

Talk Science to Me #64: Umweltschutz mit KI?

3.6.2026 - Birgit Baustädter | Technische Universität Graz

Können Recommender Systems den Klimaschutz unterstützen, während sie uns Urlaubsreisen suchen? Dieser und vielen weiteren Fragen rund um künstliche Intelligenz geht Alexander Felfernig nach.

Dieser Text ist ein Transkript der Podcast-Folge und wurde im Sinne der Verständlichkeit leicht angepasst.

Können Recommender Systems die Sustainable Development Goals unterstützen, während sie uns maßgeschneiderte Urlaubsreisen suchen? Dieser und vielen weiteren Fragen rund um künstliche Intelligenz geht Alexander Felfernig am Institute of Software Engineering and Artificial Intelligence nach. Mein Name ist Birgit Baustetter und das ist Talk Science to Me, der Wissenschaftspodcast der TU Graz.

Bitte geben Sie mir als erstes einen kurzen Eindruck, woran Sie aktuell arbeiten.

Alexander Felfernig: Also wir arbeiten aktuell und im Endeffekt schon ziemlich lange an diesem Thema Bilateral Software Engineering, nicht das, was sozusagen heutzutage in aller Munde ist, bilateral AI, wo sozusagen unterschiedliche AI-Methoden und Welten miteinander kombiniert werden. Sondern uns geht es darum, die Synergieeffekte zwischen Software Engineering und AI zu nutzen. Softwareentwicklung ist ein ganz wesentliches Thema. Software wird überall gebraucht. Und da ist es ganz wichtig, dass die Software kostengünstig entwickelt wird einerseits und andererseits auch Software entwickelt wird, die von Personen auch genutzt werden kann und effektiv genutzt werden kann. Und da gibt es die Möglichkeit, unterschiedliche AI-Methoden im Softwarebereich einzusetzen, in der Softwareentwicklung einzusetzen, um diese ganzen Prozesse einfach viel günstiger zu machen und schneller zu machen. Das ist die eine Richtung.

Aber genauso gut kann die Entwicklung von AI-Systemen vom Software Engineering profitieren, weil es natürlich auch zum Beispiel im Software Engineering schon seit Jahrzehnten unterschiedliche Formen von Testmethoden gibt. Und AI-Systeme müssen auch getestet werden auf der einen oder anderen Art und Weise. Ist ja extrem wichtig, dass man sich darauf verlassen kann, was da rauskommt, dass das auch Sinn macht. Und da kann man wieder Software-Methoden einsetzen. Für uns ist das wieder so eine Hybridisierung. Also das heißt, die wesentlichen Stärken von beiden Disziplinen zu nutzen, miteinander zu integrieren, um Systeme in der Zukunft zu schaffen. Und es gibt ja auch in Graz so ein Projekt, das nennt sich Bilateral AI. Es fokussiert aber unter Führungszeichen nur auf diesen Bereich Künstliche Intelligenz. Ignoriert oder sozusagen berücksichtigt nicht das, was sozusagen im Software-Engineering-Bereich gemacht wird.

Aber versteht es jetzt richtig, nutzen Sie künstliche Intelligenz als Tool für Software Engineering oder entwickeln Sie also Künstliche-Intelligenz-Tools, die dann das Software Engineering unterstützen?

Felfernig: Das hat unterschiedliche Dimensionen. Also man kann Künstliche-Intelligenz-Tools oder Standardalgorithmen oder Large Language Models verwenden, um sozusagen unterschiedliche Aufgabenstellungen im Software-Engineering-Bereich besser zu unterstützen. Da sind es Tools, aber

