

DEN EU AI ACT WIRKSAM UMSETZEN – KI VERANTWORTUNGSVOLL STEUERN – SICHER, TRANSPARENT, NACHVOLLZIEHBAR

ServiceNow als zentrale Plattform zur Umsetzung des EU AI Act

AUTOREN:
KENTARO ELLERT, ANDREJ GREINDL, SEBASTIAN MAYER

EINLEITUNG

Der EU Artificial Intelligence Act (kurz EU AI Act) ist der weltweit erste umfassende Rechtsrahmen für Künstliche Intelligenz. Die finale Fassung wurde am 13. Juni 2024 verabschiedet und am 12. Juli 2024 im EU-Amtsblatt veröffentlicht.¹ Seit dem 1. August 2024 ist die Verordnung in Kraft, die meisten Vorgaben gelten verbindlich ab 2026.² Der

EU AI Act definiert strenge Anforderungen an KI-Systeme, von Governance und Risikomanagement über Transparenz- und Dokumentationspflichten bis hin zu laufender Überwachung. Verstöße können mit empfindlichen Strafen geahndet werden, u. a. mit Bußgeldern bis zu 7% des weltweiten Jahresumsatzes.³ Unterschiedlichste Verantwortliche, bspw. Chief Information Officer, Chief Information Security Officer, Heads of AI, Chief Compliance Officer, Chief Risk Officer und Chief Audit Executives, stehen nun vor der Herausforderung, diese Vorgaben praktisch umzusetzen und kontinuierlich zu überprüfen. Entsprechend ist es für Unternehmen wichtig, sich frühzeitig mit der Umsetzung des EU AI Act auseinanderzusetzen und diesen in ihre Governance- und Compliance-Strukturen zu integrieren.

Dieses Whitepaper fasst die zentralen Anforderungen des EU AI Act zusammen und zeigt praxisnah, wie Unternehmen diese mit Unterstützung der Plattform ServiceNow umsetzen können. Im Mittelpunkt stehen vier Use Cases, vom Aufbau eines KI-Inventars bis zum automatisierten Monitoring, die veranschaulichen, wie ServiceNow die Einhaltung des EU AI Act effizient unterstützt.

» Die erfolgreiche Umsetzung des EU AI Act ist mehr als nur die Erfüllung von regulatorischen Pflichten, sie ist eine strategische Chance, KI nachhaltig und sicher in Unternehmen zu verankern. «

SEBASTIAN MAYER, MANAGING DIRECTOR,
PROTIVITI DEUTSCHLAND



1. ZENTRALE ANFORDERUNGEN DES EU AI ACT

Der EU AI Act verfolgt einen risikobasierten Ansatz und stellt differenzierte Pflichten auf, abhängig vom Gefährdungspotenzial eines KI-Systems. Insbesondere in den Bereichen Governance, Risikomanagement, Transparenz, Risikoklassifizierung und Überwachung definiert der EU AI Act umfangreiche Vorgaben. Unternehmen müssen diese verstehen und in ihre Governance-Strukturen übertragen, um regelkonform zu handeln. Im Folgenden werden die wichtigsten Anforderungen zusammengefasst.

1 artificialintelligenceact.eu
2 digital-strategy.ec.europa.eu
3 eur-lex.europa.eu

Risikoklassifizierung von KI-Systemen



Abbildung 1: Die vier Risikostufen des EU AI Act

Der EU AI Act teilt KI-Anwendungen in vier Risikostufen ein:

- **Verbotene Praktiken im Bereich der Künstlichen Intelligenz:** Die Verordnung benennt verbotene KI-Anwendungen, die als klare Bedrohung für Sicherheit oder Grundrechte gelten. Beispiele sind etwa manipulative KI-Systeme, die Menschen unbewusst steuern sollen, die Ausnutzung von Schwächen vulnerabler Gruppen (z.B. von Kindern), jegliche Form von staatlichem Social Scoring sowie bestimmte prognostische KI zur polizeilichen Risikoeinschätzung. Ebenfalls untersagt sind etwa die wahllose Extraktion, Kopie, Speicherung und Wiederverwendung biometrischer Daten aus dem Internet zur Gesichtserkennung, die biometrische Identifizierung von Personen in Echtzeit im öffentlichen Raum für Strafverfolgungszwecke (mit eng begrenzten Ausnahmen), diskriminierende Emotionserkennung am Arbeitsplatz oder KI zur Bestimmung geschützter Merkmale (wie z.B. Religion) durch biometrische Analyse. Unternehmen müssen sicherstellen, dass derartige KI-Praktiken und -Funktionalitäten in keinem von ihnen betriebenen oder genutzten System vorkommen.
- **Hochrisiko-KI-Systeme:** Zur Kategorie der Hochrisiko-Systeme zählen KI-Systeme, die potenziell erheblichen Einfluss auf Leben, Gesundheit oder grundlegende Rechte von Menschen haben. Der Gesetzgeber nennt hier verschiedene Einsatzbereiche, u.a. sicherheitskritische Infra-

struktur (z. B. KI-Steuerung im Verkehrswesen), Bildungs- und Personalentscheidungen (z. B. Software zur Bewerberauswahl), Zugang zu essenziellen Dienstleistungen (etwa KI-gestützte Kreditwürdigkeitsprüfung), polizeiliche und justizielle Anwendungen (z. B. Beweismittelbewertung) oder KI in der Medizin. Solche hochriskanten KI-Systeme sind nicht per se verboten, unterliegen aber strikten Auflagen, bevor sie in Verkehr gebracht oder in Betrieb genommen werden dürfen (vgl. „Anforderungen an Hochrisiko KI-Systeme“).

