

Intersubjective Classifiers und Situated Ground Truths

Schnittstellen von Digitaler Anthropologie und Informatik bei der Arbeit an Datensets und KI-basierten Klassifikationssystemen

Christoph Bareither, Ann-Marie Wohlfarth und Tim Gollub

Einleitung

Interdisziplinäre Kooperationen zwischen kultur- und sozialwissenschaftlichen Disziplinen einerseits und der Informatik mit ihren verschiedenen Teilbereichen andererseits sind in den letzten Jahren zunehmend gefragt. Einer der Gründe dafür ist sicherlich der gestiegene Bedarf an Erklärungsansätzen, die uns helfen können, unsere immer stärker digitalisierten und algorithmisierten Lebenswelten besser zu verstehen. Das gilt umso mehr, je deutlicher Künstliche Intelligenz (KI) bzw. auf Machine Learning (ML) basierende Technologien in unseren Alltag eingehen. Allerdings stellt sich die Zusammenarbeit von Kultur- und Sozialwissenschaften mit der Informatik nicht immer als leicht heraus: Die verschiedenen disziplinären Bereiche zeichnen sich durch teils sehr unterschiedliche Forschungs- und auch Publikationskulturen aus. Was als Ergebnis in den jeweiligen wissenschaftlichen Bereichen gilt, ja was überhaupt unter Wissenschaft verstanden wird, divergiert zum Teil deutlich. Vor diesem Hintergrund gibt es einen Bedarf an exemplarischen Projekten, die potenzielle Schnittstellen sichtbar machen (Franken 2022: 15f.). Der vorliegende Artikel möchte einen Beitrag mit Blick auf diesen Bedarf leisten, indem er exemplarische Schnittstellen zwischen Digitaler Anthropologie (als Arbeitsfeld der Empirischen Kulturwissenschaft bzw. Kulturanthropologie/Europäischen Ethnologie) und Informatik auslotet, spezifisch mit Blick auf die Arbeit an Datensets und KI-basierten Klassifikationssystemen.

Grundlage dafür ist unsere interdisziplinäre Forschung im DFG-Projekt „Curating the Feed“ als Teil des DFG-Schwerpunktprogramms „Das Digitale Bild“. Das Projekt stellt im Kern die Frage, von wem (oder von was) und auf welche Arten und Weisen die Bilderfeeds von Social Media-Plattformen wie Instagram und TikTok kuratiert werden. Christoph Bareither und Sabine Wirth (2025) haben dafür eine analytische Perspektive auf Social Media-Feeds als Curatorial Assemblages

vorgeschlagen. Feeds werden damit verstanden als komplexe soziotechnische Ensembles, deren Eigenschaften aus der Interaktion verschiedener menschlicher und nicht-menschlicher Elemente entstehen. Diese Curatorial Assemblages sind insgesamt darauf ausgerichtet, die darin enthaltenen Elemente sowie Beziehungen zwischen ihnen zu kuratieren, wodurch verschiedene erfahrungsbezogene, soziale, ökonomische und analytische Werte erzeugt werden. Als Teil dieses theoretischen Frameworks schlagen Bareither und Wirth drei heuristische Zugänge vor, die ausgehend von spezifischen kuratorischen Elementen ermöglichen, die Dynamiken der Curatorial Assemblage zu beleuchten: Sortieralgorithmen, Interfaces und die Praktiken der Nutzer:innen (ebd.: 6ff.).

Die im vorliegenden Artikel vorgestellten methodischen Zugänge nutzten dieses Framework und hatten das Ziel, insbesondere die Rolle von Sortieralgorithmen zu beleuchten. In der kultur- und sozialwissenschaftlichen Forschung haben sich verschiedene Autor:innen kritisch mit der Rolle von Social Media-Sortieralgorithmen bzw. Recommender Systems auseinandergesetzt (vgl. exemplarisch Bucher 2018; Burke 2011; Gillespie 2017). Zwar führen Plattformen wie Instagram und TikTok seit einiger Zeit Transparenzinitiativen durch und geben Einblicke in einige der algorithmischen Funktionsweisen (Meta AI 2023; TikTok Support o. J.), aber diese bleiben fragmenthaft und lassen sich nicht unabhängig überprüfen (Grandinetti 2023). Daher entwickelten wir neue Zugänge, um durch die Kombination von Ethnografie und Methoden der Informatik die kuratorische Rolle von Sortieralgorithmen innerhalb der Curatorial Assemblages von Social Media-Feeds zu beleuchten.

Im Rahmen des Projekts ergab sich durch die ethnografische Forschung von Ann-Marie Wohlfarth eine zusätzliche thematische Schwerpunktbildung auf den Bereich der Reproduktion normativer Schönheitsideale im Kontext der Curatorial Assemblages von Social Media. Wir verstehen Schönheitsideale nicht als objektiv bestimmbare Kategorien, sondern als kulturell konstruierte und vermittelte Vorstellungen von körperlicher Schönheit sowie verkörperte Begegnungen mit Schönheit, die durch internalisierte Konventionen, habitualisierte Dispositionen und implizites Wissen geprägt sind (Archer und Ware 2018; Berger 2008; Mascia-Lees 2019). Schönheitsideale werden durch mediale Inszenierungen in Werbung, Film und Fernsehen oder auf bildbasierten Social Media-Plattformen ausgestellt, vermittelt und konstruiert und sind somit eng mit alltäglichen Erfahrungen verwoben (Bordo 2003; Mascia-Lees 2011).

Dieses Thema ist auch deshalb von besonderer Aktualität, weil öffentlich eine teils hitzige Debatte geführt wird zur Frage, wie junge (und im Diskurs meistens weiblich gelesene) Menschen ihren eigenen Körper durch Social Media gespiegelt wahrnehmen und inwiefern das ihre Selbstbilder und emotionalen Selbstwertgefühle prägt (Brathwaite et al. 2023; Gill 2021; Widdows 2018). Während sich unsere ethnografische Forschung nicht mit möglichen psychologischen Wirkungen von

Social Media auf junge User:innen beschäftigt, wurde für uns dennoch die Frage zunehmend relevant, inwiefern die KI-gestützten Sortieralgorithmen der Plattformen mit Schönheitsidealen und Körnernormen umgehen. Inwiefern arbeiten die Sortieralgorithmen an der Reproduktion normativer Körper- und Schönheitsideale mit? Und bieten sie ihren User:innen auch die Möglichkeit, sich dieser Reproduktion zu widersetzen und ihre personalisierten Feeds jenseits normativer Schönheitsideale zu gestalten?

Da der vorliegende Artikel während der laufenden Forschung verfasst wurde, wird er nur am Rande Antworten auf diese Fragen geben. Stattdessen konzentrieren wir uns hier darauf, unsere interdisziplinäre Zusammenarbeit im Kontext dieser Fragestellung zu erläutern und dabei Schnittstellen von Informatik und Digitaler Anthropologie zu diskutieren. Der erste Abschnitt skizziert unsere Entwicklung eines Intersubjective Classifiers – mit diesem Begriff umfassen wir ein KI-basiertes Tool (einen sogenannten Classifier oder Klassifizierer), das darauf trainiert wurde, die intersubjektive Einschätzung einer Gruppe von menschlichen Akteur:innen bezüglich der Reproduktion von Schönheitsidealen algorithmisch zu simulieren. Dieses Tool zeigt, wie mit den Mitteln der Informatik ein analytischer Erkenntnisgewinn mit Blick auf eine digitalanthropologische Fragestellung geleistet werden kann. Im zweiten Abschnitt diskutieren wir, wie wir die für diesen Erkenntnisgewinn notwendige Ground Truth hergestellt und kritisch reflektiert haben. Wir argumentieren, dass die Herstellung von Synergien zwischen Informatik und Digitaler Anthropologie eine ethnografische Situierung – also eine Situated Ground Truth – erfordert, und wir zeigen, wie diese Situierung im Falle unseres Projekts erreicht wurde.