gleichzeitig geht es natürlich auch darum, dass man mit der künstlichen Intelligenz Algorithmen schneller macht. Dass sozusagen auch die Energiekonsumtion reduziert wird und so weiter. Also wir arbeiten auch auf der algorithmischen Ebene, um Algorithmen einfach effizienter zugestalten. Da gibt es beispielsweise unterschiedliche Formen von Diagnose-Algorithmen, wo es darum geht, herauszufinden, wo sind Fehler in komplexen Systemen, zum Beispiel Atomkraftwerke. Wir arbeiten nicht mit Atomkraftwerken, das war ein nettes Beispiel, weil es sofort klar macht, dass es da wichtig ist, dass die AI-Systeme, die in dem Bereich verwendet werden, auch tatsächlich korrekt sind. Und wir arbeiten sozusagen in diesem Bereich modellbasierte Diagnose und da geht es auch darum, Algorithmen in dem Bereich schneller und effizienter zu machen. Es hängt vom Anwendungsbereich ab. Es ist aber genauso gut auch der Algorithmus selbst, der dann entsprechend verbessert werden muss. Und das sind dann auch wieder unsere Beiträge zur Grundlagenforschung, wo wir so in internationalen Spitzenkonferenzen publizieren.

Wie muss ich mir das vorstellen? Sie sagen, Sie machen einen Algorithmus energieeffizienter. Wie geht denn das? Welche Möglichkeiten haben Sie da?

Felfernig: Jetzt nehmen wir zum Beispiel her, wie wird Software produziert oder ausgeführt. Da hat man irgendwo ein Tool, das übersetzt sozusagen einen Quellcode in etwas, das vom Computer ausführbar ist. Und dieses etwas, das die Übersetzung macht, das nennt sich Compiler. Und Compiler funktionieren halt nach bestimmten Parametern. Also ein Compiler, ein C-Compiler hat 100, 200 Parameter, die irgendwann mal gesetzt werden müssen und abhängig davon, wie diese Parameter gesetzt werden, kommt da irgendeine Software raus, die dann energieeffizient funktioniert oder nicht energieeffizient funktioniert. Und das, was dazwischen ist, oder was man dazwischen reinbringen kann, ist AI, Machine Learning, die sozusagen aufgrund der Struktur von dem Computerprogramm versteht, wie diese Parameter ausschauen sollen, damit die Software effizient funktioniert. Das heißt, ein Problem schneller löst. Es geht ja im AI-Bereich sehr oft darum, optimale Lösungen zu finden, aber wichtig dabei ist auch zu berücksichtigen, dass diese optimalen Lösungen schnell gefunden werden. Und dass es nicht, also im Extremfall, Wochen dauert, bis man irgendwann einmal die Lösung findet. Nett wäre es, wenn es in ein paar hundert Sekunden passiert, oder in ein paar Sekunden passiert. Und genau an dem arbeiten wir. Und das kann man mit AI machen.

Geht es Ihnen da, wenn es jetzt um Energieeinsparungen geht, ein bisschen um den Umweltaspekt? Oder geht es Ihnen um Performance? Was ist denn so der Hintergrund?

Felfernig: Also Performance korreliert ja immer mit Energieeinsparung. Wenn es jetzt um den persönlichen Aspekt geht, dann geht es natürlich auch um den Umweltaspekt. Aber es geht natürlich auch um die Performance, um sozusagen für einen bestimmten Anwendungsbereich effiziente Lösungen zu bringen. Also wenn Benutzer*innen Wochen lang auf eine Lösung warten müssen ist in unterschiedliche Richtungen nicht wirklich optimal. Firmen werden es nicht kaufen, Benutzer*innen werden es nicht verwenden und das ist natürlich auch, was den Umweltaspekt betrifft, was den Nachhaltigkeitsaspekt betrifft, ziemlich problematisch. Das heißt, Effizienz ist extrem wichtig und mit AI und Machine Learning hat man sozusagen das Potenzial, in dem Bereich wirklich sehr gute Lösungen zu bauen. Man muss natürlich auch berücksichtigen, welche AI verwendet wird, um diese Lösungen zu bauen, weil mittlerweile das, was man unter AI versteht, ist sehr oft reduziert auf dieses Thema Large Language Models und die haben durchaus genau diese Problematik, dass sie extrem viel Energie brauchen. Also geht es darum, wirklich zu überlegen, brauche ich diese Large Language Models jetzt, um das konkrete Problem zu lösen? Und wenn ich diese Large Language Models verwende, wie verwende ich sie, dass möglichst wenig Energie verbraucht wird.

