

Trainingscenter für smarte Agenten

Intelligente und nachvollziehbare KI-Systeme könnten in Zukunft die Energienetze steuern. Eine vom Bundesforschungsministerium geförderte Nachwuchsgruppe um Eric Veith entwickelt smarte Agenten, die ein ausgeklügeltes Training hinter sich haben – und auch in kritischen Situationen richtig reagieren.

Von Ute Kehse

Der 28. April 2025 wird vielen Menschen in Spanien und Portugal wohl noch lange in Erinnerung bleiben: An diesem Montag um 12.33 Uhr fiel der Strom auf der iberischen Halbinsel komplett aus, die Versorgung konnte teils erst am folgenden Tag wiederhergestellt werden. Ursache war eine zu hohe Spannung im Netz, die eine Kettenreaktion auslöste: Ein Kraftwerk nach dem anderen schaltete sich ab, bis ein kritischer Punkt überschritten war und die Stromerzeugung innerhalb von wenigen Sekunden in beiden Ländern ganz zusammenbrach. Wie das passieren konnte, hat die spanische Regierung in einem Bericht festgehalten. Für den Oldenburger Informatiker Dr. Eric Veith ist klar, dass menschliche Bediener in den Netzleitzentralen in solchen Fällen trotz Unterstützung durch automatische Systeme an ihre Grenzen geraten: „Wenn eine massive Störung auftritt, laufen im Kontrollzentrum unfassbar viele Meldungen ein. Da leuchtet auf den Bildschirmen auf einmal alles auf wie ein Weihnachtsbaum.“ Für Menschen sei es angesichts der Vielzahl möglicher Entscheidungen nahezu unmöglich, binnen Sekunden die richtige Strategie zu entwickeln, um die fatale Kaskade aufzuhalten.

Das deutsche Stromnetz gilt zwar als eins der sichersten und stabilsten weltweit, dennoch erhöhen sich durch die Energiewende und die damit verbundenen stärkeren Schwankungen bei Stromerzeugung und Verbrauch die Anforderungen an die Netzsteuerung. Ein weiteres Risiko stellen Cyberangriffe dar. Veiths Forschungsansatz: Moderne KI-Systeme sollen die Fachleute in den Netzleitzentralen im Alltag unterstützen und so dazu beitragen, dass kritische Infrastruktur widerstandsfähiger gegen unvorhergesehene Ereignisse wird. Mit herkömmlicher Software sei das allerdings schwierig, sagt er. „In modernen Energienetzen gibt es so viele Unwägbarkeiten, so viele Einflüsse, die man bei ihrer Konstruktion noch nicht auf dem Schirm hatte, dass es gar nicht mehr möglich ist, Softwarekom-

ponenten zu entwickeln, die auf alle denkbaren Fälle vorbereitet sind.“ Seine vom Bundesforschungsministerium geförderte Nachwuchsgruppe mit dem Titel „Adversarial Resilience Learning“ setzt daher auf ein lernfähiges System mit Unterstützung von Künstlicher Intelligenz. Diese Software soll vertrauenswürdig sein und nachvollziehbare Entscheidungen treffen.

Um zu erklären, wie das funktioniert, kommt Veith zunächst auf das Brettspiel Go zu sprechen – und einen anderen denkwürdigen Tag, zumindest für die Informatik. Am 13. März 2016 erreichte die KI-Software AlphaGo etwas, das bis dahin als unmöglich galt: Sie sicherte sich den Sieg in einem Match gegen den damaligen Weltmeister Lee Sedol. „Go ist das komplexeste Brettspiel, das die Menschheit jemals erfunden hat, und die schiere Mannigfaltigkeit an Möglichkeiten hat es für eine Software lange extrem schwierig gemacht, eine gute Strategie zu finden“, berichtet Veith. Das Kunststück, die besten menschlichen Spieler zu schlagen, gelang dem Programmiererteam von AlphaGo unter anderem durch einen Trick: Sie fütterten die KI-Software erst mit Spielzügen von Go-Meistern und ließen sie dann unzählige Male gegen eine Kopie von sich selbst spielen. So lernte das Programm die Schwachstellen menschlicher Gegner kennen und entwickelte sogar bis dahin unbekannte Strategien.

