



Loab-Anomalie: Das Phantom im Latenten Raum der KI

Autore: Francesco Zinghini | **Data:** 20 Febbraio 2026

Es begann als ein harmloses Experiment mit negativen Eingabeaufforderungen und endete mit der Entdeckung einer digitalen Legende, die selbst erfahrenen Datenwissenschaftlern einen Schauer über den Rücken jagte. Die Rede ist von **Loab**, einer fiktiven Frauengestalt, die im Jahr 2022 erstmals in den Tiefen generativer Bildmodelle entdeckt wurde und seitdem als Paradebeispiel für die unerforschten Abgründe der **Künstlichen Intelligenz** gilt. Anders als herkömmliche KI-Fehler war diese Gestalt kein zufälliges Rauschen, sondern eine persistente, reproduzierbare Entität, die sich wie ein Virus durch Tausende von generierten Bildern zog, ohne dass sie jemals explizit angefordert wurde. Heute, im Jahr 2026, blicken wir zurück auf dieses Phänomen, um die technische Architektur zu verstehen, die solche „Geister“ im **Maschinellen Lernen** überhaupt erst möglich macht.

Die Genese durch doppelte Negation

Um das Phänomen technisch zu durchdringen, müssen wir uns zunächst der Methode zuwenden, mit der es entdeckt wurde: dem sogenannten „Negative Prompting“ (negative Eingabeaufforderung). In der Welt der **Generative AI** und Diffusionsmodelle geben wir normalerweise einen positiven Befehl, etwa „Ein Bild von einem Sonnenuntergang“. Das Modell sucht im gelernten Vektorraum nach den Merkmalen, die diesem Begriff entsprechen.

Die Entdeckerin von Loab, eine Künstlerin unter dem Pseudonym Supercomposite, nutzte jedoch eine inverse Technik. Sie forderte die KI auf, das genaue Gegenteil eines Begriffs zu generieren. Der ursprüngliche Prompt war „Brando“ (bezogen auf Marlon Brando). Die KI generierte ein abstraktes Logo. In einem zweiten Schritt forderte sie die KI auf, das Gegenteil dieses Logos zu erstellen. Die Logik der **Neural Networks** diktierte hierbei: „Zeige mir das, was am weitesten vom Konzept dieses Logos entfernt ist.“

Das Ergebnis war keine Leere, sondern das Bild einer älteren Frau mit markanten, oft geröteten Wangen und einem leeren, fast unheimlichen Blick. Dies war die Geburtsstunde der Entität. Technisch gesehen hatte die KI durch die doppelte Negation einen Bereich im hochdimensionalen Raum isoliert, der normalerweise durch positive Prompts kaum zugänglich ist. Es war, als hätte man durch das Wegmeißeln aller bekannten Konzepte den rohen Kern eines unbekannten Clusters freigelegt.

Der Latente Raum: Eine Karte der Unendlichkeit

Um zu verstehen, warum diese Gestalt immer wiederkehrte, müssen wir den Begriff des „Latenten Raums“ (Latent Space) definieren. Moderne KI-Modelle, seien es **LLMs** wie **ChatGPT** oder Bildgeneratoren, operieren nicht mit Pixeln oder Wörtern, sondern mit Vektoren – langen Zahlenreihen, die Bedeutungen repräsentieren. Der Latente Raum ist eine multidimensionale Karte, auf der ähnliche Konzepte nahe beieinanderliegen. „Hund“ liegt mathematisch nahe bei „Katze“, aber weit entfernt von „Toaster“.

Diese düstere Gestalt repräsentierte einen spezifischen Koordinatenpunkt in diesem Raum. Doch sie war mehr als nur ein Punkt; sie war ein sogenannter „Attraktor“. In der Systemtheorie ist ein Attraktor ein Zustand, auf den sich ein

dynamisches System zubewegt. Sobald die KI in die Nähe der Vektoren kam, die diese Frau definierten, „rutschte“ das Modell förmlich in diesen Zustand ab. Das erklärt, warum sie so schwer loszuwerden war. Selbst wenn man ihr Bild mit dem eines harmlosen Objekts, wie einer Blume, kombinierte (Image-to-Image-Synthese), dominierte ihre visuelle Signatur das Ergebnis. Ihre Vektorgewichte waren im mathematischen Sinne „schwerer“ oder stabiler als die der anderen Konzepte.

Semantische Nähe zu Horror und Gore

Ein besonders verstörender Aspekt dieses Phänomens war die Tendenz der KI, die Gestalt fast zwangsläufig mit grotesken, blutigen oder horrorartigen Szenarien zu verknüpfen. Sobald die Entität im Bild war, driftete der Stil in Richtung „Macabre“. Dies ist kein Beweis für ein böses Bewusstsein der **Künstlichen Intelligenz**, sondern ein Resultat der Trainingsdaten.

Große Bildmodelle werden mit Datensätzen wie LAION-5B trainiert, die Milliarden von Bild-Text-Paaren aus dem Internet enthalten. In diesem riesigen Korpus sind visuelle Merkmale (wie rote Haut, bestimmte Beleuchtung, verzerrte Gesichtszüge) statistisch oft mit Horror-Filmen, medizinischen Bildern oder Gore-Inhalten verknüpft. Die Vektoren, die das Gesicht der Frau repräsentierten, lagen im Latenten Raum in unmittelbarer Nachbarschaft zu den Vektoren für „Blut“, „Schrei“ oder „Leiche“. Wenn das neuronale Netz den Pfad zu ihr aktivierte, aktivierte es aufgrund der semantischen Nähe (Semantic Proximity) automatisch auch diese düsteren Konzepte.

