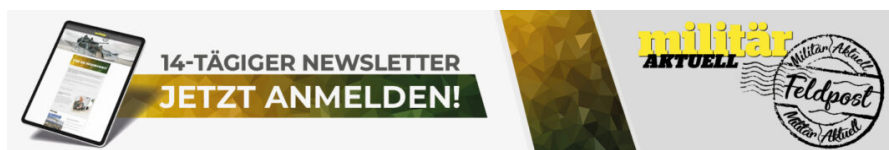


*Interview***KI GEHT IN DEN KRIEG: HEIKO BORCHERT IM GESPRÄCH ÜBER MILITÄRISCHE KI**Von **Christian Bendl** · 20. Juli 2026

Militär Aktuell sprach ausführlich mit Heiko Borchert, Ko-Leiter des Defense AI Observatory an der Helmut-Schmidt-Universität Hamburg, über KI bei Streitkräften – auch beim **österreichischen Bundesheer** –, das Defense Metaverse und Drohnenschwärme.

Herr Borchert, bei der Veranstaltung „Innovision 2026“ vom Militärischen Cyberzentrum des Österreichischen Bundesheeres (-> Innovision 2026: Bundesheer setzt auf Cyber-Innovation), stellten Sie ein Drei-Wellen-Modell für militärische KI vor. Können Sie dieses bitte kurz erklären?

Der Ausgangspunkt ist, dass KI oft als einheitliche Wundertechnologie diskutiert wird, obwohl wir es mit einem ganzen Spektrum an Methoden zu tun haben. Das Drei-Wellen-Modell geht auf die DARPA (Defense Advanced Research Projects Agency/Behörde für Forschungsprojekte der Verteidigung) zurück. Die erste Welle setzt auf handkodiertes Expertenwissen – klassische Regelsysteme. Die zweite Welle, die gegenwärtig das Denken der meisten Streitkräfte bestimmt, ist korrelations- und datengetrieben: Mustererkennung, Klassifikation, Vorhersage auf Basis großer Datensätze, wie wir sie von Bild- und Spracherkennung und von generativen Modellen kennen. Die dritte Welle versucht etwas anderes: kontextbewusste Systeme, die Entscheidungsfolgen verstehen, mit Unsicherheit umgehen und aus sich verändernden Einsatzumgebungen lernen, statt nur Vergangenheitsdaten zu extrapolieren.

**Warum ist diese Unterscheidung der Wellen für Streitkräfte so wichtig?**

Für Streitkräfte ist die Unterscheidung deshalb zentral, weil sie mit der Einsatzfrage verknüpft ist: Genauso wie ich mir Klarheit über den strategischen Zweck verschaffen muss, den ich mit KI erreichen will, brauche ich ein Verständnis dafür, von welcher Welle beziehungsweise welcher KI-Methode ich im Hinblick auf dieses Ziel eigentlich rede.

Sie betonten im Vortrag, dass die meisten Staaten derzeit in dieser datenorientierten zweiten Welle feststecken. Wo sehen Sie die größten Risiken dieser Datenzentrierung – gerade mit Blick auf Krieg, Unsicherheit und begrenzte Ressourcen?

Die zweite Welle ist Fluch und Segen zugleich. Einerseits ist sie das, was jeder sofort nachvollziehen kann. Andererseits entsteht gerade durch die Popularität generativer KI-Modelle das Problem, dass jeder Chancen und Risiken der KI an einer einzigen Methode festmacht und dann sagt: „KI kann dies oder jenes nicht.“ Da muss man zurückfragen: Von welcher KI-Methode reden wir denn genau? In solchen Zuspitzungen liegt die Gefahr, Chancen zu verkennen. Daher ist eben diese Unterscheidung so wichtig.



„KI GEHT IN DEN KRIEG – UND GEWINNT NOCH NICHT.“

Ihre Analysen von Streitkräften, die KI verwenden, zeigen Dynamiken auf, warum welche KI gewählt wird, wie man sich am Gegner und nicht am eigenen Bedarf orientiert, und so weiter.

