlichsten Output, basierend auf ihren Trainingsdaten. Dabei achtet diese nicht auf Fakten oder Logik, sondern arbeitet mit Wahrscheinlichkeiten. Dies erzeugt zwar oft ein zuverlässiges Ergebnis, jedoch nicht immer (Universität Tübingen o.J.).
Strukturierte Verifikation	Strukturierte Verifikation ist ein geplanter, systematischer Prozess zur Überprüfung, ob ein Produkt (z.B. Software, Hardware) die definierten Anforderungen und Spezifikationen korrekt und vollständig erfüllt (Universität Tübingen, o.J.).
Graphneuronale Netzwerke (GNNs)	Graph Neuronale Netze (GNNs) sind spezielle Deep-Learning-Modelle, die Daten verarbeiten, die als Graphen (Knoten und Kanten) strukturiert sind, anstatt in regelmäßigen Gittern oder Sequenzen, und lernen die Beziehungen zwischen Entitäten (z. B. in sozialen Netzwerken, Molekülen) zu verstehen und vorherzusagen, indem sie Informationen iterativ zwischen verbundenen Knoten austauschen (Message Passing). Sie sind eine leistungsstarke Erweiterung traditioneller neuronaler Netze für vernetzte Daten, die komplexe Abhängigkeiten erfassen, wo die Struktur genauso wichtig ist wie die Daten selbst (IBM, o.J.).
Spatial-Heterophily-aware GNN	Graphneuronales Netzwerk, das die Wirkungsstärke von Nachbarschaftsbeziehungen integriert und mit Gewichtungen versieht (Xiao et al. 2023).
Rotations-Skalierungs-Aggregationsmodul	Methode, um räumlich nah beieinander liegende Nachbarn richtig zu gruppieren und jede Gruppe mit geringerer Vielfalt innerhalb der Gruppe separat zu verarbeiten (Xiao et al. 2023).

Multimodal Large Language Model (MLLM)	Ein großes generatives KI-Modell, das multimodale Informationen wie Texte, Bilder, Videos, Audioinhalte und andere Informationen sowie Kombinationen dieser verstehen, verarbeiten und generieren kann (Universität Tübingen, o.J.).
Zero-Shot-Prompt	KI-Methode, bei der ein Sprachmodell eine Aufgabe oder Frage ohne vorherige Beispiele oder Demonstrationen im Prompt lösen soll (IBM, o.J.).
Prompt-Chain-Architektur	Architekturmuster im Bereich des Prompt Engineering, bei dem eine komplexe Aufgabe in eine Abfolge kleinerer, sequenzieller Prompts (Eingabeaufforderungen an ein KI-Modell) zerlegt wird. Die Ausgabe eines Prompts dient dabei als Eingabe oder Kontext für den nächsten Prompt in der Kette (IBM, o.J.).
Convolutional Neural Networks	Convolutional (dt.: „faltende“) neuronale Netze sind biologischen Prozessen nachempfunden. Sie arbeiten im Bereich des maschinellen Lernens mit einer zwei- oder dreidimensionalen Anordnung von Neuronen, die lokal miteinander verbunden sind und gewichtet werden (vgl. Ketkar and Moolayil 2021).
Retrieval Augmented Generation (=RAG)-Architektur	RAG-Architekturen (dt. abrufgestützte Generierung) verbessern die Leistungsfähigkeit großer Sprachmodelle. Sie ermöglichen Zugriff auf externe, aktuelle und spezifische Informationen während der Antwortgenerierung. Ohne RAG verlassen sich LLMs ausschließlich auf die Daten, mit denen sie trainiert wurden (vgl. Gao et al. 2023)