- **Begrenztes Risiko & Transverpflichtungen:** Hierunter fallen KI-Anwendungen, die zwar kein hohes Risiko darstellen, aber dennoch gewisse Transparenzpflichten nach sich ziehen. Insbesondere muss bei solchen Systemen der Einsatz von KI gegenüber Nutzern offengelegt werden. Ein typisches Beispiel sind Chatbots oder virtuelle Assistenten: Nutzer müssen informiert werden, dass sie mit einer Maschine interagieren. Auch für generative KI (etwa Texte oder Bilder erzeugende Modelle) schreibt der EU AI Act vor, dass KI-generierte Inhalte als solche erkennbar sein müssen, z. B. durch Kennzeichnung von synthetischen Medien oder dem Hinweis „AI-generiert“. Auch sogenannte Deepfakes – also KI-generierte Inhalte, die realistische menschliche Darstellung vortäuschen, müssen klar und sichtbar als künstlich gekennzeichnet werden. Diese Transparenzanforderungen sollen verhindern, dass Menschen

unwissentlich KI-Ergebnissen Glauben schenken (vgl. „Transparenzpflichten für Anbieter und Betreiber bestimmter KI-Systeme“).

- **Minimales oder kein Risiko:** Eine Vielzahl alltäglicher KI-Anwendungen (z. B. Spam-Filter, KI in Videospielen) fällt in die Kategorie des minimalen Risikos bzw. wird als gänzlich nicht risikobehaftet klassifiziert. Für solche Systeme sieht der EU AI Act keine zusätzlichen Vorgaben vor.

Anforderungen an Hochrisiko KI-Systeme

Für Hochrisiko-KI-Systeme schreibt der EU AI Act ein robustes internes Kontrollsystem vor. Anbieter (Entwickler oder Inverkehrbringer) solcher KI-Systeme **müssen ein systematisches Risikomanagement etablieren**, das den gesamten Lebenszyklus der KI abdeckt. Risiken sind fortlaufend zu identifizieren, zu analysieren und zu mindern, von der Planung über Entwicklung und Test bis zum laufenden Betrieb eines KI-Systems. Konkret stellt der EU AI Act in diesem Zusammenhang verschiedene Anforderungen an Hochrisiko-KI-Systeme (Kapitel 2 EU AI Act) und verlangt die Einhaltung unterschiedlicher Pflichten durch die Anbieter und Nutzer von Hochrisiko-KI-Systemen:

- **Risikomanagementsystem (Art. 9 EU AI Act):** Anbieter von Hochrisiko-KI-Systemen müssen ein kontinuierliches Risikomanagementverfahren etablieren, welches die systematische Identifizierung und Analyse sowie die Minderung potenzieller Risiken aus dem Einsatz des KI-Systems für die Gesundheit, die Sicherheit und die in der Charta der Europäischen Union verankerten Grundrechte, einschließlich Demokratie, Rechtsstaatlichkeit und Umweltschutzes, sicherstellt. Das Risikomanagement erstreckt sich über den gesamten Lebenszyklus des KI-Systems, von der Planung über Tests bis zum laufenden Betrieb, und ist regelmäßig zu aktualisieren, da sich Risiken im laufenden Betrieb der KI-Anwendung ändern können.
- **Daten und Daten-Governance (Art. 10 EU AI Act):** Für die Nutzung von Daten für Trainings-, Validierungs- und Testzwecken sind umfangreiche Qualitätskriterien zu berücksichtigen. Es müssen bspw. Verfahren etabliert werden, sodass verwendete Datensätze möglichst relevant, repräsentativ, fehlerfrei und vollständig sind, um die Wahrscheinlichkeit

für Verzerrungen oder Fehler bei der Ausführung eines KI-Systems zu minimieren. Gegebenenfalls sind Datensatzprüfungen und -vorbehandlungen (z. B. Bereinigung oder Bias-Checks) durchzuführen, um den Vorgaben des EU AI Act zu entsprechen. Der EU AI Act fordert zudem explizit Maßnahmen zur Vermeidung diskriminierender Ergebnisse durch mangelhafte Daten. Diese Maßnahmen sollen gewährleisten, dass KI-Systeme zuverlässig und gerecht arbeiten, ohne bestimmte Gruppen systematisch zu benachteiligen.

- **Technische Dokumentation (Art. 11 & 18 EU AI Act):** Anbieter von Hochrisiko-KI haben ausführliche technische Dokumentationen (siehe ebenfalls Annex IV EU AI Act) zu erstellen. Darin werden u. a. Systemzweck, Funktionsweise, Modell-Architektur, Trainingsdaten, Performance-Metriken und alle getroffenen Maßnahmen dokumentiert. Diese Dokumentation muss es Dritten (z. B. der notifizierende Behörde) ermöglichen, das System nachzuvollziehen und auf die Erfüllung der Vorschriften zu prüfen. Die technische Dokumentation ist für einen Zeitraum von mindestens zehn Jahre ab dem Inverkehrbringen bzw. der Inbetriebnahme des Hochrisiko-KI-Systems aufzubewahren.
- **Aufzeichnungspflichten (Art. 12 & Art. 19 EU AI Act):** Hochrisiko-KI-Systeme müssen so gestaltet sein, dass relevante Ereignisse automatisiert protokolliert werden. Insbesondere Entscheidungen oder Nutzungsvorgänge sollen nachvollziehbar festgehalten werden. Die Protokolle müssen von Anbietern und Betreibern über angemessene Zeiträume, jedoch mindestens sechs Monate, aufbewahrt werden, um z. B. die Ursache von Fehlentscheidungen untersuchen zu können.
- **Transparenz und Bereitstellung von Informationen für die Betreiber (Art. 13 EU AI Act):** Hochrisiko-KI muss so konzipiert und entwickelt werden, dass ihr Betrieb hinreichend transparent ist, damit die Betreiber den Output richtig interpretieren und angemessen verwenden können. Anbieter müssen für jede Hochrisiko-KI verständliche Betriebsanleitungen für den Gebrauch zu Verfügung stellen. Diese enthalten u. a. Angaben zu Zweck und Funktionsumfang des Systems, seiner Leistung (z. B. Genauigkeitsmetriken, Robustheit, Sicherheitsmaßnahmen), bekannten Einschränkungen

und möglichen Risiken sowie Hinweise für den korrekten Einsatz. Hierdurch soll sichergestellt werden, dass Nutzer die Ausgaben eines KI-Systems nachvollziehen und Fehlanwendungen vermeiden können.