Unser Thema dieser Staffel ist ja auch, was hat sich in den vergangenen Jahren im Bereich Artificial Intelligence getan. Sie haben jetzt schon die Large Language Models erwähnt. Das ist ja irgendwie

so das, was man am meisten sieht oder am meisten spürt im Alltag. Was sagen Sie, hat sich so in den letzten, sagen wir, vier Jahren getan im Bereich Artificial Intelligence?

Felfernig: Es gibt positive Veränderungen, es gibt negative Veränderungen. Die positive Veränderung ist, dass diese Large Language Models natürlich Routinearbeiten ersetzen, einfacher machen. Man kann produktiver arbeiten durch den Einsatz dieser Technologien. Was manchmal missverstanden wird, dass diese Large Language Models eine ganz tolle und intelligente AI sind. Aber wenn man die Technologie dahinter versteht, merkt man relativ schnell, dass genau das nicht der Fall ist, weil Large Language Models basieren auf dem Prinzip, dass einfach nur immer das nächste Wort auf Basis von dem, was bis jetzt generiert worden ist, vorhergesagt wird. Das heißt, das, was dahinter ist, ist hochgradig wahrscheinlichkeitsbasiert. Das heißt, man muss relativieren. Also die AI hat sich weiterentwickelt. Ein wesentlicher Punkt dabei ist natürlich auch die Rechenleistung, die dahinter ist, die sich sukzessive enorm verbessert und mit dieser Rechenleistung, mit den Datenquellen, die sozusagen verwendet werden, um diese AI zu trainieren, verbessert sich natürlich auch die Vorhersagequalität der AI. Also es ist ein Unterschied. Vor zehn Jahren hat man noch AI-Systeme gehabt, die basieren auf lokalen Datensätzen. Heutzutage auch aufgrund der Rechenleistung hat man AI-Systeme, die basieren, was die Datensätze betrifft, auf dem Weltwissen. Ja, und das ist natürlich ein wesentlich größerer Datensatz, um diese AI zu trainieren, auf einfach nur lokale Datensätze. Damit steigert sich auch die Output-Qualität der AI. Aber man muss es trotzdem relativieren. Ja, es gibt ein großes Potenzial, das da dahinter ist, aber man ist natürlich weit davon entfernt, nur auf irgendeine Art und Weise menschliche Intelligenz nachzubilden. Also da ist man auch sehr, sehr weit davon entfernt. Das wird irgendwann einmal passieren. Hängt einerseits ab von den Kapazitäten, die irgendwann einmal bestehen werden und andererseits natürlich auch von der Forschung, die gefordert ist, qualitativ bessere, intelligentere AI-Architekturen zu entwickeln, die dann ein intelligentes Verhalten, so wie man es heute von Menschen kennt, entsprechend nachbilden.

Das ist jetzt eine sehr spannende Frage. Was heute als künstliche Intelligenz gilt und was Sie jetzt als die tatsächliche Intelligenz erwähnt haben, also menschliche Intelligenz, was fährt da noch dazwischen? Also wo ist da noch der Aufholbedarf quasi oder Entwicklungsbedarf?

Felfernig: Das sind genau die zwei Punkte. Also wahrscheinlich sind es mehrere Punkte. Aber einerseits ist es sicher mal die Rechenteistung, weil einfach sozusagen die größten Supercomputer der Welt das noch nicht schaffen, diese ganze Verstricktheit des menschlichen Gehirns entsprechend nachzubilden. Das ist mit Hardware aktuell noch nicht möglich. Und auf der anderen Seite ist die Forschung gefragt, entsprechende intelligentere, verbesserte Lernarchitekturen zu entwickeln. Also Lernarchitekturen, das sind sozusagen Architekturen oder Systeme, die dabei helfen, ein AI-System zu bauen. Die Architekturen müssen verbessert werden. Also wenn man die Architektur nicht verbessert, wird man sozusagen nie zu einer wirklich absolut intelligenten AI finden. Also das sind aus meiner Sicht einmal zwei wesentliche Komponenten. Also einerseits die Logik, wie so ein AI-System gebaut wird, und andererseits die entsprechende Rechenkapazität, die einfach jetzt noch nicht vorhanden ist. Und man sieht ja die Limitierung von AI-Systemen, die bestehen. Also es gibt sowas wie Halluzinieren und so weiter. Der Grund ist einfach, dass sozusagen einfach nur stupide, immer wieder Texte vorhergesagt werden. Das hat bei weitem noch nichts mit einem intelligenten Verhalten zu tun.