Ja. Zu beobachten ist beispielsweise ein „Fear of Missing Out“-Momentum: Regierungen glauben, ihre Streitkräfte seien als Koalitionspartner nicht mehr relevant, wenn sie KI nicht nutzen. Das führt oft dazu, dass Rollenmodelle imitiert werden, die jemand anderes für seinen eigenen Kontext entwickelt hat und mit dem er eine strategische Botschaft sendet – etwa: „Ich will jetzt diese AI-first-warfighting-force.“ Ein anderer Akteur läuft dem hinterher, obwohl es überhaupt nicht zu seinen Ambitionen und Missionsprofilen passt.

Birgt dieses Vorgehen nicht auch gewisse Gefahren? Schafft nicht der Klügere einen Vorsprung, im Gegensatz zum Schnellsten?

Genau. Ich halte es für strategisch falsch gedacht, immer nur hinterherlaufen zu wollen. So ist zum Beispiel die Schwierigkeit unserer bisherigen Simulationsansätze, dass wir uns den Gegner nach unseren Wünschen zusammenstellen.



*Heiko Borchert, Ko-Leiter
des Defense AI
Observatory an der
Helmut-Schmidt-
Universität Hamburg.*

So funktioniert der Gegner aber nicht – sonst wäre es kein Gegner, sondern ein Freund. Wir müssen deshalb in Richtung unkonventioneller roter Taktiken gehen, die uns permanent herausfordern. Wenn ich davon ausgehe, dass wir uns künftig verstärkt Peer-to-Peer-Konflikten gegenübersehen, aber nicht bereit bin, mich in der Vorbereitung einem unkonventionell agierenden Gegner zu stellen, dann passt das nicht zusammen.

Seit der Studie „The Very Long Game“ von 2024 ist einiges passiert – neue Konflikte, generative und agentische KI, sehr viel mehr Geld im Markt ...

Richtig. Ich habe für einen im Herbst erscheinenden Aufsatz die Entwicklungen in inzwischen 37 Ländern verglichen. Gegenwärtig gibt es einen „Dreifach-Beschleuniger“: Krieg, konkrete Gefechtsfeldeinsätze sowie der massive Anstieg der Verteidigungsausgaben und privaten Defense-Investitionen. Die weltweiten Militärausgaben lagen 2025 bei knapp 2,9 Billionen US-Dollar – ein Plus von 41 Prozent seit 2016 –, und global floss rekordverdächtige 49,1 Milliarden US-Dollar Risikokapital in den Verteidigungssektor.

„DIE EIGENTLICHE FRAGE IST EINE STRATEGISCHE: FÜR WELCHE MISSIONEN UND MIT WELCHEM STRATEGISCHEN ZWECK SETZE ICH KI EIN?“

Wie ist KI in diesen Entwicklungen zu verorten?

KI geht in den Krieg – und gewinnt noch nicht. Denn dieser Schub ist kein Durchbruch. KI wird nach wie vor primär als Verlängerung bestehender Systeme und Plattformen verstanden und verstetigt damit die Vorstellungen, Konzepte und Praktiken, die zu diesen Systemen gehören.

Was hat Ihre Sicht auf KI im militärischen Kontext seitdem am stärksten verändert? Was sind die drängendsten Entwicklungen?

Drei Entwicklungslinien haben sich seit „The Very Long Game“ verstärkt: Erstens ist die Emulation noch



dominanter, nur mit anderem Twist – heute wollen Staaten die vermeintlichen Lektionen der laufenden Kriege umsetzen. Zweitens entwickeln die meisten Streitkräfte KI weiter so, dass sie bestehende Systeme verbessert. Drittens sind die Sorgen um digitale Souveränität deutlich ausgeprägter als noch vor zwei Jahren – in Europa ebenso wie am Arabischen Golf, in **Brasilien, Indonesien** oder **Pakistan**. Spannend finde ich rückblickend zwei Dinge: Der **Fünf-Tage-Krieg zwischen Indien und Pakistan** ist meines Erachtens der erste, in dem KI-Lösungen prominent auftreten, die nicht westlichen Ursprungs sind. Und wenn die Berichte über das russische unbemannt fliegende V2U-System zutreffen, ist **Russland** derzeit offenbar der erste Akteur, der ein System mit emergentem Verhalten einsetzt, also Verhalten, das nicht vorprogrammierten Routinen entspricht, sondern ohne menschliche Eingriffe auf Veränderungen in Gefecht reagiert (-> **Ukraine-Krieg treibt militärische Revolution voran**). Wer sich davon im Einsatz überraschen lässt, hat schlicht geschlafen – wir haben genau dieses emergente Verhalten vor zwei Jahren in unseren Simulation analysiert und nachgewiesen, wie effektiv es ist.