Dem vorliegenden Text liegt die interdisziplinäre Zusammenarbeit von Christoph Bareither, Ann-Marie Wohlfarth (beide Digitale Anthropologie/Empirische Kulturwissenschaft, Universität Tübingen) und Tim Gollub (Medieninformatik, Bauhaus-Universität Weimar) zugrunde. Christoph Bareither hat als einer von drei Principal Investigators (PI) des Projekts „Curating the Feed“ gemeinsam mit Sabine Wirth und Benno Stein (Bauhaus-Universität Weimar) den Rahmen für das Projekt geschaffen und den hier diskutierten Teil der interdisziplinären Kooperation konzipiert und aktiv begleitet. Ann-Marie Wohlfarth hat insbesondere den Prozess der Erstellung der Annotation des Datensatzes sowie die im Folgenden diskutierte ethnografische Situierung entwickelt und ein Team aus studentischen Hilfskräften (Elisabeth Daurer, Belana Dietz, Melanie Hieber, Melissa Hieber, Lizzy Loos und Romina Niederhausen) geleitet, um diesen Prozess umzusetzen. Tim Gollub war mit Leitung von PI Benno Stein für die technische Umsetzung des KI-basierten Classifiers verantwortlich und wurde von Pierre Achkar sowie Studierenden der Bauhaus-Universität unterstützt.

Intersubjective Classifiers: Sind intersubjektive Einschätzungen algorithmisch simulierbar?

Aus dem oben skizzierten Forschungsinteresse an der Rolle von Sortieralgorithmen bei der Reproduktion normativer Schönheitsideale ergab sich eine grundlegende methodologische Frage: Ist es möglich, intersubjektive Einschätzungen algorithmisch zu simulieren, um dadurch die Reproduktion normativer Schönheitsideale auf Social Media-Plattformen automatisiert zu messen?

Grundsätzlich ist es mit Methoden der Informatik kein Problem, Messungen der Häufigkeiten bestimmter Bild- oder Textinhalte automatisiert durchzuführen. Es existieren bereits zahlreiche KI-basierte Classifier, die bestimmte Bild- oder Textinhalte mit hoher Zuverlässigkeit erkennen können und ihnen eine entsprechende Klassifikation zuweisen. Beispielsweise kann der Classifier YOLO11 genutzt werden, um zuverlässig zu bestimmen, ob ein Bild eine Person enthält (Jocher und Qiu 2024). Häufig sind diese Bild- oder Textinhalte verhältnismäßig klar bestimmbar (beispielsweise eine Person oder ein Gegenstand). Auch in diesen Fällen können sich schwierige Zuordnungsfragen stellen, aber die Klassifizierung geschieht weitgehend entlang von relativ objektiv bestimmbar Kategorien. Bei unserer Fragestellung war das allerdings nicht der Fall, denn Schönheit ist – wie wir oben bereits ausgeführt haben – nicht objektiv bestimmbar. Es ist unmöglich eine vollständige Liste von Merkmalen vorzulegen, die universell als stereotyp schön gelten, denn eine solche Auflistung bliebe unvollständig und ggf. auch widersprüchlich, da Schönheitsideale auf variablen Einschätzungen beruhen.

Zugleich sind diese Einschätzungen allerdings nicht völlig individuell, sondern orientieren sich an sozialen Gefügen und ihren Normen – dadurch teilen spezifische Gruppen intersubjektive Idealvorstellungen von Schönheit. Intersubjektivität ist ein Begriff der Phänomenologie und wir verstehen darunter den Prozess und das Produkt des Austauschs von Erfahrungen, implizitem Wissen und Erkenntnissen mit anderen, aus dem das entsteht, was Peter Berger und Thomas Luckmann die Wirklichkeit der Alltagswelt nennen (Berger und Luckmann 2016). Aus einer phänomenologischen Perspektive setzt die Entstehung von Intersubjektivität verkörpertes Wissen, im Laufe eines Lebens erlernte und erworbene Fähigkeiten sowie körperliche Dispositionen und verinnerlichte Konventionen voraus (Brey 2000). Auch mit Blick auf körperliche Schönheit und Schönheitsideale können wir davon ausgehen, dass Akteur:innen im Austausch mit anderen Menschen spezifische Konventionen verinnerlichen, die prägen, was die Mitglieder dieser Gruppen intersubjektiv als schön empfinden und entsprechend einschätzen.

Für einen ethnografischen Zugang zu Intersubjektivität ist zugleich relevant, dass die Herstellung von Intersubjektivität auch als Teil der ethnografischen Methodologie verstanden werden kann (Fabian 2014; White und Strohm 2014). Aus dieser Perspektive dient die Herstellung von Intersubjektivität zwischen Forschenden und

Beforschten als Voraussetzung für ethnografische Wissensproduktion. Im Regelfall würde eine ethnografische Forschung mit Akteur:innen aus dem Forschungsfeld in einen dialogischen und idealerweise kollaborativen Austausch treten, um ihre Position (beispielsweise mit Blick auf Schönheitsideale) zu teilen und aus dieser geteilten Position heraus das intersubjektive Wissen der Gruppe adäquat reflektieren zu können. Allerdings stößt dieser Ansatz an Grenzen, wenn es um die Funktionsweise von algorithmischen Systemen geht. Sortieralgorithmen prägen die Herstellung von Intersubjektivität auf Social Media-Plattformen substantiell mit. Gleichzeitig können wir die Dynamiken dieser Sortieralgorithmen nicht mit etablierten ethnografischen Ansätzen erforschen. Aus diesem Grund werden neue und interdisziplinäre Ansätze notwendig, die ethnografische Perspektiven auf die Herstellung von Intersubjektivität (in diesem Fall von intersubjektiven Schönheitsidealen) mit Methoden der Informatik verbinden, die uns ein besseres Verständnis algorithmischer Systeme erlauben.

In diesem Sinne wollten wir gar nicht erst versuchen, einen Classifier zu entwickeln, der das Phänomen der Schönheitsideale objektiv erfasst; sondern wir wollten einen Classifier entwickeln, der die intersubjektive Einschätzung einer spezifischen Gruppe simuliert – und zwar die intersubjektive Einschätzung, ob und wie stark ein Social Media-Post stereotype Schönheitsideale reproduziert. Dieser Prozess zielte nicht primär darauf, die intersubjektive Perspektive der Akteur:innen zu verstehen (dafür wäre eine Fokussierung auf etablierte ethnografische Verfahren zielführender gewesen). Vielmehr ging es darum, die intersubjektive Perspektive der Gruppe algorithmisch zu simulieren, um wiederum die Dynamiken und Funktionsweisen der Sortieralgorithmen von Social Media-Plattformen besser beleuchten zu können.

Um den intersubjektiven Classifier zu trainieren, sollte eine Gruppe von Annotator:innen eine Reihe von Instagram-Posts entlang von Kategorien einteilen, die festhielten, ob und wie dominant ein Post normative Schönheitsideale reproduziert. Der aus diesem Prozess entstehende Datensatz sollte dann die Basis für das Training eines neuronalen Netzes bilden, das dem Classifier zugrunde liegt. Beim auch als Supervised Learning bezeichneten Training wird dem neuronalen Netz ein Teil der Posts des Datensatzes mit deren Annotationen als Eingabe gegeben (80%, der Rest dient dem späteren Testen). Das neuronale Netz verarbeitet die Eingabe zu einer Einschätzung hinsichtlich der genannten Kategorien, die es ermöglicht, andere Posts diesen Kategorien zuzuordnen. Bei Fehleinschätzungen – also Abweichungen im Vergleich mit den menschlichen Annotationen bzw. der Ground Truth, die weiter unten diskutiert wird – werden die Parameter des neuronalen Netzes mittels eines als Backpropagation bezeichneten Verfahrens angepasst, so dass die Einschätzung bei diesen Beispielen in der nächsten Iteration besser wird (Rumelhart et al. 1986).