- **Menschliche Aufsicht (Art. 14 EU AI Act):** Bereits bei der Entwicklung muss berücksichtigt werden, dass Hochrisiko-KI-Systeme durch geeignete Maßnahmen von Menschen überwacht und kontrolliert werden können („Human Oversight“). KI sollte z. B. Warnanzeigen oder Schnittstellen bieten, damit menschliche Aufsichtspersonen eingreifen können, falls das System fehlerhaft reagiert oder unerwünschte Ergebnisse liefert. Die technischen Vorkehrungen unterstützen den Betreiber dabei, die vom Anbieter vorgesehenen Human-Oversight-Maßnahmen praktisch umzusetzen, etwa die Möglichkeit, automatisierte Empfehlungen zu überstimmen oder das System bei Bedarf abzuschalten.
- **Genauigkeit, Robustheit und Cybersicherheit (Art. 15 EU AI Act):** Hochrisiko-KI-Systeme müssen mit angemessener Genauigkeit arbeiten, konsistente Leistung erbringen und robust gegenüber Störungen, Fehlern oder Missbrauch sein. Dies gilt für alle vorgesehenen Einsatzbedingungen über den gesamten Lebenszyklus. Wo erforderlich, sind technische Redundanzen (z. B. Failsafe-Modi oder Back-up-Systeme) vorzusehen, um sicherheitskritische Ausfälle aufzufangen. Zudem müssen Anbieter angemessene Cybersicherheits-Maßnahmen treffen, damit das KI-System gegen Manipulation oder Angriffe geschützt ist. Insbesondere sollen bekannte Angriffsmuster wie Data Poisoning (gezielte Manipulation von Trainingsdaten), Modellmanipulation (gezielte negative Veränderung des KI-Modells selbst) oder adversariale Inputs (gezielt manipulierte Eingabedaten) proaktiv verhindert, erkannt und abgewehrt werden. Auch selbstlernende KI-Modelle müssen so konzipiert sein, dass negative Feedback-Schleifen (z. B. eine fortschreitende Verzerrung durch eigene fehlerhafte Ausgaben) vermieden oder durch Korrekturmechanismen begrenzt werden.
- **Qualitätsmanagement-System (Art. 17 EU AI Act):** Anbieter von Hochrisiko-KI-Systemen müssen ein wirksames Qualitätsmanagementsystem (QMS) implementieren. Dieses QMS

stellt sicher, dass alle oben genannten Anforderungen und Prozesse systematisch eingehalten und dokumentiert werden. Dazu gehören klare Verfahren für Designkontrollen, Tests, Validierung und Qualitätssicherung des KI-Systems sowie Maßnahmen zum Änderungsmanagement, falls am System nachträglich Anpassungen vorgenommen werden. Das QMS muss außerdem vorsehen, dass die Verantwortlichkeiten im Unternehmen klar geregelt sind und Personal angemessen geschult wird. Die Dokumentation des QMS ist für einen Zeitraum von mindestens zehn Jahre ab dem Inverkehrbringen bzw. der Inbetriebnahme des Hochrisiko-KI-Systems aufzubewahren.

Über die internen Anforderungen hinaus unterliegt die Bereitstellung von Hochrisiko-KI weiterer aufsichtsrechtlicher Kontrolle. So ist vor dem Inverkehrbringen eine Konformitätsbewertung durchzuführen (u. a. Art. 6 Abs. 1 b EU AI Act, ggf. unter Einbindung einer unabhängigen notifizierten Stelle, falls keine harmonisierten Normen oder spezifische Ausnahmen greifen). Außerdem müssen Hochrisiko-KI-Systeme vor dem Einsatz in einer EU-Datenbank registriert werden. Diese Registrierung (Art. 49 EU AI Act) erfasst Kerninformationen zum System (u. a. den Anbieter, Zweck, Einstufung, etc.) und dient der Transparenz gegenüber Behörden und der Öffentlichkeit. Die EU-Kommission betreibt hierzu gemäß Art. 71 EU AI Act eine zentrale Datenbank für Hochrisiko-KI. Durch diese zusätzlichen formalen Pflichten, einschließlich möglicher Beteiligung notifizierter Stellen, CE-Kennzeichnung, Konformitätserklärung und Registrierung, wird sichergestellt, dass Hochrisiko-KI erst nachweislich regelkonform und behördlich nachvollziehbar in den Markt gelangt.

Transparenzpflichten für Anbieter und Betreiber bestimmter KI-Systeme (Art. 50 EU AI Act)

Neben den Vorschriften für Hochrisiko-KI definiert der EU AI Act auch Transparenzpflichten für bestimmte sonstige KI-Anwendungen, die zwar kein hohes Risiko darstellen, aber im Alltag der Nutzer besondere Bedeutung haben. Diese Vorgaben sollen vor allem sicherstellen, dass Menschen erkennen, wann sie es mit KI zu tun haben oder KI-generierte Inhalte konsumieren.

Betroffen sind insbesondere:

- **KI, die mit Menschen interagiert (z.B. Chatbots):** Wenn ein KI-System dazu bestimmt ist, direkt mit natürlichen Personen zu interagieren, muss der Anbieter dafür sorgen, dass die betroffenen Personen darüber informiert sind, dass es sich um eine KI handelt. Praktisch bedeutet dies, dass Chatbots, virtuelle Assistenten oder ähnliche Systeme sich dem Nutzer gegenüber als KI-System zu erkennen geben, es sei denn, dies ist ohnehin offensichtlich durch die Umstände.
- **Generative KI für synthetische Inhalte:** Anbieter von KI-Systemen, die synthetische oder manipulierte Inhalte (Text, Bild, Audio, Video) erzeugen, müssen technische Vorkehrungen treffen, damit die künstlich generierten Inhalte kenntlich sind. In der Praxis sollen solche Inhalte mit dauerhaft erkennbaren Hinweisen oder Wasserzeichen versehen werden (maschinell auslesbar und für Menschen wahrnehmbar), die anzeigen, dass z. B. ein Bild oder Text künstlich erstellt wurde. Keine Kennzeichnung ist nötig, wenn das KI-System lediglich redaktionell unterstützt, etwa Autokorrektur oder Bildfilter, und den menschlichen Input nicht wesentlich verändert.
- **Emotionserkennung und biometrische Kategorisierung:** Setzt ein Akteur ein KI-System zur Erkennung von Emotionen oder zur biometrischen Kategorisierung von Personen ein (bspw. Systeme, die aus Mimik auf Gemütszustände schließen oder Menschen bestimmten demografischen Gruppen zuordnen), so muss der Betreiber sicherstellen, dass die betroffenen Personen darüber informiert werden.
- **Deepfakes:** Wer KI-Systeme einsetzt, um realistisch wirkende künstliche Medieninhalte („Deepfakes“) zu erzeugen oder zu verbreiten, muss klarstellen, dass es sich um künstlich erzeugte bzw. manipulierte Inhalte handelt. Beispielsweise sollte ein KI-generiertes Video einen Hinweis darauf tragen, dass es nicht real ist. Dies gilt insbesondere, wenn derartige Inhalte der öffentlichen Meinungsbildung dienen. Hier muss offengelegt werden, dass der Text oder das Material künstlich erstellt wurde.

Diese Transparenzanforderungen für begrenzt risikante KI-Anwendungen sollen das Vertrauen der Nutzer stärken und Missbrauch vorbeugen, ohne die Innovationsfähigkeit unnötig zu hemmen. Anwender wissen aufgrund dieser Vorgaben, wann sie mit einer KI kommunizieren und können Inhalte besser beurteilen.

KI-Modelle mit allgemeinem Verwendungszweck

Ein weiterer wichtiger Aspekt der EU AI Act sind besondere Vorgaben für KI-Modelle mit einem sogenannten allgemeinen Verwendungszweck („General-Purpose AI“), insbesondere solche, die als Grundlagenmodelle breite Anwendung finden. Hierunter fallen etwa große Sprachmodelle oder multimodale Modelle, die nicht für einen einzigen spezifischen Zweck entwickelt wurden, sondern in verschiedenen Bereichen einsetzbar sind. Die Verordnung trägt der Rechnung, dass solche Modelle von vielen nachgelagerten Anbietern weiterverwendet werden und potenziell breite Auswirkungen haben können.

Die zentralen Pflichten für Anbieter allgemeiner KI-Modelle sind:

- **Technische Dokumentation (Art. 53 Abs. 1 a EU AI Act):** Der Anbieter muss eine umfassende technische Dokumentation des Modells erstellen und aktuell halten. Diese soll u. a. Einblicke in den Trainings- und Testprozess, die Modellarchitektur, Leistungsbewertungen und getroffene Schutzmaßnahmen geben. Ziel ist, dass relevante Behörden bei Bedarf die Dokumentation anfordern und prüfen können, ob das Modell den Vorschriften entspricht.
- **Weitergabe von Informationen an nachgelagerte Nutzer (Art. 53 Abs. 1 b EU AI Act):** Allgemeine KI-Modelle werden oft von Dritten eingebunden, um daraus spezifische KI-Systeme zu entwickeln. Daher muss der Modell-Anbieter ausreichende Informationen und Dokumentation für diese „Downstream“-Provider bereitstellen. Konkret sollen Entwickler, die das Modell weiterverwenden, über Leistungsfähigkeiten, Grenzen, empfohlene Einsatzbereiche und bekannte Risiken des Modells informiert werden. So können diese ihre eigenen Verpflichtungen einhalten und einschätzen, ob das Modell für eine geplante Hochrisiko-Anwendung geeignet ist.

- **Urheberrechts-Compliance (Art. 53 Abs. 1 c EU AI Act):** Es ist sicherzustellen, dass das Modell geltende Urheberrechte respektiert. Insbesondere muss der Anbieter darauf achten, auch durch technische Mittel, dass er keine Daten verwendet, die mit einem „Maschinenlernverbot“ versehen sind. Die Entwicklung allgemeiner KI-Modelle darf also nicht über geltende Copyright-Bestimmungen hinweggehen.
- **Offenlegung des Trainingsmaterials (Art. 53 Abs. 1 d EU AI Act):** Anbieter müssen eine Zusammenfassung der im Training verwendeten Inhalte veröffentlichen. Dieser Bericht (in einem vom Büro für Künstliche Intelligenz vorgegebenen Format) soll in verständlicher Weise aufzeigen, welche Datenquellen für das Modelltraining genutzt wurden.
- **Anforderungen für ausländische Anbieter (Art. 54 EU AI Act):** Anbieter allgemeiner KI-Modelle, die nicht in der EU ansässig sind und ihr Modell in der EU auf den Markt bringen, müssen vorab einen in der EU niedergelassenen Bevollmächtigten benennen. Dieser vertritt den Anbieter gegenüber Aufsichtsbehörden und stellt sicher, dass z.B. die technische Dokumentation verfügbar gehalten wird. Hierdurch soll die Durchsetzung der Regeln bei Anbietern außerhalb Europas erleichtert werden.

Allgemeine KI-Modelle können aufgrund ihrer Reichweite und Leistungsfähigkeit zudem systemische Risiken mit sich bringen, d.h. mögliche großflächige negative Auswirkungen auf wichtige Bereiche (etwa kritische Infrastrukturen, demokratische Prozesse oder die Gesellschaft insgesamt). Der EU AI Act führt dafür die Kategorie „Allgemeine KI-Modelle mit systemischem Risiko“ ein (Art. 51 EU AI Act), worunter ein Modell fällt, wenn es Fähigkeiten mit hohem Wirkungsgrad aufweist (gemessen an Kriterien in Annex XIII) und damit ein entsprechend hohes Missbrauchs- oder Schadenspotenzial besteht (bspw. große und weit verbreitete Sprachmodelle, worunter potentiell GPT-4/5 oder Claude fallen). Zudem wird angenommen, dass ein allgemeines KI-Modelle mit systemischem Risiko vorliegt, wenn die für das Training verwendeten Berechnungen, gemessen in Gleitkommaoperationen (Rechenvorgänge mit Dezimalzahlen, die zeigen, wie viel Rechenleistung ein KI-System verbraucht, etwa beim Berechnen neuronaler Netzwerke), mehr als 10 beträgt. Für

solche Modelle gelten die nachfolgenden zusätzlichen Auflagen (Art. 55 EU AI Act):