Ein „Evil Twin“ fordert den Betreiber- agenten heraus

Zwei identische Softwareprogramme gegeneinander antreten zu lassen, damit immer ausgeklügeltere Taktiken entstehen – das macht Veith auch in seinem Forschungsprojekt. Beim Lernverfahren setzen er und sein fünfköpfiges Team auf die gleiche Methode wie AlphaGo, das sogenannte „Autocurricular Deep Reinforcement Learning“. Dabei erhält ein Computerprogramm – die Forschenden sprechen von einem Agen-

ten – über Sensoren Informationen über ein System. „Das kann ein Spielfeld, aber auch ein Energienetz sein“, erklärt Veith. Zum Agenten gehört außerdem ein Trainingsalgorithmus. Dieser basiert auf einem sogenannten neuronalen Netz, also einem KI-Programm, das biologischen Nervenarchitekturen nachempfunden ist. Dieser Aufbau macht es möglich, dass das neuronale Netz Entscheidungen auf ähnliche Weise trifft wie das menschliche Gehirn – es arbeitet kein vorprogrammiertes Schema ab, sondern programmiert sich gewissermaßen selbst, lernt aus Erfahrungen und ist in der Lage, mit neuen Situationen umzugehen.

Veiths Agent erhält die Vorgabe, einen bestimmten Zustand im System herzustellen. Im Falle des Stromnetzes könnte die Vorgabe darin bestehen, Netzfrequenz und Spannung innerhalb gewisser Grenzen zu halten. „Die Strategie ist aber vorher nicht festgelegt“, betont der Informatiker. Der Agent nimmt etwas wahr, reagiert und berechnet dann ein Feedback-Signal, um herauszufinden, ob die Vorgabe erreicht wurde. Um das Ziel erreichen zu können – etwa, die Netzfrequenz zu stabilisieren –, hat der Agent viele Möglichkeiten: Er kann zum Beispiel Kraftwerke zuschalten, Verbraucher vom Netz trennen und Reglereinstellungen ändern. Außerdem kann er Blindleistungen, die notwendig sind für die Spannungsstabilisierung, verschieben oder Schutzeinrichtungen aktivieren. Mithilfe des Trainingsalgorithmus lernt er, wie er sein Ziel am besten erreicht. „Solche Agentensysteme handeln proaktiv statt nur reaktiv. Und: Wir müssen gar nicht festlegen, wie dieses System etwas erreicht, sondern nur, was der gewünschte Zustand ist“, erklärt Veith.

Um ihren Agenten zu trainieren, mussten die Forschenden zunächst eine passende Simulationsumgebung entwickeln. Diese realitätsnahe Nachbildung eines Energienetzes zählte zu den aufwändigsten Teilen des Projekts. „Unsere an AlphaGo angelehnte Idee war es dann, nicht nur einen Agenten herzunehmen, der das Netz stabilisiert, sondern einen weiteren, quasi einen Evil

Twin, der genau das Gegenteil bewirkt. Die beharren sich dann gegenseitig“, erklärt Veith. Die Idee dahinter: Auf diese Weise wird der „gute“ Agent mit immer neuen Problemen konfrontiert und lernt schneller dazu. Je nachdem, wie die Forschenden den „bösen“ Agenten konfigurierten, konnte dieser Cyberangriffe, extreme Wetterlagen oder auch einen sprunghaft ansteigenden Strombedarf durch zahlreiche gleichzeitig angeschaltete Smart-Home-Geräte simulieren, um sich gegen seinen Gegenspieler durchzusetzen und das Netz aus dem Gleichgewicht zu bringen.

„Mit dieser Grundidee haben wir am Informatikinstitut OFFIS in der Arbeitsgruppe von Sebastian Lehnhoff bereits 2018 angefangen und sind recht weit gekommen“, erzählt Veith. Doch wie viele andere KI-Programme hatte auch das ursprüngliche Agentensystem des Teams ein grundlegendes Problem: Es lieferte keine Informationen darüber, wie es zu seinen Ergebnissen kommt. Für den Einsatz in kritischer Infrastruktur ist eine Nachvollziehbarkeit jedoch zwingend erforderlich, insbesondere, da KI-Programme im

Training manchmal auch Dinge lernen, die wenig sinnvoll sind. „Daher mussten wir das ursprüngliche Konzept fundamental erweitern, verbunden mit der Frage, warum ein Agent etwas tut“, berichtet Veith.