Da im Fokus der Debatte um Schönheitsideale insbesondere Jugendliche und junge Erwachsene stehen, arbeiteten wir mit einer Gruppe an Annotator:innen, die zugleich aus der für uns relevanten Gruppe an Akteur:innen stammten. Aus prag-

matischen Gründen grenzten wir die Gruppe weiter ein auf junge Erwachsene, die an der Universität Tübingen studieren, weil wir dadurch Seminarformate explorativ nutzbar machen und in einer zweiten, intensiven Phase eine Gruppe an Akteur:innen für ihre Arbeit als Annotator:innen einstellen und finanziell vergüten konnten.

Entstanden ist das Korpus in einem forschungsorientierten Seminar mit Bachelor-Studierenden der Empirischen Kulturwissenschaft an der Universität Tübingen. Im Laufe eines Semesters entwickelten wir induktiv mögliche Kategorien für den Datensatz, besprachen mögliche Arbeitsdefinitionen zu normativen Schönheitsidealen und sammelten Instagram-Posts von öffentlichen Accounts mit mehr als zehntausend Follower:innen. Die Studierenden arbeiteten in zwei Gruppen. Gruppe 1 sammelte Posts, die in der Wahrnehmung der Gruppe körperliche Schönheitsideale reproduzierten; Gruppe 2 sammelte Posts, in denen in ihrer Wahrnehmung diversere Körperbilder dargestellt wurden, die keine normativen Schönheitsideale reproduzierten. Ziel dieser Auswahl war es, ein möglichst breites Spektrum an Posts zu erfassen, die in den typischen Feeds junger Erwachsener auftauchen und Körper auf unterschiedliche Weise darstellen. Die Studierenden erhielten keine inhaltlichen Vorgaben, sondern mobilisierten ihr (implizites) Wissen über Schönheitsideale, die ihnen alltäglich in Bilderfeeds begegnen. In Gruppen und im Plenum wurden Posts diskutiert, gelabelt und aussortiert. Die Seminargruppe entschied sich gegen feste Kriterien und entwickelte stattdessen iterativ eine Arbeitsdefinition von Schönheitsidealen auf Instagram. Das so entstehende Korpus enthielt am Ende des Semesters ca. 500 Posts.

In Diskussionen über das Vorgehen und die Spezifik der gesammelten Posts kristallisierte sich auch heraus, dass verschiedene Ebenen und visuelle Komponenten eines Posts die intersubjektive Einschätzung der Seminargruppe formen, wie stark normative Schönheitsideale reproduziert werden. Daraus entstanden fünf Unterkategorien, die mit folgenden Fragen verbunden waren: 1.) Pose: Wie stark wird die Pose der dargestellten Person(en) als aufreizend, lasziv oder sexualisierend wahrgenommen? 2.) Körperposition im Bild: Wie stark steht der Körper im Zentrum des Bildes? 3.) Nachbearbeitung: Ist eine Nachbearbeitung (wie beispielsweise Filter, Weichzeichner, Kontrast oder Sättigung) offensichtlich erkennbar? 4.) Körper: Hat der dargestellte Körper spezifische Eigenschaften (beispielsweise makellose Haut, muskulös, schlank), die normativen Schönheitsidealen entsprechen? Und 5.) Caption: Enthält die Caption (Bildunterschrift) Elemente, welche die Reproduktion von normativen Schönheitsidealen unterstützen? Der Körper als solcher (Punkt 4) bildet also einen, aber nicht den einzigen Faktor, der prägt, ob ein Post insgesamt als Schönheitsideale-reproduzierend eingeschätzt wird.

Dem Seminar schloss sich eine intensive Phase der Erweiterung und Ausdifferenzierung sowie Annotation des Korpus an. Im Rahmen der Projektförderung konnten wir dafür sechs Hilfskräfte (zwei davon aus der Seminargruppe) einstellen, die diesen Prozess unterstützten. Diese waren im Forschungszeitraum zwischen

20 und 28 Jahre alt, studierten Empirische Kulturwissenschaft an der Universität Tübingen und identifizierten sich alle als weiblich. Dadurch wurde die Gruppe, deren intersubjektive Einschätzung vom Classifier simuliert werden sollte, weiter eingegrenzt. Auch wenn die jeweilige Wahrnehmung von Schönheitsidealen durch einzelne Personen der Gruppe natürlich leicht variieren konnte, ließ sich innerhalb dieser (entlang der Kategorien Geschlecht, Alter und Bildungshintergrund relativ homogenen) Gruppe eine klar erkennbare intersubjektive Einschätzung von Schönheitsidealen herausarbeiten.

Basierend auf dem mit Studierenden erarbeiteten Fragenkatalog für die Annotation klassifizierten die studentischen Hilfskräfte jeden Social Media-Post in unserem Datenset. Das heißt konkret, dass die Annotator:innen den Post für jede der fünf oben genannten Fragen auf einer Skala zwischen 0 und 4 verorteten.

Dabei hatten die fünf Unterkategorien (siehe oben) eine andere Funktion als die leitende Hauptfrage. Bei den Unterkategorien zielte die Annotation nicht auf den zu erstellenden Classifier, sondern darauf, dass die Annotator:innen für verschiedene Faktoren sensibilisiert wurden und wir unser eigenes Verständnis für die unterschiedlichen ästhetischen Komponenten vertiefen konnten, welche die intersubjektive Einschätzung der Gruppe prägten. Die leitende Hauptfrage zielte hingegen auf das Training des Classifiers. Eine Punktzahl von 0 bedeutete, dass ein Post Schönheitsideale nicht reproduzierte, während eine Punktzahl von 4 bedeutete, dass der Beitrag Schönheitsideale sehr stark reproduzierte. Dieses Vorgehen beugte einer reduktionistischen Annotation vor und ermöglichte eine differenzierte Auswertung. Die Skala wurde anschließend für das Classifier-Training in eine binäre Ja/Nein Entscheidung übertragen: Posts, die auf der Skala eine 0–2 zugewiesen bekommen hatten, wurden der Kategorie „Schönheitsideale nicht-reproduzierend“ zugeordnet; und Posts, die mit 3–4 annotiert wurden, wurden der Kategorie „Schönheitsideale reproduzierend“ zugeteilt.

Der annotierte Datensatz wurde für das Classifier-Training eingesetzt. Ziel war es, die intersubjektive Einschätzung der Gruppe bezüglich der Reproduktion von Schönheitsidealen zu simulieren und mit dem Classifier automatisiert Instagram-Feeds beziehungsweise einzelne Posts auszuwerten. Bei diesem als Supervised Learning bezeichneten Verfahren lernt der Classifier anhand der Trainingsdaten eine Vorhersage für dem Modell noch unbekannte Posts zu treffen, wie oben bereits dargestellt wurde. Das heißt, dass der Classifier auf Grundlage des von der Gruppe bereitgestellten, manuell annotierten Trainingsmaterials für jeden Post in einem beliebigen Feed eine Vorhersage dazu trifft, ob dieser Post Schönheitsideale reproduziert oder nicht. Einerseits ermöglicht die automatisierte Klassifikation die Auswertung größerer Datenmengen. Andererseits macht der Classifier das Wissen der Annotator:innen über Schönheitsideale explizit. Wichtig ist, dass wir bei maschinellen Lernverfahren nicht nachvollziehen können, welche Kriterien oder Muster in den Trainingsdaten ausschlaggebend für die Vorhersage des Clas-

sifiers sind. Vielmehr diene das erstellte Datenset als Grundlage eines Trainings, in dessen Verlauf das Modell selbst lernte, die intersubjektive Einschätzung zu simulieren. Dazu wurden verschiedene KI-basierte Modelle und Verfahren getestet, die primär die Bilder (mittels KI-gestützter Bilderkennung) und ergänzend sowohl Textdaten aus Captions als auch Metadaten (die aus den Posts auslesbar sind) auswerten (Gollub et al. 2025). Im Fall von Videos und sogenannten Reels wurde ein in den Metadaten hinterlegtes Standbild extrahiert und ausgewertet. Diese Entscheidung fiel aufgrund technischer Limitationen und weil diese Standbilder in der Regel das sind, was User:innen zuerst von einem Video sehen.