- Der Anbieter muss verschärfte Evaluierungen des Modells durchführen, worunter auch Angriffstests zählen, um Schwachstellen oder Missbrauchsmöglichkeiten frühzeitig zu identifizieren. Die Testergebnisse und gegebenenfalls ergriffenen Gegenmaßnahmen sind zu dokumentieren.
- Es sind gezielte Risikominderungsmaßnahmen umzusetzen, die etwaige systemweite Risiken auf EU-Ebene adressieren. Der Anbieter muss daher bspw. beurteilen, welche negativen Effekte sein Modell auf Unionsebene haben könnte (z.B. für die Informationssicherheit oder öffentliche Debatte) und geeignete Vorkehrungen treffen, um diese Risiken zu begrenzen.
- Der Anbieter muss schwerwiegende Vorfälle und ergriffene Abhilfemaßnahmen erfassen und dokumentieren und diese unverzüglich an die Aufsichtsbehörden (das Büro für Künstliche Intelligenz und gegebenenfalls nationale Behörden) melden.
- Es ist ein angemessenes Maß an Cybersicherheit sowohl für das KI-Modell, aber auch die physische Infrastruktur zu gewährleisten. Das umfasst bspw. den Schutz des KI-Modells selbst sowie der dazugehörigen Infrastruktur vor Angriffen, Manipulation oder Spionage.

Um die Einhaltung dieser Pflichten zu erleichtern, fördert die Verordnung die Entwicklung sogenannter Praxisleitfäden für Anbieter allgemeiner KI-Modelle (Art. 56 EU AI Act), die als branchenweiten Standards ein EU AI Act konformes Maß an Transparenz, Sicherheit und Risikomanagement zu etablieren. Anbieter, die sich an solch anerkannte Kodizes halten, können damit ihre Konformität nachweisen.

Überwachung und Berichtspflichten (Monitoring & Reporting)

Auch nach dem Inverkehrbringen und der Inbetriebnahme einer KI-Anwendung endet die Verantwortung der Akteure nicht. Der EU AI Act schreibt ein kontinuierliches Monitoring und proaktive Berichtspflichten vor, um die laufende Sicherheit und Regelkonformität von KI-Systemen sicherzustellen. Wesentliche Aspekte hierbei sind:

- **Beobachtung nach dem Inverkehrbringen durch Anbieter (Art. 72 EU AI Act):** Jeder Anbieter eines Hochrisiko-KI-Systems muss ein System zur **Beobachtung nach dem Inverkehrbringen** einrichten. Konkret sollen einschlägigen Daten zur Leistung des KI-Systems, die von den Anbietern oder den Betreibern bereitgestellt oder aus anderen Quellen erhoben werden können, über ihre gesamte Lebensdauer hinweg aktiv und systematisch erhoben, dokumentiert und analysiert werden. Falls das KI-System mit anderen Systemen interagiert, ist auch diese Wechselwirkung einzubeziehen. Das Monitoring-System muss sich auf einen strukturierten Plan (Teil der technischen Doku) stützen und sollte an das Risikoniveau des Systems angepasst sein (umfangreichere Überwachung bei höherem Risiko).

- **Meldepflicht bei schwerwiegenden Vorfällen (Art 73 EU AI Act):** Trotz aller Vorsichtsmaßnahmen können schwerwiegende Vorfälle oder Fehlfunktionen auftreten. Der EU AI Act verpflichtet Anbieter von Hochrisiko-KI, jeden schwerwiegenden Vorfall unverzüglich der Marktüberwachungsbehörde zu melden. Als „schwerwiegend“ gelten etwa Situationen, in denen das KI-System zu einem tödlichen oder lebensbedrohlichen Ereignis, erheblichen wirtschaftlichen Schäden oder schweren Verletzungen von Grundrechten führt. Die Meldung soll erfolgen, sobald ein kausaler Zusammenhang mit dem KI-System bekannt oder wahrscheinlich ist, spätestens innerhalb von 15 Tagen nachdem der Anbieter (oder Betreiber) davon erfahren hat. In besonders kritischen Fällen (z.B. bei Todesfällen oder eine schwere bzw. unumkehrbare Störung der Verwaltung oder des Betriebs kritischer Strukturen) gelten verkürzte Fristen. Der Anbieter muss im Anschluss an die Meldung eines schwerwiegenden Vorfalls diesen untersuchen, eine Risikobewertung durchführen und Korrekturmaßnahmen ergreifen.

» Erfolgreiches Risikomanagement im Kontext des EU AI Act erfordert eine konsequente Umsetzung und kontinuierliche Überwachung. Ein integriertes Kontrollsystem hilft, Risiken effektiv zu managen und frühzeitig zu erkennen. «

ANDREJ GREINDL, MANAGING DIRECTOR,
PROTIVITI DEUTSCHLAND



Zusammenfassend lässt sich festhalten, dass der EU AI Act ein umfassendes, risikobasiertes Rahmenwerk schafft, um den verantwortungsvollen Einsatz von KI in der EU zu gewährleisten. Organisationen, die KI-Systeme entwickeln oder einsetzen, stehen nun vor der Herausforderung, eine Vielzahl neuer regulatorischer Vorgaben, insbesondere im Hinblick auf Risikomanagement, Datenqualität, Transparenz, menschliche Aufsicht und kontinuierliches Monitoring, strukturiert umzusetzen. Während Hochrisiko-KI-Systeme mit besonders strengen Anforderungen verbunden sind, stellt die Verordnung auch klare Regeln für Anwendungen mit begrenztem Risiko sowie KI-Modelle mit allgemeinem Verwendungszweck auf. Diese umfassenden Vorgaben sollen letztlich sicherstellen, dass KI-Technologien vertrauenswürdig, sicher und gesellschaftlich akzeptiert genutzt werden, gleichzeitig soll die Innovationskraft Europas hierdurch nicht geschwächt werden, sondern das Vertrauen in den Wirtschaftsstandort Europa weiter gefestigt werden.