Aber ich habe ein bisschen rausgehört, dass Sie schon der Meinung sind, dass es das einmal geben wird.

Felfernig: Ja, absolut.

Gehen wir wieder zurück zu Ihrem ganz höchstpersönlichen Forschungsgebiet. Sie arbeiten ja auch

stark mit Recommender Systems. Was ist das genau? Können Sie mir das erklären?

Felfernig: Also ich möchte jetzt da keine Definition bringen, aber wenn man es allgemein beschreibt, sind es Entscheidungsunterstützungssysteme. Und da gibt es unterschiedliche Beispiele. Ein Entscheidungsunterstützungssystem ist beispielsweise der Recommender in Amazon, also sehr industriegetrieben, nämlich diese Vorhersagequalität immer weiter zu verbessern. Also Amazon oder Netflix, die Firmen wollen einfach wissen, was sind die Produkte, die für bestimmte Kunden interessant sind und die werden dann auch entsprechend angeboten. Das nennt man auch so the dark side of Recommender Systems, weil es einfach nur darum geht, immer wieder Kund*innen dazu zu bringen, wieder neue Produkte zu konsumieren. Genau das machen Recommender Systeme, die beschleunigen im Endeffekt oder erhöhen den Umsatz, den Produktumsatz. Bei Amazon geht man davon aus, dass Recommender Systeme ungefähr für 30 Prozent der Umsatzsteigerung verantwortlich sind. Bei Netflix habe ich jetzt nicht die genauen Zahlen, aber genau in diese Richtung geht es. Und das hat natürlich negative Konsequenzen, weil Leute einfach dazu gebracht werden, Produkte zu kaufen, die sie vielleicht gar nicht brauchen oder gar nicht wollen. Und deswegen sagt man auch die dark side of Recommender Systems.

Und andererseits, weil es ja Entscheidungs- und Unterstützungssysteme sind, gibt es auch die Möglichkeit, das Ganze ins Positive zu bringen. Das ist ein bisschen mein Background bei dem Ganzen. Also ich bin zu dem Thema Recommender Systeme gekommen durch Anwendungen im Bereich Financial Services. Also da ist es darum gegangen, Kund*innen zu betreuen, welches Finanzportfolio ist für bestimmte Familiensituationen, also Familien mit drei Kindern und so weiter, wirklich geeignet. Also in welche Produkte sollte die Familie investieren, ist es eher blöd beispielsweise ein teures Auto zu kaufen und gleichzeitig ein neues Haus zu finanzieren und solche Dinge. Und da sagt das Recommender System schon, das ist geeignet, da würde ich eher aufpassen und in die Richtung sollte das gesamte Finanzportfolio gehen. Also da geht es eher in die positive Richtung. Das ist das, was ich vorher schon erwähnt habe, wo es darum geht, Software zu entwickeln. Recommender Systeme und generell AI helfen dabei, jetzt diese Software effizienter zu machen. Also es wird dann weniger Strom konsumiert für die Ausführung einer bestimmten Logik. Das sind positive Dinge und das hat auch mit diesen United Nations Nachhaltigkeitskriterien zu tun. Da geht es in eine positive Richtung. Aber generell haben diese Recommender Systeme eben zwei Seiten. Die dunkle Seite und die positive Seite.

Sie haben es jetzt schon erwähnt, die Nachhaltigkeitsziele der UN. Wie können Recommender Systems die unterstützen?