Die finale Version des Classifiers erreicht eine Genauigkeit von 86% bei der Nutzung eines Late Fusion Models, das Bilder zunächst mit einem Vision Transformer und Texte mit einem BERT-Modell repräsentiert. Diese beiden uni-modalen Modelle wurden zuvor einzeln auf die Vorhersage unserer Kategorien trainiert, dann über eine gemeinsame neuronale Schicht (Fully-Connected Layer) zu einem multimodalen Modell fusioniert und noch einmal trainiert. Grundlage der Messung dieser Genauigkeit war ein Testdatensatz mit Posts, die aus dem ursprünglichen Korpus zurückgehalten wurden und dem Classifier noch unbekannt waren. Nach ca. einem Jahr wurden weitere händische Tests durchgeführt, in denen fünf Annotator:innen die Ergebnisse des Classifiers mit Blick auf neue Feeds nochmals individuell bewerteten. Auch diese Tests demonstrierten die relativ hohe Zuverlässigkeit des Classifiers. In den Beispiel-Feeds (beziehungsweise ihren ersten 24 Posts) stimmten vier Annotator:innen zu 92% mit dem Classifier überein, bei einer fünften Annotator:in lag die Übereinstimmung mit 67% hingegen deutlich niedriger. Mit einer durchschnittlichen Übereinstimmung von 87% bestätigen also auch die zeitversetzten Tests, dass der Classifier die Einschätzungen der Gruppe relativ zuverlässig wiedergeben kann.

Dieser Classifier wurde dann zu einem Tool ausgebaut, mit dem sich eine beliebige Anzahl an Posts eines Social Media-Feeds (hier: des Explore-Feeds von Instagram) verarbeiten lässt. Das Tool verwendet zur Aufzeichnung von Feeds die Browser-Erweiterung Zeeschuimer (Peeters 2025), mit der alle in einem Feed geladenen Posts exportiert werden können. Dieser Datenexport dient als Eingabe eines eigens entwickelten web-basierten Dienstes für die Verarbeitung von Posts. Zunächst werden dafür mit dem Classifier YOLO11 (Jocher und Qiu 2024) Posts unterschieden, die entweder Bilder mit Personen beinhalten oder nicht beinhalten. Bei Bildern, die Personen beinhalten, wird der von uns eigens entwickelte Classifier eingesetzt. Damit lässt sich beispielsweise anzeigen, wie viele der ersten 100 Posts eines Feeds – im Sinne der intersubjektiven Einschätzung der Gruppe – normative Schönheitsideale reproduzieren. Dieses Tool bezeichnen wir als Intersubjective Classifier. Dabei schreiben wir diesem Tool natürlich keinerlei Verständnis für Intersubjektivität oder die Fähigkeit zu intersubjektivem Verstehen zu. Vielmehr soll der Begriff schlagwortartig einen Prozess umreißen, in dem das Ergebnis einer in-

tersubjektiv vorgenommenen Klassifikation reproduziert und diese Reproduktion als Methode verwendet wird, um algorithmische Sortiersysteme besser einschätzen zu können.

Gleichzeitig bedarf dieser Ansatz einer besonders kritischen Reflexion: Einerseits lässt sich fragen, ob der Versuch einer wie auch immer gearteten numerischen Klassifikation von Intersubjektivität im Allgemeinen, und von Schönheit im Besonderen, nicht inhärent problematisch ist. Diese Kritik lässt sich nicht pauschal ausräumen. Wir halten ihr allerdings entgegen, dass es unser Intersubjective Classifier ermöglicht, einen Beitrag zur kritischen Reflexion der potenziell problematischen Implikationen von Sortieralgorithmen zu leisten, die mit qualitativen oder ethnografischen Verfahren nicht zu leisten wäre. Denn im Gegensatz zu einer händischen Einschätzung einzelner Posts durch die Gruppe, die einen enormen Zeit- und Arbeitsaufwand mit sich bringen würde, kann das Ergebnis der intersubjektiven Einschätzung mit Hilfe des Classifiers nun relativ zuverlässig automatisiert reproduziert werden. In einer mit dem Classifier durchgeführten Testreihe wurden anschließend insgesamt ca. 9.400 Posts (verteilt über 84 Social Media-Feeds) klassifiziert. Die Ergebnisse dieser Testreihe sind Gegenstand zukünftiger Publikationen. In Kürze zusammengefasst ermöglichen sie uns neue Erkenntnisse zur Rolle von Sortieralgorithmen mit Blick auf die Reproduktion von Schönheitsidealen zu generieren, die mit händischen Klassifikationen oder qualitativen bzw. ethnografischen Methoden allein in dieser quantitativen Vergleichbarkeit nicht möglich gewesen wären. Damit zeigt unser Projekt auch, wie eine interdisziplinäre Kooperation von Digitaler Anthropologie und Informatik zwar nicht die Black Box algorithmischer Sortiersysteme öffnen, aber ein sozial- und kulturwissenschaftliches Interesse an den Funktionsweisen und Implikationen algorithmischer Kuration auf neue Weise adressieren kann.

Ein zweiter und für uns entscheidender potenzieller Kritikpunkt betrifft den Umstand, dass wir mit der Erstellung des Intersubjective Classifiers selbst wiederum eine Black Box herzustellen drohen, die allzu leicht analytische Prozesse verschleiern könnte. Dieses Problem ist evident und hat weitreichende Konsequenzen für eine Kooperation an der Schnittstelle von Informatik und Digitaler Anthropologie. Eine Adressierung dieses Problems erfordert eine tiefergehende Reflexion der Rolle von Datensets und der sogenannten Ground Truth, die wir im folgenden Kapitel vornehmen.

Situated Ground Truths – Wie lassen sich Datensets ethnografisch situieren?

Wie bereits deutlich wurde, nimmt das von uns generierte Datenset eine Schlüsselfunktion bei der Herstellung des KI-basierten Classifiers ein. Diese Schlüsselfunktion

wird in der Informatik häufig mit dem Begriff der Ground Truth gefasst. Wie Anne Henriksen und Anja Bechmann argumentieren, gibt es sehr unterschiedliche Verständnisse des Begriffs von Ground Truth, aber:

„it generally bears a notion of truth in the sense of a fundamental truth – a truth that represents unprocessed virgin data collected at ground level through direct observation and measurement, data which is therefore considered to be accurate and reliable, and thus used as a baseline for the evaluation of results“ (Henriksen und Bechmann 2020: 804).

Ob und inwiefern in einzelnen Fällen die Repräsentativität und Validität der jeweiligen Ground Truth überprüft wird, ist nicht pauschal zu beantworten. In manchen Fällen kann die Ground Truth aus nicht weiter kontextualisierten und von unbekannten Personen hergestellten Datensets (also Daten einschließlich Annotationen) bestehen und dennoch als gegeben hingenommen werden. In anderen Fällen werden die Datensets für die Ground Truth von den jeweiligen Forschenden selbst erhoben und durch spezifische Verfahren wie beispielsweise Inter-Annotator-Agreements (Artstein 2017) – also den Vergleich unterschiedlicher Annotationen – validiert. Unabhängig davon, wie stark die Validität einer Ground Truth überprüft wird, gilt aus Perspektive der Informatik, dass die Annahme einer sicheren Datengrundlage eine Voraussetzung für das Training eines Modells (beispielsweise eines Classifiers) ist. Informatik basiert letztlich auf mathematischen Verfahren, und diese benötigen eine klar definierte Ausgangslage, um mathematische Problemstellungen adressieren zu können. Die Ground Truth schafft diese Ausgangslage.