2. OPERATIVE UMSETZUNG DES EU AI ACT MIT SERVICENOW

Vielen Unternehmen stellt sich jedoch die Frage, wie sich die beschriebenen Anforderungen nun konkret, effizient und nachhaltig umsetzen lassen. Eine bewährte Methode ist die Nutzung einer integrierten Governance, Risk und Compliance-Plattform. Im Folgenden werden mittels vier relevanter Anwendungsfälle des EU AI Act beschrieben, wie Organisationen die Anforderungen des EU AI Act mit den umfangreichen Funktionalitäten der ServiceNow-Plattform umsetzen können. Der Fokus liegt hierbei u.a. auf dem ServiceNow AI Control Tower, der als zentrale Steuerungsinstanz für KI-Systeme dient und in Zusammenspiel mit bestehenden Modulen (z.B. IT Operations Management (ITOM), Integrated Risk Management (IRM) und Third Party Risk Management (TPRM)) die operative Compliance sicherstellt.

USE CASE 1: KI-INVENTAR – ZENTRALES VERZEICHNIS ALLER KI-SYSTEME BASIEREND AUF CSDM 5.0

Unternehmen müssen im Zuge des EU AI Act einen vollständigen Überblick über ihre eingesetzten KI-Systeme und -Modelle bereithalten. Ein KI-Inventar dient als zentrales Register aller KI-Anwendungen, inklusive ihrer Zweckbestimmung, Risikoklassifikation, Datenquellen und Verantwortlichkeiten. Dieses Inventar ist die Grundlage für Transparenz und Zuständigkeiten, da nur bekannte und vollständig erfasste KI-Systeme effektiv überwacht und gesteuert werden können. Zudem erleichtert ein zentrales Verzeichnis die Erfüllung von Dokumentationspflichten und Meldeauflagen aus der KI-Regulierung (z.B. Registrierung hochriskanter KI-Systeme bei Behörden).

Die ServiceNow-Plattform ermöglicht es, durch den AI Control Tower ein solches KI-Inventar bspw. unter Nutzung des neuen Common Service Data Model (CSDM) 5.0 aufzubauen. Konkret werden KI-Assets in der Configuration Management Database (CMDB) hinterlegt und nach CSDM klassifiziert. CSDM 5.0 hat eigens Strukturkomponenten für KI eingeführt, etwa AI System Digital Asset (für

» Bei der Umsetzung des EU AI Act zeigt sich, dass technologische und regulatorische Anforderungen Hand in Hand gehen müssen. Eine zentrale Plattformlösung ist hier entscheidend, um beide Welten effektiv zusammenzubringen.“ «

KENTARO ELLERT, ASSOCIATE DIRECTOR,
PROTIVITI DEUTSCHLAND



KI-Anwendungen) sowie zugehörige Objekte für KI-Modelle und KI-Datensätze. Dadurch lässt sich jedes KI-System als Configuration Item mit standardisierten Attributen und Beziehungen abbilden. Insbesondere können KI-Systeme direkt Geschäftsprozessen oder zugehörigen (IT-)Dienstleistungen zugeordnet werden. Auf diese Weise wird KI in den Kern der Geschäfts- und IT-Services eingebettet, sodass Abhängigkeiten transparent sind. Das Inventar ist nicht nur eine technische Liste, sondern gilt als die Business-Architektur: Alle KI-Komponenten und

deren Beziehungen zu Anwendungen, Services und Datenflüssen sind zentral erfasst. Dies bildet die Basis, um Auswirkungen von KI auf Geschäfts-services nachzuvollziehen und strategische Entscheidungen faktenbasiert zu treffen.

Schlüsselfunktionen des KI-Inventars:

- **ServiceNow AI Control Tower als zentrale Plattform für KI-Governance.** Hier werden alle KI-Systeme des Unternehmens erfasst (AI Asset Inventory) und mit relevanten Informationen (Einsatzzweck, Verantwortliche, Risiko-Einstufung) hinterlegt, **selbst wenn diese nicht in der ServiceNow Plattform entwickelt wurden.** Über das AI Stewardship-Konzept wird jedem KI-Modell ein verantwortlicher KI-Eigentümer zugewiesen, der für die Überwachung und Freigabe zuständig ist. Zudem sorgen Freigabe-Workflows dafür, dass kein KI-gestützter Service live geht, ohne vorherige gründliche Prüfung und Genehmigung.
- **CSDM-gestützte Struktur:** Jedes KI-System ist nach einheitlichen Standards im CMDB modelliert und mit Unternehmensservices verknüpft. Dies gewährleistet konsistente Daten und Kontext zu Geschäftsprozessen.
- **Integration in die Plattform:** Das Inventar ist in die Now-Plattform integriert, andere Module wie IRM oder Policy & Compliance können diese Daten nutzen, um risikoreiche KI zu identifizieren und entsprechende Kontrollmaßnahmen zuzuordnen.
- **Lebenszyklus-Verfolgung:** Von der Planung über die Entwicklung bis zum Betrieb wird der Status jeder KI-Lösung nachvollziehbar dokumentiert. Änderungen (z. B. neue Modelle oder Versionsupdates) sind im Inventar nachverfolgbar, was die Grundlage für Change Management und kontinuierliche Compliance-Prüfungen bildet. Die Möglichkeit zur Integration mit anderen ServiceNow-Modulen, wie bspw. IT Service Management und Security Operations, liefert zusätzlichen Mehrwert und erweitert das KI-Inventar. Dadurch wird die Governance dieser Lösungen weiter optimiert.