Felfernig: Durch Beratung, durch Entscheidungsunterstützung. Also beispielsweise Financial Services. Das ist ja durchaus auch ein Bereich bei diesen Sustainable Development Goals der United Nations. Also Leute sollen nicht in eine Schuldenfalle tappen. Das hat dann Konsequenzen fürs Berufsleben und so weiter. Das sind Dinge, die vermieden werden sollen. Wir empfehlen das System, helfen dabei und bieten entsprechende Beratung. Energieberatung ist auch anders. So typisch, wenn man an Nachhaltigkeit denkt, dann geht es sofort in Richtung Energie. Also welche Konfiguration von Auto, Photovoltaikanlage ist für mich überhaupt geeignet? Oder soll ich jetzt weiter meinen Diesel fahren oder macht es vielleicht doch Sinn, in andere Technologien irgendwie umzusteigen und Recommender Systeme bieten einfach Lösungen, indem erklärt wird für die aktuelle Situation, was ist geeignet, was sollte man machen, was sollte man eher nicht machen. Und da gibt es eben bei den United Nations 17 unterschiedliche Sustainable Development Goals und die Recommender Systeme können die alle unterstützen. Also wir haben auch mal so ein Überblicks-Paper geschrieben, für welchen Bereich ist ein Recommender System, auf welche Art und Weise geeignet. Und was bei diesen Recommender Systemen ein ganz wesentlicher Punkt ist, ist der Aspekt von Erklärungen. Also das ist immer wieder das große Problem von AI-Systemen, dass irgendein Output generiert wird. Du solltest das machen oder jemand sollte einen Kredit bekommen

oder nicht bekommen. Aber das, was bei diesen AI-Systemen und AI-Modellen immer das Problem ist, dass keine Transparenz besteht. Also für Enduser*innen von einem AI-System ist nicht klar, warum das AI-System einen bestimmten Vorschlag macht. Und da kommen diese Erklärungen ins Spiel. Also Recommender Systeme, generell gibt es Technologien, um Erklärungen zu berechnen für einen AI-Output. Also da machen wir auch an der TU Graz eine Lehrveranstaltung zu dem Thema. Aber ohne Erklärungen gibt es kein Vertrauen. Also wenn wir bekommen irgendeinen Vorschlag von einem Recommender System bekommt, den sie nicht verstehen, dann ist das Vertrauen relativ gering und die Wahrscheinlichkeit, dass diese Empfehlung dann umgesetzt wird, ist auch relativ gering. Das heißt, um überhaupt das umzusetzen, was von einem Recommender System empfohlen wird, braucht man Vertrauen und diese Erklärungen sind ein ganz wesentlicher Aspekt dabei, um dieses Vertrauen zu steigern.

Aber weil Sie jetzt vorher schon gesagt haben, dass man die Systeme eben einsetzt, um diese Ziele zu unterstützen. Sind wir da im Bereich der Beeinflussung oder wie kann man das dann irgendwie trennen?

Felfernig: Es gibt im Wirtschaftsbereich dieses Thema Nudging, ob Sie davon schon gehört haben. Es gibt Nudging und es gibt Persuasion und es gibt eine Grauzone dazwischen. Also Nudging bedeutet, dass sozusagen die Entscheidungsbasis, die Entscheidungsgrundlage für eine Person, die gleiche bleibt, nur die Art und Weise der Aufbereitung der Entscheidungsgrundlage mit den gleichen Parametern ein bisschen anders ist. Nudging ist anstoßen, also jemand wird angestoßen, tendenziell eher eine andere Entscheidung zu treffen. Und dieses Nudging passiert sehr oft auf unterschiedlichen psychologischen Prinzipien des Entscheidungsverhaltens. Da kommt die Entscheidungspsychologie ins Spiel und aus der Entscheidungspsychologie wissen wir, dass Menschen generell nicht optimal entscheiden. Geht nicht. Ja, weil der Grund ist immer noch in der Steinzeit, wenn man sozusagen mit einem Säbelzahn tiger konfrontiert ist, hat man nicht die Zeit drüber nachzudenken, was die Entscheidungsoptionen sind, was man jetzt machen sollte. Und nach zehn Minuten kommt man vielleicht drauf, das wäre jetzt die richtige Entscheidung. Das ist in der Regel schon tot, wenn man so vorgeht. Und deswegen ist das, was in unseren Hirnen passiert, das sind sozusagen Shortcuts, also Abkürzungen, Entscheidungsabkürzungen. Also wenn wir in einer bestimmten Entscheidungssituation sind, ist es sehr oft so, dass wir einfach nicht darüber nachdenken oder irgendeinen optimalen Algorithmus starten, um die beste Entscheidung zu treffen, sondern diese Shortcuts verwenden, um diese Entscheidung zu treffen. Es gibt so hunderte von Shortcuts. Das ist extrem vielfältig und variantenreich. Einfaches Beispiel, das vielleicht eher allgemein bekannt ist, sind diese Köder-Effekte oder Decoy-Effekte. Das heißt, wenn man zwei Produkte präsentiert, gibt man noch ein drittes dazu, das absolut das schlechteste ist, aber eines der anderen beiden Produkte besser darstellt. Und dadurch wird man manipuliert, sozusagen jetzt eher ein bestimmtes Produkt, das sogenannte Zielprodukt oder Budgetprodukt zu kaufen. Das wird auch sehr oft in Supermärkten verwendet, aber da fällt man halt darauf ein, aber die wenigsten kennen die Mechanismen dahinter.