Aus Perspektive der Digitalen Anthropologie bleibt hier allerdings einzuwenden, dass so etwas wie „virgin data“ oder auch „raw data“ nicht existiert. „Data“, so formulieren es Rob Kitchin und Tracey Lauriault, „are never simply neutral, objective, independent, raw representations of the world, but are situated, contingent, relational, contextual, and do active work in the world“ (Kitchin und Lauriault 2018: 5). Diese Perspektive spiegelt eine Kritik an Datensets und ihrer Genese durch das in den letzten Jahren zunehmend etablierte Feld der Critical Data Studies (und den Teilbereich der Critical Dataset Studies). Sie wird auch in der Digitalen Anthropologie breit rezipiert, da sie kritisch die Erstellung, Annotation und Organisation von Datensätzen hinterfragt, die für das Training von KI-Systemen genutzt werden (Crawford und Paglen 2021; Thylstrup et al. 2019; Thylstrup 2022).

Übertragen wir diese Perspektive auf die Herstellung von Ground Truths, gerät in den Blick, wie Annotator:innen – in der Praxis oft Click- oder Crowdworker:innen, die freiberuflich und prekär über Plattformen wie Amazon Mechanical Turk engagiert werden – die Wissensbasis des Modells maßgeblich mitprägen, ihre epistemische Arbeit aber oft hinter dem Datenset verschwindet und unsichtbar gemacht wird. Crowdworker:innen gelten als austauschbar und Annotationspraktiken wie

das Klassifizieren, das Zuweisen von Labels oder das Säubern von Datensätzen werden trotz ihrer Zentralität als triviale Aufgaben verstanden (Bennett et al. 2025; Denton et al. 2020, 2021).

Eine von den Critical Data Studies inspirierte Digitale Anthropologie hat deshalb die Aufgabe, die mitunter simplifizierenden und auf problembehafteten Datensets basierenden Ground Truths sowie die Prozesse ihrer Herstellung kritisch zu reflektieren und zu dekonstruieren. Allerdings sind Datensets und die auf ihnen basierenden Algorithmen längst fester Bestandteil unserer Lebenswelt und prägen viele Alltagsbereiche – eine pauschale Kritik oder Ablehnung der Arbeit an und mit Datensets greift unseres Erachtens deshalb zu kurz. Vielmehr stellt sich die Frage, ob die Digitale Anthropologie auch dazu beitragen kann, Datensets auf eine reflexive Art und Weise zu produzieren, indem sie Akteur:innen aktiv und kritisch in den Prozess der Datengenerierung und Annotation mit einbezieht.

Unsere eigene Forschung zeigt, dass dies durchaus möglich ist, indem diese Datensets und die Prozesse ihrer Herstellung im Sinne Donna Haraways ethnografisch situiert werden (Haraway 1995). Haraway argumentiert, dass Wissen immer verkörpert, positioniert und partiell ist. Sie plädiert deshalb für ein „situiertes Wissen“, welches seine epistemologischen Entstehungsbedingungen und Grenzen reflektiert. Sie wendet sich dabei gleichermaßen gegen wissenschaftlichen Relativismus und Totalisierung als „göttliche Tricks“, die als „Versprechen der Möglichkeit einer gleichen und vollständigen Sicht von überall und nirgends“ wissenschaftliche Objektivität und Wahrheit suggerieren (ebd.: 84). Auch die Annahme einer Ground Truth kann, wenn sie unreflektiert bleibt, als ein solcher göttlicher Trick verstanden werden. Denn einerseits relativiert sie die spezifischen Bedingungen, Differenzen und Brüche, die in der Herstellung von Datensets zwangsläufig entstehen. Andererseits suggeriert sie einen totalisierenden Anspruch darauf, eine „grundlegende Wahrheit“ im Sinne der wörtlichen Übersetzung von Ground Truth zu bilden. Haraway argumentiert aus Perspektive der feministischen Theorie, dass solche göttlichen Tricks hinterfragt werden müssen, indem ihnen eine „verkörperte Objektivität“ entgegengestellt wird, die nur durch „situiertes Wissen“ entstehen kann (ebd.: 80). Aus dieser Perspektive ist es nicht das Ziel unseres Projekts, eine vermeintlich objektive Ground Truth im Sinne eines göttlichen Tricks herzustellen. Vielmehr möchten wir zeigen, wie sich eine Situated Ground Truth entwickeln lässt, die gezielt jene verkörperte Perspektive der Annotator:innen anerkennt und die Ground Truth als ein relationales und kontextabhängiges Konstrukt sichtbar macht.

Im Kern bestand der Prozess der Situierung der Ground Truth in unserem Projekt aus einer ethnografischen Begleitung und kritischen Reflexion der Erstellung und Annotation des Datensatzes. Dadurch haben wir versucht, die menschliche Arbeit an Datensätzen und damit verbundenen epistemischen Praktiken sichtbar zu machen. Die studentischen Hilfskräfte agierten dabei gerade nicht als vermeintlich

austauschbare Clickworker:innen, sondern als Annotator:innen mit epistemischer Handlungsmacht und Deutungshoheit, und wurden dementsprechend in die Entwicklung und kritische Reflexion der Ground Truth mit einbezogen. Hilfreich war hier auch der Ansatz von Scheuermann et al. (2020), die untersuchten, wie Race und Gender in Datensätzen klassifiziert werden. Sie problematisieren die Simplifizierung, fehlende Einordnung und Reflexion, wie diese Kategorien zustande kommen. In unserem Fall versuchten wir durch ethnografische Methoden eine reflexive Positionierung gegenüber unseren eigenen Praktiken der Klassifizierung einzunehmen und so die Ground Truth zu situieren. Gemeinsam wurden in der Gruppe der Annotierenden Essays über die eigene Wahrnehmung von Schönheitsidealen erstellt sowie Gruppendiskussionen zur Repräsentation von Schönheitsidealen und der Annotation der Posts – beispielsweise auch zu Ambivalenzen in der Annotation – durchgeführt. Außerdem wurden begleitend zur Annotation Feldnotizen angefertigt und Lücken bzw. unterrepräsentierte visuelle Kategorien und Merkmale im Datensatz identifiziert und ergänzt. Im Folgenden gehen wir darauf ein, wie durch diese Verfahren die Überlegungen und Konflikte im Prozess der Klassifizierung auf mehreren Ebenen sichtbar gemacht, kollaborativ diskutiert und nachgezeichnet werden konnten.

In den Essays reflektierten die Annotator:innen, wie sie normative Schönheitsideale auf Instagram wahrnahmen und wie ihre Positionalität ihre Einschätzung prägte, die sich im Prozess der Annotation in das Datenset einschrieb. Die Annotator:innen assoziierten normative Schönheitsideale mit Perfektion und Makellosigkeit und beschrieben, dass das Aufwachsen in einer mehrheitlich weißen Gesellschaft ein eurozentristisches Bild von Schönheit vermittelt hat, welches sie hinterfragten und aktiv zu dezentrieren versuchten. Als weiblich gelesene Personen erkannten und benannten die Annotator:innen einerseits körperliche Übereinstimmungen mit gängigen Schönheitsidealen (wie z. B. Jugendlichkeit, Schlankheit), aber auch das Gefühl der Unzulänglichkeit, da sie das Ideal nicht selbst verkörperten (z. B. durchtrainierter Körper, Hautreinheit). In den Essays wird deutlich, dass Einschätzungen von Schönheit einerseits von subjektiven Einstellungen abhängig sind und zwischen Selbst- und Fremdwahrnehmung verhandelt werden. Andererseits zeigen sie, dass diese Einschätzungen intersubjektiv hergestellt werden und von bestimmten sozialen Gruppen – in diesem Fall eine Gruppe von sich als weiblich identifizierenden, weißen und westlich sozialisierten Studentinnen der Empirischen Kulturwissenschaft – geteilt werden.