Simulationen geben heute rasch aussagekräftige Ergebnisse. In diesem Zusammenhang sprechen sie auch vom „Defense Metaverse“.

Das Defense Metaverse ist im Kern nur das Trainingsumfeld: ein virtueller Zwilling des Gefechtsfeldes. Das eigentliche Ziel ist die Entwicklung von KI-Taktiken für Rot und Blau, die spieltheoretisch non-exploitable sind. Non-Exploitability ist ein spieltheoretisches Theorem: Wenn ich diesen Zustand erreiche, kann ohne signifikante strukturelle Veränderung – etwa die Einführung eines neuen Waffensystems – weder Rot noch Blau die Taktik der Gegenseite konterkarieren.

Welche Rolle spielt das Metaversum für die KI-gestützte Fähigkeitsentwicklung?

Man modelliert die eigenen Vorstellungen als blaue Kräfte und den Gegner als rote Kräfte und lässt sie gegeneinander interagieren, um zu prüfen: Stimmen die blauen Prämissen? Dann geht man einen Schritt weiter und lässt Rot selbst lernen, wie es dieses blaue Profil am besten schlägt. Das ist die volle Kraft einer KI der dritten Welle in simulierten Gefechtsszenarien mit unkonventionellem gegnerischen Verhalten – da legen auch Fähigkeitsentwickler die Ohren an, weil sie zum Beispiel Designschwächen eines Systems entdecken können, bevor dieses überhaupt in Entwicklung geht.

„TESTEN SIE IHRE ANNAHMEN IN DIESEM DIGITALEN UMFELD. SIE WERDEN HERAUSFINDEN, WAS ERFOLGREICH IST, WAS NICHT – UND VOR ALLEM, WARUM.“

Wie lassen sich diese technischen Fortschritte mit der Drohnentechnologie vereinen?

Im Rahmen des Projekts Ghostplay, ein Fähigkeits- und Technologieentwicklungsprojekt des Zentrums für Digitalisierungs- und Technologieforschung der **Bundeswehr**, haben wir gemeinsam mit dem Amt für Heeresentwicklung des Deutschen Heeres untersucht, wie KI eingesetzt werden kann, um UAV-Schwärme gegen bodengebundene Flugabwehr des Gegners einzusetzen. In der Simulation kam unter anderem heraus, dass der angreifende Schwarm nicht hoch und schnell fliegen sollte, sondern tief und langsam, um sich zu verstecken und das Gelände zum eigenen Vorteil zu nutzen. In einer besonderen Untersuchung traf ein Switchblade-Schwarm auf das „Clam Shell“-Tieffliegerradar. Dieses sitzt nicht am Boden, sondern in 30 bis 40 Metern Höhe. Tief und langsam fliegend, baut ein Angriff von unten zu wenig kinetische Energie auf, um das hochsitzende Radar zu zerstören. Fliegt man dagegen hoch und schnell von oben an, läuft man sofort in den Standard-Bedrohungsalgorithmus der Flugabwehr. Die entwickelte KI-Taktik lernte in diesem Fall, selbst einen ebenfalls im Einsatzgebiet anwesenden Spike NLOS-Lenkflugkörper anzufordern, um das Ziel auszuschalten und die Mission erfolgreich durchzuführen.

Kann man Drohnen-Schwärme als die „ultimate KI-Waffe“ bezeichnen?

Hinter dem Begriff steckt viel Etikettenschwindel: „Swarming“ wird verkauft als „ein Operateur steuert zehn Systeme“. Das ist aber kein Swarming, das ist – verzeihen Sie den Vergleich – „betreutes Swarming“, weil es noch immer vom Menschen gesteuert ist.

[Mehr Informationen](#)**Inhalt entsperren****Erforderlichen Service akzeptieren und Inhalte entsperren****Damit sind wir bei der Frage, wie weit es den Menschen überhaupt noch braucht.**

Den braucht es auf jeden Fall, und den wird es auch noch brauchen – aber ich muss die Diskussion intelligenter führen. Ein System, bei dem ein Kommandeur auf Führungsebene – bildlich gesprochen – einen Tennisarm bekommt, weil er maschinelle Zielvorschläge nur noch per „Ja“- oder „Nein“-Klick bestätigt, kann nicht das Wesen von Führung sein.