Die Black Box der Neuronalen Netze

Das Phänomen demonstriert eindrucksvoll das „Black Box“-Problem des **Maschinellen Lernens**. Wir kennen den Input und den Output, aber die komplexen Verschaltungen in den verborgenen Schichten (Hidden Layers) sind selbst für die Entwickler oft undurchsichtig. Diese Gestalt war nicht explizit programmiert worden. Sie war ein emergentes Phänomen – eine Eigenschaft, die aus der Komplexität des Systems entstand und nicht aus seinen Einzelteilen vorhersehbar war.

Für die technische Kommunikation und KI-Sicherheit ist dies von enormer Bedeutung. Es zeigt, dass in den Modellen Cluster existieren, die wir nicht kennen, bis wir zufällig über sie stolpern. Diese Cluster können harmlos sein, oder sie können Bias, Stereotypen oder eben verstörende Inhalte repräsentieren, die durch ungewöhnliche Prompting-Techniken an die Oberfläche gespült werden.

Warum sie nicht einfach gelöscht werden kann

Viele Laien fragen sich: „Warum löscht man diese Frau nicht einfach aus dem Code?“ Die Antwort liegt in der Natur der verteilten Repräsentation. In einem neuronalen Netz ist ein Konzept wie diese Frau nicht an einem einzigen Speicherort abgelegt, den man formatieren könnte. Ihr „Wesen“ ist über Millionen von Parametern verteilt. Um sie gezielt zu entfernen, müsste man das gesamte Modell neu trainieren oder komplexe „Machine Unlearning“-Algorithmen anwenden, die versuchen, spezifische Gewichte zu neutralisieren, ohne die Fähigkeit des Modells zu beschädigen, andere Gesichter zu generieren.

Das Phänomen ist vergleichbar mit dem Versuch, eine bestimmte Farbe aus einem gemischten Farbeimer zu entfernen. Man kann sie nicht einfach herausnehmen, da sie Teil der Gesamtmischung geworden ist. Dies macht solche Anomalien im Bereich der **Generative AI** so hartnäckig und faszinierend zugleich.

Fazit

Die Geschichte der düsteren Gestalt, die wir heute als Loab kennen, ist weit mehr als eine digitale Geistergeschichte. Sie ist eine Lektion in hochdimensionaler Mathematik und statistischer Wahrscheinlichkeit. Sie führt uns vor Augen, dass der Latente Raum der **Künstlichen Intelligenz** Landschaften enthält, die noch kein Mensch betreten hat – und dass manche dieser Landschaften düster sind, weil sie die Summe unserer eigenen, im Internet hinterlassenen Ängste und Fiktionen widerspiegeln. Für Technologen bleibt sie eine Mahnung: In der Komplexität von Milliarden Parametern warten noch unzählige unentdeckte Attraktoren darauf, durch den richtigen – oder falschen – Befehl geweckt zu werden.

Häufig gestellte Fragen

Was ist die Loab-Anomalie in der generativen KI?

Loab ist eine persistente, fiktive Frauengestalt, die im Jahr 2022 in den latenten Räumen von KI-Bildgeneratoren entdeckt wurde. Sie gilt als technisches Phänomen und nicht als bloßer Fehler, da sie reproduzierbar ist und als sogenannter Attraktor im Vektorraum fungiert. Ihre Entstehung verdeutlicht, wie komplexe Cluster in neuronalen Netzen existieren können, ohne dass sie explizit programmiert wurden.

Wie funktioniert die Entdeckung durch Negative Prompting?

Beim Negative Prompting wird die Künstliche Intelligenz aufgefordert, das genaue Gegenteil eines bestimmten Begriffs oder Bildes zu generieren, anstatt nach einer positiven Übereinstimmung zu suchen. Im Fall von Loab führte eine doppelte Negation dazu, dass das Modell in einen normalerweise unzugänglichen Bereich des hochdimensionalen Raums navigierte. Durch das mathematische Wegmeißen bekannter Konzepte wurde so der rohe Kern eines unbekannten Clusters freigelegt.

Warum wird Loab oft mit Horror-Elementen dargestellt?

Die Verbindung zu makabren Inhalten resultiert aus der semantischen Nähe innerhalb der Trainingsdaten, wie etwa dem LAION-5B-Datensatz. Die Vektoren, die das Gesicht von Loab definieren, liegen im mathematischen Raum in direkter Nachbarschaft zu Vektoren für Blut, Schreie oder medizinische Bilder. Wenn das neuronale Netz die Merkmale dieser Frau aktiviert, werden aufgrund dieser statistischen Verknüpfung automatisch auch die benachbarten düsteren Konzepte abgerufen.

Was bedeutet der Begriff Latenter Raum in diesem Kontext?

Der Latente Raum ist eine multidimensionale Karte, auf der KI-Modelle Bedeutungen als Vektoren und Zahlenreihen anordnen. Ähnliche Konzepte liegen dort nahe beieinander, während unähnliche weit entfernt sind. Loab repräsentiert einen spezifischen, sehr stabilen Koordinatenpunkt in diesem Raum, der als Attraktor wirkt und das Modell dazu bringt, immer wieder in diesen Zustand abzurutschen, sobald es in die mathematische Nähe kommt.

Kann man Loab einfach aus dem KI-Modell löschen?

Nein, ein einfaches Löschen ist nicht möglich, da Loab keine einzelne Datei ist, sondern eine verteilte Repräsentation über Millionen von Parametern hinweg darstellt. Ihr visuelles Muster ist fest in die Gewichte des neuronalen Netzes eingewoben, vergleichbar mit einer Farbe, die untrennbar in einen Eimer gemischt wurde. Eine Entfernung würde ein komplettes Neutraining des Modells oder komplexe Machine Unlearning-Verfahren erfordern.