In der Gruppendiskussion zeigten sich verschiedene Ambivalenzen und Unklarheiten in der Annotation. Ein zentrales Beispiel waren Unterschiede in der Sensibilisierung von Queerness bzw. Transpersonen innerhalb der Gruppe. Diese Ambivalenzen warfen Fragen danach auf, inwiefern normative Schönheitsideale auf einem binären Geschlechtssystem aufgebaut sind und durch Darstellungen von nicht-binären Personen oder Uneindeutigkeit herausgefordert werden. Insbe-

sondere Bilder von durchtrainierten, oberkörperfreien Transmännern mit kaum sichtbaren Narben einer Mastektomie wurden teilweise als stark Schönheitsideale-reproduzierend eingeschätzt und kontrovers diskutiert. Zum Teil erfüllten die Posts alle visuellen Komponenten von Schönheitsideale-reproduzierenden Posts (beispielsweise muskulöser Körper, zentrale Bildposition, erotische Pose, Zeichnungsfilter) und nur ein Transgender Flag Emoji in der Caption lieferte den Kontext, den nicht alle Gruppenmitglieder wahrgenommen hatten oder deuten konnten. Nach sorgfältigem Abwägen beschlossen die Annotator:innen, die Posts trotz ihres Transgender-Kontexts als Schönheitsideale-reproduzierend zu klassifizieren, wenn sie die Mehrheit der in Betracht zu ziehenden Komponenten als erfüllt ansahen.

Während der Gruppendiskussion bestand die Möglichkeit, Änderungen an der Annotation vorzunehmen und Lücken sowie unterrepräsentierte visuelle Kategorien zu identifizieren. Da vier der Annotator:innen nicht in die Herstellung des Korpus involviert waren, wurde der Datensatz um 140 zusätzliche Posts ergänzt, die in der intersubjektiven Wahrnehmung der Annotator:innen unterrepräsentiert waren. Die Erweiterung des Datensets umfasste Posts mit künstlerischer Inszenierung, solche, die das binäre Geschlechtssystem herausfordern, sowie Darstellungen körperlicher Merkmale, wie z. B. Frauenkörper im Kraftsport, vielfältige Haartexturen und -farben, Personen mit Brille oder Sommersprossen. Ziel dieses Verfahrens war, das von Studierenden kuratierte Datenset zu diversifizieren und das durch die Annotation und ihre Situierung erworbene Wissen über die Spezifik des Datensets produktiv einzusetzen.

Ein zweites Beispiel für Ambivalenzen und Unklarheiten in der Annotation sind die Unterschiede zwischen dem im Rahmen des Forschungsprojekts stattfindenden Rezeption von Social Media-Posts (die von den Annotations-Guidelines geleitet und geprägt wurde) im Vergleich zur tatsächlichen Alltagsnutzung der Annotator:innen. So schätzten die Annotator:innen beispielsweise Captions, die Dienstleistungen oder Produkte bewarben und Selbstoptimierung anpriesen, als tendenziell Schönheitsideale-reproduzierend ein. In der Gruppendiskussion stellten sie aber fest, dass Captions für die Einschätzung eines Posts im Prozess der Annotation relevanter sind als in ihrer alltäglichen Social Media-Nutzung. Während in der alltäglichen Nutzung nicht alle Bildunterschriften bewusst wahrgenommen werden, lenkte die Annotations-Guideline die Aufmerksamkeit darauf, reproduzierende Elemente zu identifizieren und einzuschätzen.

Die begleitenden Feldnotizen zeigten außerdem, dass sich mit dem Verlauf der Annotation ein routiniertes Abscannen oder geschultes Sehen und damit auch eine gewisse Blindheit gegenüber Details in Bildern und Captions einstellte, welche die Annotator:innen in der Gruppendiskussion reflektierten. Die Annotator:innen tauschten sich außerdem zu einer subjektiv empfundenen Abstumpfung und Ermüdung aus, die während des repetitiven Prozesses der Annotation und der

Verinnerlichung der Annotations-Guideline auftrat. Sie debattierten dies anhand eines Beispiels: Das Bild zeigt einen gebräunten, durchtrainierten Mann mit einer weißen Badehose, der sein rechtes Knie auf einem durchsichtigen Stuhl abstützt. Die Bildkomposition lenkt den Blick auf die weiße Badehose und den sich darunter abzeichnenden Penis. Während die Annotatorin, die das Beispiel eingebracht hatte, aufgrund der extremen Sexualisierung den Post in ihren Feldnotizen als besonders herausstechend vermerkte, war den anderen Annotatorinnen das entsprechende Detail im Strom der vielen Posts entgangen. Das Beispiel und die sich daran anschließende Diskussion verdeutlichte anschaulich, dass mit zunehmendem Voranschreiten im Datenset die Aufmerksamkeit für Details variierte und die Einschätzung eines Posts in Relation zu vorangegangenen Annotationen erfolgte.

Die ethnografische Situierung des Datensatzes machte also einerseits Lücken im Datenset sichtbar und ermöglichte, diese zu füllen. Andererseits band sie die Annotatorinnen aktiv in die Konstruktion des Datensets ein und unterstützte unseren kollaborativen Ansatz, durch den die sonst oft unsichtbare epistemische Arbeit von Annotator:innen sichtbar gemacht und wertgeschätzt wurde. Entscheidend war über alle Arbeitsschritte hinweg, dass die hier skizzierten Überlegungen, Konflikte, Unsicherheiten und Ambivalenzen der Kategorisierung und die daran geknüpften Aushandlungsprozesse rund um das Datenset gerade nicht hinter dem fertigen Datensatz und einer vermeintlich objektiven Ground Truth verschwinden sollten. Vielmehr wollten wir herstellen, was wir als *Situated Ground Truth* bezeichnen: Ein Datenset, dessen Herstellungsprozess mit ethnografischen Methoden ausführlich dokumentiert und kritisch reflektiert wurde, wodurch die Positionalität der Annotator:innen einbezogen und die vom Datenset generierte Objektivität als eine verkörperte Objektivität im Sinne Haraways sichtbar gemacht werden kann.

Im Falle unseres Projekts wurde in diesem Prozess wesentlich besser verstanden, wie die jeweils subjektiven Einschätzungen einzelner Annotator:innen zustande kamen, aber auch, wie die Annotator:innen geteilte bzw. intersubjektive Wissensbestände kollaborativ sichtbar zu machen versuchten. Dabei zeigten sich zahlreiche Ambivalenzen und Unklarheiten: von den unterschiedlichen Einschätzungen bei der Klassifikation queerer Körperbilder, über die Differenzen zwischen alltäglicher Nutzung und zielgerichteter Klassifikation, bis hin zu Ermüdungserscheinungen und Veränderungen routinierter Wahrnehmung, welche die Klassifikation prägen. In einem von Crowdworker:innen klassifizierten Datenset erhalten solche Probleme und Herausforderungen im Regelfall keine Aufmerksamkeit. Sie sind aber zentral, um die Validität der Ground Truth adäquat einschätzen zu können. Als Konsequenz können einige der Probleme ggf. durch die Ergänzung des Datensets adressiert, aber niemals vollständig gelöst werden. Denn die ethnografische Situierung einer intersubjektiven Ground Truth macht vor allem sichtbar, dass letztere eben keine objektive Wahrheit repräsentiert, sondern eine von vielen

Spannungen, Ambivalenzen und Unklarheiten gezeichnete Perspektive einer aus individuellen Menschen bestehenden Gruppe.