Das ist ja auch ein Forschungsgebiet von Ihnen, die Schnittstelle zwischen Psychologie und künstlicher Intelligenz und Software Engineering.

Felfernig: Also wenn man konfrontiert ist mit Empfehlungen und mit irgendwelchen Erklärungen, ist das immer das, was im Hintergrund ist. Es könnte sein, dass die*der Entwickler*in von dem System eine gewisse Manipulation in das Ganze reinbringt, wie beispielsweise mit diesen Decoy-Effekten. Und im positiven Sinne kann man eben diese unterschiedlichen psychologischen Theorien dazu verwenden, um Nudging zu betreiben, also jemanden anzustupfen, vielleicht nicht diesen BMW zu kaufen, vielleicht ein günstigeres Auto zu kaufen, weil es ja keinen Sinn macht, die Familie irgendwie in finanzielle Risiken reinzubringen und solche Dinge. Also es geht in beide Richtungen. Es gibt sozusagen unterschiedliche Modelle, was menschliches Entscheidungsverhalten betrifft und

die Modelle muss man kennen, wenn man AI-Systeme baut. Und natürlich abhängig von der Zielsetzung dieses AI-Systems kann es in eine positive Richtung gehen. Aber leider muss man auch sagen, dark side of AI, das Ganze kann genauso gut in eine negative Richtung gehen.

Wir haben jetzt auch schon darüber gesprochen, dass Sie sagen, künstliche Intelligenz kann beim Software Engineering helfen. Da geht es ja auch ums Requirement Engineering zum Beispiel. Was ist das genau und wie kann künstliche Intelligenz dabei unterstützen? Also was ist Requirement Engineering?

Felfernig: Da geht es darum, irgendwer muss definieren, was die Software können soll, was sind die Features, was will man mit der Software tun. Also jede Software hat einen Business Case. Das heißt, die Software muss bei irgendwas unterstützen. Also Software zum Selbstzweck wird ja nie gebaut. Das muss man irgendwo definieren. Was sind die Anforderungen? Was soll die Software können? Was sind die Features dahinter? Und da ist es ganz wichtig, dass genau die Formulierung dieser Anforderungen passt, strukturiert ist und eindeutig ist.