USE CASE 2: RISIKOBEWERTUNG UND KONTROLLSTEUERUNG:

Die Arbeit endet nicht mit der Inbetriebnahme einer KI, vielmehr verlangt der EU AI Act eine laufende Überwachung und Qualitätssicherung von KI-Systemen. Jeder KI-Anwendungsfall sollte noch vor der Entwicklung einer standardisierten Risikobewertung unterzogen werden. Zunächst wird das Risikoprofil des KI-Systems ermittelt, etwa durch Einordnung in Risikoklassen (z. B. minimales, begrenztes oder hohes Risiko gemäß EU AI Act) und Identifizierung use-case-spezifischer Risiken. Auf Basis dieser Einstufung wird dann ein Katalog geeigneter Maßnahmen zur Risikominderung abgeleitet. Beispielsweise erfordern Hochrisiko-KI-Systeme umfangreichere Kontrollen wie strengere Validierung, ausführliche Dokumentation und menschliche Überwachung, während bei geringeren Risiken eher Basiskontrollen genügen. Wichtig ist dabei, den unternehmensindividuellen Risikoappetit zu berücksichtigen, sollten identifizierte KI-Risiken nicht tragbar sein, muss ein Mechanismus zur Projektpause oder Anpassung greifen. Das heißt, bevor ein KI-Modell produktiv geht, soll klar nachweisbar sein, welche Risiken (z. B. Bias, Datenschutz, Fehlentscheidungen) bestehen und welche Kontrollen implementiert wurden, um diese Risiken zu mitigieren.

Zur Umsetzung können Unternehmen, neben den bereits beschriebenen KI-Inventar und Freigabe-Workflows im ServiceNow AI Control Tower, folgende ServiceNow-Funktionalitäten nutzen:

- **Risk Management:** Das Risikomanagement-Modul von ServiceNow wird genutzt, um für jede KI-Anwendung systematisch Risiken zu erfassen und zu bewerten. Basierend auf dem EU AI Act kann so etwa festgelegt werden, welche KI-Systeme als hochriskant gelten und priorisierte Aufmerksamkeit erhalten. Risiken wie z. B. diskriminierende Ergebnisse, Datenlecks oder Modellfehler werden im Risk Register dokumentiert und einer Bewertung unterzogen. Für identifizierte Risiken können direkt Maßnahmen oder Kontrollen hinterlegt werden. Dank der Integration des Risk Management mit dem AI Control Tower werden neu erfasste KI-Risiken automatisch der jeweiligen KI-Asset-Information im Control Tower zugeordnet.

So behält das Governance-Team in Echtzeit den Überblick über den Risikostatus jeder KI, von der Entwicklung bis zur Freigabe.

- **Policy and Compliance Management:** Das Policy & Compliance-Modul ermöglicht es, unternehmensweite KI-Policies und -Kontrollen zu definieren und mit regulatorischen Vorgaben abzugleichen. ServiceNow stellt hierfür vordefinierte Inhalte bereit, etwa ein AI Risk & Compliance Content Pack, das Anforderungen aus wichtigen KI-Regulierungen (z. B. EU AI Act, NIST AI RMF) als Kontrollziele in das System einspeist. Dadurch kann man frühzeitig eine Gap-Analyse durchführen: Welche bestehenden Kontrollen erfüllen schon die neuen KI-Vorgaben und wo bestehen Lücken? Neue oder angepasste Kontrollen werden dann im System hinterlegt.

Diese kombinierte Lösung ermöglicht es, schon vor einem KI-Go-Live ein umfassendes KI-Governance-Rahmenwerk anzuwenden. Alle relevanten Informationen und Nachweise (Inventar, Zuständigkeiten, Risikobewertungen, Kontrollen) sind in ServiceNow zentral verfügbar, wodurch relevante Verantwortliche vollständige Transparenz und Kontrolle über alle KI-Initiativen erhalten, sodass KI-Systeme von Anfang an im Einklang mit den internen Richtlinien und den Anforderungen des EU AI Act stehen.

USE CASE 3: KONTINUIERLICHE KI-ÜBERWACHUNG UND COMPLIANCE IM BETRIEB

Sobald KI-Systeme produktiv im Einsatz sind, müssen Unternehmen deren Performance und Veränderungen kontinuierlich im Blick behalten. Modelle selbst können sich im Laufe der Zeit verändern (z. B. durch sog. Model-Drift bei selbstlernenden Systemen), aber auch der Use Case für ein KI-System entwickelt sich weiter, was Änderungen an der KI nach sich ziehen kann. Der EU AI Act schreibt daher vor, dass Hochrisiko-KI laufend überwacht wird und bei neuen Erkenntnissen oder Änderungen eine erneute Evaluierung erfolgen muss. Das bedeutet, dass das Risikoprofil eines KI-Systems und somit die Überprüfung der Ein-

haltung von Compliance- und Kontrollanforderungen nicht nur einmalig vor Inbetriebnahme durchgeführt werden muss, sondern über den gesamten Lebenszyklus hinweg. Treten Zwischenfälle auf, etwa Fehlentscheidungen der KI oder Abweichungen von Richtlinien, so sind diese zu dokumentieren und zu beheben.

Für diese reaktive KI-Governance im laufenden Betrieb können folgende ServiceNow-Funktionen einen wesentlichen Beitrag leisten:

- **AI Governance Workspace:** Als Teil des AI Control Tower bietet ServiceNow einen dedizierten AI Governance Workspace. Dieser bietet einen Echtzeit-Überblick über den Status von KI-Modellen innerhalb des Unternehmens: Compliance-Kennzahlen, offene Risiken und Ereignisse je KI-System sind auf einen Blick verfügbar. Der Workspace unterstützt dabei, dass KI-Systeme den Unternehmensrichtlinien und regulatorischen Anforderungen entsprechen, wobei ein besonderer Fokus auf Datenschutz, Data Governance und ethischer KI-Nutzung liegt. Außerdem verwaltet der Workspace den gesamten Lebenszyklus von KI-Assets, vom Onboarding und Offboarding bis hin zur fortlaufenden Leistungsüberwachung. Er ermöglicht Unternehmen einen risikobasierten Ansatz für das Management von KI-Technologien und die Implementierung von KI-Lösungen. Damit erhalten alle relevanten Stakeholder einen ganzheitlichen Überblick über die Situation des Unternehmens und können fundierte Entscheidungen auf allen Ebenen treffen. Dies unterstützt eine optimierte Nachvollziehbarkeit und Erklärbarkeit von KI-Ergebnissen, was eine zentrale Anforderung an vertrauenswürdige und ethische KI ist und regulatorisch gefordert wird.