Fazit

Im vorliegenden Artikel haben wir anhand einer Kooperation zwischen Informatik und Digitaler Anthropologie zwei unterschiedliche Konzepte vorgeschlagen und diskutiert. Das Konzept des Intersubjective Classifiers bezeichnet ein KI-basiertes Tool, das die intersubjektive Einschätzung einer Gruppe menschlicher Akteur:innen (in unserem Fall mit Blick auf spezifische Qualitäten von Social Media-Posts) relativ zuverlässig automatisiert simulieren kann. Das Konzept der Situated Ground Truth beschreibt ein Verfahren, mit dem die Herstellung von Datensets im Sinne Donna Haraways ethnografisch situiert werden kann, um die Spannungen, Ambivalenzen und Unklarheiten sichtbar zu machen, die diesen Herstellungsprozess prägen.

Beide Konzepte verhalten sich relational zueinander. Grundsätzlich erscheint uns die Situierung von Ground Truth gewissermaßen die Voraussetzung dafür zu sein, um überhaupt von einem intersubjektiven Classifier sprechen zu können. Nehmen wir zum Vergleich einen Classifier, der den Unterschied von Hund und Katze erkennt, dann mag auch hier eine Situierung der Ground Truth hilfreich und gewinnbringend sein – sie ist aber nicht in gleichem Maße grundlegend. Bei einem intersubjektiven Classifier, der bewusst mit nicht-objektivierbaren Kategorien wie Schönheit operiert, ist es aus unserer Sicht hingegen absolut essenziell, die Herstellung der entsprechenden Ground Truth kritisch zu reflektieren. Dadurch wird einerseits ein Mehrwert erzeugt, indem blinde Flecken erkannt und fehlende Daten ergänzt werden können. Vor allem aber hilft dieser Ansatz dabei, den epistemischen Mehrwert, aber genauso die epistemischen Grenzen des Classifiers adäquat einzuschätzen. In Konsequenz macht erst die Situierung der Ground Truth diesen Classifier zu einem Intersubjective Classifier, denn erst in ihr wird sichtbar, ob und inwiefern dieser eine intersubjektive Einschätzung angemessen wiedergeben kann. Die Konzepte des Intersubjective Classifiers und der Situated Ground Truth bedingen sich hier also gegenseitig.

Damit kehren wir abschließend noch einmal zur übergeordneten Frage nach dem Mehrwert von Schnittstellen zwischen Informatik und Digitaler Anthropologie zurück. Was wir mit den hier vorgestellten Konzepten zeigen möchten, sind potenzielle Synergien zwischen den beiden Disziplinen. Diese erfordern zugegebenermaßen, dass alle Beteiligten bereit sind, über so manchen epistemologischen Schatten zu springen und sich auf die epistemischen Logiken anderer Fächer einzulassen. Wir hoffen aber, zeigen zu können, dass damit ein potenzieller Mehrwert verbunden ist. In unserem konkreten Fall ermöglichte uns die Kombination der Zugänge, eine für unsere digitalanthropologische Forschung relevante Fragestellung

– hier die Frage nach den Dynamiken der Sortieralgorithmen von Social Media-Feeds mit Blick auf die Reproduktion von Schönheitsidealen – zumindest in Teilen beantworten zu können. Insofern die algorithmischen Sortiersysteme von Social Media-Plattformen ein zunehmend einflussreiches Element alltagskultureller Aushandlung werden, halten wir das angewandte Verfahren auch über die Frage nach Schönheitsidealen hinaus für vielversprechend – anwendbar wäre es beispielsweise auch im Bereich der politischen Meinungsbildung, die zunehmend abhängig von den Infrastrukturen digitaler Plattformen wird (Bareither et al. 2023; Postill 2021).

Gleichzeitig hoffen wir, dass dieses Verfahren auch für die Perspektive der Informatik gewinnbringend sein kann. Denn wir zeigen erstens, dass eine Klassifikation und relativ zuverlässige Reproduktion intersubjektiver Einschätzungen mit aktuellen KI-Modellen grundsätzlich möglich ist. Zweitens adressieren wir aber auch ein Problem, das in der Informatik zunehmend Beachtung findet: Wie oben bereits beschrieben, muss die Informatik als gesichert geltende Datengrundlagen (Ground Truths) voraussetzen, um mit ihren Verfahren mathematische Problemstellungen effektiv adressieren zu können. Das heißt aber nicht, dass die Informatik dieser Notwendigkeit zwangsläufig unkritisch gegenübersteht. In den letzten Jahren gibt es verstärkte Tendenzen innerhalb der Disziplin, die Validität von Datensets nicht einfach als gegeben hinzunehmen, sondern sie kritisch zu hinterfragen (Denton et al. 2020, 2021; Buolamwini und Gebru 2018). Mit unserem Projekt demonstrieren wir anhand eines konkreten Fallbeispiels, wie eine ethnografische Situierung dabei helfen kann, die Herstellung von Datensets kritisch zu begleiten und damit auch den epistemischen Mehrwert einer durch dieses Datenset generierten Ground Truth adäquat zu reflektieren und im Idealfall zu erweitern. Zwar ist diese Art der Situierung aufwendig und dementsprechend kostspielig – und deshalb ist sie sicher nicht in allen Fällen leistbar. Aber gerade dort, wo es sich um besonders komplexe Datensets handelt, in denen zentrale Kategorien nur intersubjektiv bestimmbar sind, scheint uns der hier vorgeschlagene Ansatz vielversprechend. Er könnte dazu beitragen, dass auch die Informatik von digitalanthropologischen bzw. sozial- und kulturwissenschaftlichen Ansätzen verstärkt profitieren kann.

Literatur

- Archer, Alfred, und Lauren Ware. 2018. „Beyond the Call of Beauty: Everyday Aesthetic Demands Under Patriarchy“. *The Monist* 101 (1): 114–27. <https://doi.org/10.1093/monist/onx029>.
- Artstein, Ron. 2017. „Inter-Annotator Agreement“. In *Handbook of Linguistic Annotation*, herausgegeben von Nancy Ide und James Pustejovsky. Springer Netherlands: 297–313. https://doi.org/10.1007/978-94-024-0881-2_11.