Anforderungen sind typischerweise textuell definiert. Wenn die Qualität der Anforderungen schlecht ist, bedeutet das, die Anforderungen sind nicht eindeutig, sie sind unverständlich, inkonsistent. Das bedeutet, dass die Software, die rauskommt, wahrscheinlich nicht das tut, was Kund*inne gerne hätte. So, jetzt hat man die Möglichkeit, am Anfang von einem Projekt genau in das Thema Requirement Engineering, also Anforderungsmanagement auf Deutsch, zu investieren und sicherzustellen, dass die gewünschten Eigenschaften der Software genau die sind, die die*der Kund*in haben will. Wenn das so ist, wird eine Software gebaut, mit der die*der Kund*in zufrieden ist. Wenn das nicht so ist, wird zuerst eine Software gebaut, mit der die*der Kund*in unzufrieden ist. Dann wird eine neue Software gebaut. Unter Umständen wiederholt sich dieser Prozess zehn oder fünfzehn Mal, bis irgendwann eine Lösung rauskommt, die vielleicht geeignet ist. Das kostet dann das fünfzehnfache. So im Schnitt sagen Studien, wenn man 2% mehr, was die Gesamtkosten eines Softwareprojekts betrifft, in Requirement Engineering investieren würde, wird man sich am Ende 50% der Projektkosten sparen. Das heißt, wenn man sozusagen hunderte von Softwareprojekten betrachtet, haben die im Schnitt ein Einsparungspotenzial von 50%. Das ist Wahnsinn. Überall dort, wo es unter Umständen passieren kann, dass Requirement Engineering nicht optimal ist. Und das ist so. Da gibt es empirische Studien, dass eben speziell in diesem Bereich Requirement Engineering, weil es keine zu 100% technische Disziplin ist, genau in dem Bereich sehr viele Fehler gemacht werden. Da geht es ja darum, Texte zu definieren, die spezifizieren, was soll eine Software können. Also in der Entwicklung haben wir weniger Fehler gemacht, im Requirement Engineering sehr sehr viel. Und viele glauben auch, dass das gar nicht wichtig ist, sondern ich weiß eh schon, was Kund*innen wollen, setze mich hin und entwickle die Software und das funktioniert überhaupt nicht. Also das ist sozusagen ein Riesenpotenzial in dem Bereich, einfach durch Verbesserungen, Strukturierungen, mehr Qualität zu schaffen und das bedeutet für das Gesamtprojekt wesentlich geringere Kosten und in dem Bereich Requirement Engineering spielen Large Language Models eine wesentliche Rolle, weil sie dabei helfen, Texte zu analysieren, Inkonsistenzen in Texten zu erkennen, Texte zu vereinfachen, vor allem auch Texte zu personalisieren. Das ist in Softwareprojekten sehr oft so, dass ein unterschiedlicher Wissensstand bezüglich unterschiedlicher Technologien besteht. Also man muss vielleicht mit jemandem, der einen betriebswirtschaftlichen Hintergrund hat, die Requirements anders diskutieren als mit Techniker*innen, die diese Requirements dann umsetzen. Und da bieten diese Large Language Models ihr sinniges Potenzial, auch das Management dieser Requirements Texte enorm zu verbessern.

Wie sind Sie eigentlich selbst in dieses Themengebiet gekommen? Also was interessiert Sie daran so?

Felfernig: Der erste Zugang was das Thema Financial Services. Ich komme ursprünglich von der

Uni Klagenfurt und wir haben dort so Kooperationen gestartet im Bereich Financial Services, wo es genau um diese kundenzentrierte Finanzberatung gegangen ist. Und da ist eben rausgekommen, da geht es ja nicht nur darum, dass sozusagen einfach festgestellt wird, was Kund*innen wollen, sondern Kund*innen müssen einmal die Thematik selbst verstehen, bevor sie überhaupt eine Entscheidung treffen können, in welche Richtung es weiter geht. Und das war sozusagen der Zugang zu dem Thema Recommender Systeme. Und Recommender Systeme sind Software. Das heißt, wir haben dann auch in dem Bereich eine eigene Firma gegründet, damals in Klagenfurt, die solche Recommender Systeme entwickelt und vertrieben hat. Also das war der erste Zugang zu dem Thema Recommender Technologien. Und eben dieser wesentliche Punkt, man muss einfach mal verstehen, was Kund*innen wollen. Und Kund*innen müssen selbst verstehen, was sie wollen. Und da sind Recommender Systeme sozusagen eine perfekte Spielwiese, auch aus Forschungssicht.

Wie war Ihr Werdegang bis zu diesem Punkt?

Felfernig: Also ich habe an der Universität Klagenfurt studiert, ich habe dort mein Doktorat gemacht, habe auch in Klagenfurt habilitiert und bin 2009 nach Graz gekommen. Das ist der Werdegang. Und habe aber schon in Klagenfurt an dem Thema Recommender Systeme und Entscheidungsunterstützung gearbeitet. Also habe das dann einfach in Graz weiterentwickelt.

Gibt es für Sie irgendeine Forschungsfrage, die Sie gerne beantworten würden?