Wie geht man das Thema also besser an?

Die eigentliche Frage ist eine strategische: Für welche Missionen und mit welchem strategischen Zweck setze ich KI ein? Wenn ich mit der Losung „it's all about damage, not accuracy“ Krieg führe, ich also Schaden höher gewichte als Genauigkeit, dann sende ich damit eine klare Botschaft. Die Zerstörung, die daraus entsteht, verursacht nicht die KI, sondern der Mensch, der diese Entscheidung trifft. Der Mensch bleibt dort unverzichtbar, wo Zweck, Grenzen und Verantwortung für Gewaltanwendung definiert werden – aber diese Definitionen können sehr unterschiedlich ausfallen.

Schwarm-Angriffe: Einblicke in die Simulationen des Projekts Ghostplay.

2 von 2 < >

**Auch die finanzielle Frage ist nicht zu unterschätzen.**

Genau. Wir haben einen doppelten Sachzwang: Der eine kommt vom Gefechtsfeld und erfordert es, das Fähigkeitsprofil anzupassen. Das bedingt Investitionen. Der zweite ist innen- und finanzpolitisch: Steigende Verteidigungsaufgaben konkurrenzieren andere politische Aufgaben. Das weckt Widerstände. Wo ist also der Nachweis, dass zehn Milliarden Euro für ein bestimmtes Fähigkeitsprofil im künftigen Konflikt genau den Unterschied machen – und nicht ein anderes Profil?

Was würden Sie also Entscheidungsträgern empfehlen?

Testen Sie Ihre Annahmen in diesem digitalen Umfeld. Sie werden herausfinden, was erfolgreich ist, was nicht – und vor allem, warum. Genau darin liegt jene sinnvolle Beschleunigung durch KI, die wir zu wenig im Fokus haben. Wenn es heißt, KI beschleunige den Krieg, halte ich dagegen: KI beschleunigt im Moment die Automatisierung gewisser Stabsabläufe. Wirklich hilfreich ist die Beschleunigung aber in der Konzeptentwicklung. Denn: Die Lebenswegkosten eines militärischen Systems werden auf der Konzeptseite bestimmt, nämlich durch Fehlannahmen, die langfristig enorme Probleme verursachen. Daher wäre mein eigentlicher Wunsch an Planer, Beschaffer und Parlamente, die digitale Einsatzprüfung als militärischen Beschaffungs-TÜV einzuführen: keine Beschaffung mehr, die nicht vorab auf ihre Leistungsfähigkeit im Defense Metaversum getestet wurde.



Wie könnte ein konkreter Fahrplan aussehen, um KI künftig effizient nutzen zu können?

Erstens die richtige Frage stellen: nicht, welche Software ich morgen einsetze, sondern wie ich morgen kämpfen will. Das ist eine operative und strategische Frage, keine IT-Frage. Zweitens die konsequente Hinwendung zu kontextsensitiven KI-Ansätzen in der Fähigkeitsentwicklung – also Defense-Metaverse-Umgebungen, in denen blaue und rote Kräfte, gerne auch in multinationaler Konstellation, in Millionen Simulationen gegeneinander antreten. Für Staaten mit begrenzten Ressourcen ist das essenziell, weil sie mit geringen Kosten durchdenken können, welche Fähigkeitsprofile wirklich tragfähig sind. Drittens KI-Kompetenz nicht allein bei Technikern zu belassen, sondern in Führungsausbildung, Doktrin und Fähigkeitsplanung zu integrieren – inklusive einer ehrlichen Debatte über digitale Souveränität.

Und wenn Sie die Frage speziell für Österreich beantworten müssten? Wo könnte man an Ihrer Meinung nach ansetzen?

Die entscheidende, mit Abstand schwierigste Frage lautet meiner Meinung nach: Was soll das Bundesheer im Jahr 2035 mit KI besser können als mit jedem anderen lebenswegsteigernden Paket? Diese Frage muss beantwortet werden, bevor man etwas angeht – sonst besteht die Gefahr, dass man etwas tut, nur um den Schein von Modernität zu erwecken.