- Bareither, Christoph, Alexander Harder, und Dennis Eckhardt. 2023. „Special Issue: Digital Truth-Making: Anthropological Perspectives on Right-Wing Politics and Social Media in ‚Post-Truth‘ Societies“. *Ethnologia Europaea* 53 (2). <https://doi.org/10.16995/ee.9594>.
- Bareither, Christoph, und Sabine Wirth. 2025. „Social Media Feeds as Curatorial Assemblages: A Conceptual Framework.“ *Media, Culture & Society*: 1–16. <https://doi.org/10.1177/01634437251360372>.
- Bennett, S. J., Benedetta Catanzariti, und Fabio Tollon. 2025. „Everybody Knows What a Pothole Is: Representations of Work and Intelligence in AI Practice and Governance“. *AI & SOCIETY*. <https://doi.org/10.1007/s00146-024-02162-0>.
- Berger, John. 2008. „*Ways of Seeing*.“ British Broadcasting Corporation and Penguin Books.
- Berger, Peter L., und Thomas Luckmann. 2016. „*Die gesellschaftliche Konstruktion der Wirklichkeit*.“ 26. Aufl. Fischer-Taschenbuch-Verl.
- Bordo, Susan. 2003. „*Unbearable Weight: Feminism, Western Culture, and the Body*.“ Tenth Anniversary Edition. University of California Press.
- Brathwaite, Kyla N, David C DeAndrea, und Megan A Vendemia. 2023. „Non-Sexualized Images and Body-Neutral Messaging Foster Body Positivity Online“. *Social Media + Society*: 1–15. <https://doi.org/doi.org/10.1177/20563051231207852>.
- Brey, Philip. 2000. „Technology and Embodiment in Ihde and Merleau-Ponty“. In *Metaphysics, Epistemology, and Technology. Research in Philosophy and Technology*, herausgegeben von Carl Mitcham: 45–58. Elsevier/JAI Press.
- Bucher, Taina. 2018. „*If...Then: Algorithmic Power and Politics*.“ Oxford University Press.
- Buolamwini, Joy, und Timnit Gebru. 2018. „Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification“. *Proceedings of the 1st Conference on Fairness, Accountability and Transparency*: 77–91. <https://proceedings.mlr.press/v81/buolamwini18a.html>.
- Burke, Robin, Alexander Felfernig, und Mehmet H. Göker. 2011. „Recommender Systems: An Overview“. *AI Magazine* 32 (3): 13–18. <https://doi.org/10.1609/aimag.v32i3.2361>.
- Crawford, Kate, und Trevor Paglen. 2021. „Excavating AI: The Politics of Images in Machine Learning Training Sets“. *AI & SOCIETY*: 1105–1116. <https://doi.org/10.1007/s00146-021-01162-8>.
- Denton, Emily, Alex Hanna, Razvan Amironesei, Andrew Smart, und Hilary Nicole. 2021. „On the Genealogy of Machine Learning Datasets: A Critical History of ImageNet“. *Big Data & Society* 8 (2): 1–14. <https://doi.org/10.1177/20539517211035955>.
- Denton, Emily, Alex Hanna, Razvan Amironesei, Andrew Smart, Hilary Nicole, und Morgan Klaus Scheuerman. 2020. „Bringing the People Back In: Contesting Benchmark Machine Learning Datasets“. *Proceedings of ICML Workshop on Participatory Approaches to Machine Learning*. arXiv. <http://arxiv.org/abs/2007.07399>.

- Fabian, Johannes. 2014. „Ethnography and Intersubjectivity: Loose Ends“. *HAU: Journal of Ethnographic Theory* 4 (1): 199–209. <https://doi.org/10.14318/hau4.1.008>.
- Franken, Lina. 2022. „Digitale Methoden für qualitative Forschung: Computationelle Daten und Verfahren.“ utb GmbH. <https://doi.org/10.36198/9783838559476>.
- Gill, Rosalind. 2021. „Changing the perfect picture: Smartphones, social media and appearance pressures.“ University of London. https://www.citystgeorges.ac.uk/__data/assets/pdf_file/0005/597209/Parliament-Report-web.pdf.
- Gillespie, Tarleton. 2017. „#trendingistrending. Wenn Algorithmen Zu Kultur Werden“. In *Algorithmenkulturen: Über die rechnerische Konstruktion der Wirklichkeit*, herausgegeben von Robert Seyfert und Jonathan Roberge. Transcript Verlag. <https://doi.org/10.1515/9783839438008>.
- Gollub, Tim, Pierre Achkar, Martin Potthast, und Benno Stein. 2025. „Call for Research on the Impact of Information Retrieval on Social Norms“. In *Advances in Information Retrieval. 47th European Conference on IR Research (ECIR 2025)*, herausgegeben von Claudia Hauff, Craig Macdonald, Dietmar Jannach, et al., vol. 15575. Lecture Notes in Computer Science. Springer. https://doi.org/10.1007/978-3-03-1-88717-8_27.
- Grandinetti, Justin. 2023. „Examining Embedded Apparatuses of AI in Facebook and TikTok“. *AI & SOCIETY* 38 (4): 1273–1286. <https://doi.org/10.1007/s00146-021-01270-5>.
- Haraway, Donna. 1995. „Die Neuerfindung der Natur: Primaten, Cyborgs und Frauen.“, herausgegeben von Carmen Hammer und Immanuel Stieß. Campus Frankfurt / New York.
- Henriksen, Anne, und Anja Bechmann. 2020. „Building Truths in AI: Making Predictive Algorithms Doable in Healthcare“. *Information, Communication & Society* 23 (6): 802–816. <https://doi.org/10.1080/1369118X.2020.1751866>.
- Jocher, Glenn, und Jing Qiu. 2024. „Ultralytics YOLO11.“ V. 11.0.0. Released. <https://github.com/ultralytics/ultralytics>. <https://docs.ultralytics.com/models/yolo11/#what-are-the-key-improvements-in-ultralytics-yolo11-compared-to-yolov8>.
- Kitchin, Rob, und Tracey P. Lauriault. 2018. „Towards critical data studies: Charting and unpacking data assemblages and their work“. In *Thinking Big Data in geography: New regimes, new research*. University of Nebraska Press.
- Mascia-Lees, Frances E., Hg. 2011. „A Companion to the Anthropology of the Body and Embodiment.“ Blackwell Companions to Anthropology 13. Wiley-Blackwell.
- Mascia-Lees, Frances E. 2019. „The Body and Embodiment in the History of Feminist Anthropology: An Idiosyncratic Excursion through Binaries“. In *Mapping Feminist Anthropology in the Twenty-First Century*, herausgegeben von Ellen Lewin und Leni M. Silverstein. Rutgers University Press: 146–67. <https://doi.org/10.36019/9780813574318>.

- Meta AI. 2023. „Introducing 22 system cards that explain how AI powers experiences on Facebook and Instagram“. <https://ai.meta.com/blog/how-ai-powers-experiences-facebook-instagram-system-cards/>.
- Peeters, Stijn. 2025. „Zeeschuimer.“ 25. April 2025. <https://doi.org/10.5281/zenodo.15281614>.
- Postill, John. 2021. „Digital politics“. In *Digital Anthropology*, herausgegeben von Hannah Knox und Haidey Geismar: 159–177. Routledge.
- Rumelhart, David E., Geoffrey E. Hinton, und Ronald J. Williams. 1986. „Learning Representations by Back-Propagating Errors.“ *Nature* 323 (6088): 533–536. <https://doi.org/10.1038/323533a0>.
- Scheuerman, Morgan Klaus, Kandrea Wade, Caitlin Lustig, und Jed R. Brubaker. 2020. „How We’ve Taught Algorithms to See Identity: Constructing Race and Gender in Image Databases for Facial Analysis“. *Proceedings of the ACM on Human-Computer Interaction* 4 (CSCW1): 1–35. <https://doi.org/10.1145/3392866>.
- Thylstrup, Nanna Bonde. 2022. „The Ethics and Politics of Data Sets in the Age of Machine Learning: Deleting Traces and Encountering Remains“. *Media, Culture & Society* 44 (4): 655–671. <https://doi.org/10.1177/01634437211060226>.
- Thylstrup, Nanna Bonde, Mikkel Flyverbom, und Rasmus Helles. 2019. „Datafied Knowledge Production: Introduction to the Special Theme“. *Big Data & Society* 6 (2): 205395171987598. <https://doi.org/10.1177/2053951719875985>.
- TikTok Support. o. J. „How TikTok recommends content“. Zuletzt abgerufen am 12. Februar 2025. <https://support.tiktok.com/en/using-tiktok/exploring-videos/how-tiktok-recommends-content#>.
- White, Bob W., und Kiven Stroh. 2014. „Preface: Ethnographic Knowledge and the Aporias of Intersubjectivity“. *HAU: Journal of Ethnographic Theory* 4 (1): 189–197. <https://doi.org/10.14318/hau4.1.007>.
- Widdows, Heather. 2018. „Perfect Me: Beauty as an Ethical Ideal.“ Princeton University Press.