Felfernig: Naja, viele. Also nicht eine, weil da kommt sozusagen das nächste High-Quality-Paper raus. Das, was in unserem Forschungsbereich extrem wichtig ist, da sollte man wirklich früher oder später eine Antwort darauf finden. Also der heilige Gral des Programmierens, nämlich Kund*innen spezifizieren, was sie haben wollen und die Software liefert das Ergebnis. Das ist genau das, was jetzt im Bereich Large Language Models bis zu einem bestimmten Grad schon passiert. Und das muss natürlich weiterentwickelt werden. Also die Idee ist einfach zu sagen, Kund*innen spezifiziere die Anforderungen. Und was rauskommt, ist eine Software, die effizient eine Lösung für genau das anbietet, was die Kund*innen spezifiziert haben. Das ist so eine etwas generellere Fragestellung, an der wir arbeiten.

Vielleicht können wir zum Schluss noch einen kurzen Ausblick machen. In welche Richtung kann es noch gehen? Also was wird in den nächsten Jahren da alles möglich sein?

Felfernig: Also die Möglichkeiten werden sich einerseits durch erhöhte Hardwarekapazität weiterentwickeln. Die Large Language Models werden intelligenter werden. Die AI-Architekturen werden sich ändern und noch intelligentere Lösungen für All-Talks-Probleme liefern, als es jetzt schon der Fall ist. Also in die Richtung geht es.

Speziell in unserem Forschungsbereich geht es immer mehr in Richtung Automatisierung und Effizienzsteigerung. Im Bereich Softwareentwicklung ist einfach eine wesentliche Zielsetzung, unnötigen Aufwand rauszubringen und vor allem aber gleichzeitig auch die Outputqualität zu maximieren. Also Language Models machen jetzt schon eine automatisierte Generierung von Software, aber es ist immer noch eine Qualitätsfrage und eine Frage, wie effizient ist die Lösung, die da rauskommt. Also das ist sozusagen das, was AI in Zukunft machen wird. Allgemein wird es so sein, dass AI wesentlichen Einfluss auf immer mehr Arbeitsbereiche haben wird. Bestimmte Arbeitsbereiche werden wegfallen, Routinearbeiten werden sukzessive wegfallen. Also die Berufsbilder von Menschen werden sich in Zukunft ändern. Also der Rechtsanwalt oder der Notar, wie er heute besteht, wird in Zukunft vielleicht auf andere Art und Weise AI-unterstützt und vielleicht effizienter funktionieren, als es heutzutage der Fall ist. Also speziell im Rechtsbereich gibt es so irrsinnig viel Potenzial, Prozesse zu optimieren. Ja, und das geht natürlich auch in Richtung agentenbasierte Systeme. Also die ganze AI ist übrigens auch interessant, weil das hat vor 20, 30

Jahren auch schon agentenbasierte Systeme in der AI gegeben, nur weiß das niemand mehr und jetzt wird es mittlerweile wieder ein populärer Begriff. Das heißt einfach, dass Agenten, also Software, übernehmen dann integrierte oder ganze Aufgaben und lösen die. Also die Programmierung ist dann nicht mehr, ich entwickle ein System, sondern ich entwickle einen Agenten, der für mich bestimmte Aufgaben erledigt, beziehungsweise auch Agenten, die miteinander kombiniert werden oder miteinander kommunizieren, um eine bestimmte Aufgabe zu lösen. Vielleicht habe ich dann irgendwann einmal ein Reisebüro und einen Reiseführer und die kommunizieren miteinander, um für mich auf die Art und Weise die optimale Reise herauszufinden, was ich jetzt nicht unbedingt will, dass Reisebüros ersetzt werden, aber nur als Beispiel, was es in Zukunft sein kann. Also Agenten mit zunehmender Intelligenz werden immer mehr Routineaufgaben übernehmen. Das ist das, was in Zukunft definitiv passieren wird.

Vielen Dank für das Interview.

Felfernig: Danke für die Möglichkeit.

Vielen Dank, dass ihr heute zugehört habt. In der nächsten Folge spreche ich mit Bettina Könighofer über Bilateral AI.

<https://www.tugraz.at/news/artikel/talk-science-to-me-64-umweltschutz-mit-ki>