„Unser ‚böser‘ Agent hat schnell gelernt, das Energienetz anzugreifen“

Dieses Ziel verfolgte der Forscher seit 2022 – zum einen in seiner an der Universität angesiedelten Nachwuchsgruppe, zum anderen in zwei EU-geförderten Verbundprojekten, angedockt an das An-Institut OFFIS der Universität. Daran beteiligt waren auch Praxispartner wie der österreichische Energieversorger Wiener Netze und die in Stuttgart beheimatete Netze BW. Um Nachvollziehbarkeit herzustellen, entwickelte das Team einen Algorithmus, der die Strategie des Agenten in einen sogenannten „äquivalenten Entscheidungsbaum“ umwandelt. So gelang es den Forschenden, die Regeln, denen das Programm bei seinen Entscheidungen folgt, Schritt für Schritt abzubilden. „Man sieht dort explizit, welche Schwellwerte bei welchen Sensoren zu einer bestimmten Entscheidung führen und kann nachvollziehen, ob das auch physikalisch sinnvoll ist oder ob der Agent beim Lernen einer statistischen Anomalie auf den Leim gegangen ist“, erläutert Veith.

Im nächsten Schritt optimierte das Team das Training des Agenten. „Ihn von Null auf zu trainieren ist sehr ressourcenintensiv, man braucht viele Millionen Simulationsschritte, bis er eine sinnvolle, übertragbare Strategie gelernt hat“, erklärt der Forscher. Um diesen Prozess zu beschleunigen, entwickelte das Team eine Methodik, um das Anwendungswissen der Praxispartner für das Agententraining nutzbar zu machen. Ein solcher Anwendungsfall war eine Straße in Wien, in der viele Bewohner ein E-Auto fahren. „Wenn die alle in etwa gleichzeitig nach der

Arbeit nach Hause kommen und ihr Auto laden wollen, ist das schlecht für das Netz“, berichtet Veith. Dieses ungünstige Muster machte das Team für den „bösen“ Agenten nutzbar. „Der hat schnell gelernt, wie er damit einen Angriff auf das Energienetz durchführen kann“, sagt Veith. Der Betreiberagent wiederum war gezwungen, eine Gegenstrategie zu entwickeln, auf die wiederum der andere Agent reagierte – ein Spiel, das die Forschenden eine Weile weitertrieben. „Am Ende haben wir einen Entscheidungsbaum gehabt und konnten erkennen: Die Strategie ergibt tatsächlich Sinn.“

Mithilfe der Entscheidungsbäume gelang es dem Team auch, ein weiteres typisches Problem von KI-Software zu lösen – das sogenannte „katastrophale Vergessen“: Wenn eine KI mit neuen Daten trainiert, werden die vorher gelernten Muster manchmal gewissermaßen überschrieben. Als Gegenmaßnahme programmierten Veith und seine Promovierenden ihren Agenten so, dass er Entscheidungen seines eigenen Trainingsalgorithmus mit den früher gefundenen Regeln aus den Entscheidungsbäumen abgleichen kann. Falls sich dabei Diskrepanzen ergeben, kann er dann ein erneutes Training mit den alten Daten einleiten. „Das nennt man Rehearsal, das neuronale Netz muss dann quasi noch einmal die alten Fälle durchspielen, weil es sie offensichtlich vergessen hat“, sagt Veith. Durch diese Kombination lerne der Agent wesentlich schneller und komme zu zuverlässigen Ergebnissen.

Bis das Programm tatsächlich in einer Leitzentrale zum Einsatz kommen kann, werde es wohl noch eine Weile dauern, so der Forscher. „Mit einem echten Energienetz kann man zum jetzigen Zeitpunkt schlecht einen Feldtest machen, daher sprechen wir gerade über andere möglichen Anwendungsfälle für unser System, etwa im Bereich der Krisenvorsorge“, berichtet er. Erklärtes Ziel sei dennoch ein System, das tatsächlich in der Praxis zum Einsatz kommen kann – eine vertrauenswürdige Software, die in kritischen Situationen die richtige Entscheidung trifft.



Eric Veith hat mit seinem Team eine KI zum Steuern von Stromnetzen entwickelt, die auch mit unerwarteten Problemen fertig wird.