- **Kontinuierliches Compliance Monitoring:** Das Policy & Compliance-Modul von ServiceNow unterstützt regelmäßige Überprüfungen und Audits von KI-Systemen. Beispielsweise lassen sich wiederkehrende Kontrolltests einrichten, die automatisch prüfen, ob ein KI-Service weiterhin alle definierten Kontrollen einhält. Ergebnisse solcher Tests oder Audits werden im System dokumentiert. Wenn ein Kontrollziel verfehlt wird (etwa ein Verstoß gegen eine vorgeschriebene Datenschutzmaßnahme durch die KI), erzeugt das System ein Compliance-Issue zur Nachverfolgung.

- **Integration mit Risk Management:** Werden Abweichungen oder neue Risiken identifiziert, fließen diese direkt in das Risikomanagement ein. ServiceNow kann automatisch einen Risikoeintrag oder Vorfall erstellen, der den Befund beschreibt (z. B. „KI-Modell XY zeigt systematische Verzerrung gegenüber Gruppe Z“). Dieser Risk- oder Compliance-Issue wird im Risk Management Modul weiterbearbeitet indem Verantwortlichen Aufgaben zugewiesen werden, um Gegenmaßnahmen zu ergreifen (z. B. Retraining des Modells, Anpassung der Datenbasis oder zusätzliche Kontrollschritte). Durch die enge Verzahnung zwischen AI Control Tower und dem IRM-Modul sind solche Vorgänge transparent: Im AI Governance Workspace sieht man den aktuellen Risikostatus jeder KI und im Risk Management Modul sind alle technischen Details aus dem Control Tower verlinkt.
- **Integration mit IT Service Management:** Wird ein Fehler oder Zwischenfall im Betrieb einer KI-Anwendung erkannt, kann ServiceNow automatisch ein Ticket im IT Service Management (ITSM) erstellen (z. B. einen Incident). Dadurch wird sichergestellt, dass etwaige Störungen umgehend von den IT-/AI-Operations-Teams nachverfolgt und im Rahmen der etablierten Service-Management-Prozesse behoben werden. So schlägt das Unternehmen eine Brücke zwischen der KI-Governance und den operativen ITSM-Abläufen, damit KI-Störungen gezielt und zeitnah bearbeitet werden.

Mit diesem Zusammenspiel stellt ServiceNow sicher, dass KI-Systeme auch nach dem Go-Live verantwortungsvoll und nachhaltig betrieben werden. Es entsteht ein kontinuierlicher Kreislauf aus Überwachen, Erkennen, Reagieren: KI wird laufend überwacht, Erkenntnisse über Regelverstöße oder Performance-Probleme werden nahtlos ins Governance-System überführt und es erfolgen strukturierte Korrekturmaßnahmen. Dieser kontinuierliche Optimierungsprozess gewährleistet, dass KI-Anwendungen dauerhaft zuverlässig innerhalb definierter Risiko-Toleranzen agieren.

USE CASE 4: STEUERUNG VON KI-LIEFERANTENRISIKEN

Viele Organisationen beziehen KI-Technologien oder -Services von Drittanbietern, sei es der Zukauf von vortrainierten Modellen, die Nutzung von externen KI-APIs oder der Einsatz von Beratungsleistungen für KI. Trotz Auslagerung bleibt das Unternehmen letztlich dafür verantwortlich, dass auch diese externen KI-Komponenten den regulatorischen Anforderungen genügen. Gemäß EU AI Act müssen Unternehmen folglich die Risiken in ihrer Lieferkette im Auge behalten, insbesondere wenn ein externer KI-Anbieter z. B. keine ausreichenden Kontrollen vorweist oder seine KI mangelhaft dokumentiert ist. Die Herausforderung besteht darin, frühzeitig solche Lieferantenrisiken zu erkennen und geeignete Maßnahmen abzuleiten, um ggf. entgegenzusteuern bzw. ggf. einen zu risikobehafteten KI-Services nicht zu nutzen.

ServiceNow ermöglicht dies durch den Einsatz des Third Party Risk Management Moduls im Zusammenspiel mit den KI-Governance-Funktionen:

- **Third Party Risk Management (TPRM):** Das TPRM-Modul bietet einen standardisierten Prozess, um Lieferantenrisiken zu bewerten und zu managen. Für jeden (KI-)Lieferanten wird in ServiceNow ein Profil mit Risikoeinstufung geführt. Das Unternehmen kann maßgeschneiderte Risikobewertungen für KI-Zulieferer konfigurieren, etwa in Form von Fragebögen und Assessments, die vom Anbieter auszufüllen sind. Diese Fragebögen können u. a. KI-spezifische Kontrollfragen (z. B. Verfügt der Anbieter über ein eigenes KI-Risikomanagement?, Werden Trainingsdaten dokumentiert?, Existieren Bias-Prüfungen? etc.) enthalten. ServiceNow automatisiert die Aussendung und Auswertung solcher Assessments. Die Antworten des Vendors werden im TPRM-Modul validiert und fließen in einen Risiko-Score ein. So lässt sich objektiv nachvollziehen, welcher KI-Zulieferer potentiell hohe Risiken mit sich bringt. Bei auffälligen Antworten oder fehlenden Nachweisen kann das Unternehmen im TPRM-Modul Folgemaßnahmen definieren, etwa zusätzliche Audits beim Lieferanten oder vertragliche Auflagen.

- **Policy and Compliance für Drittanforderungen:** Über das Compliance-Modul können auch Anforderungen an Drittparteien formalisiert werden. Zum Beispiel lässt sich eine interne Policy definieren, dass alle KI-Drittanbieter bestimmte Mindeststandards erfüllen müssen (etwa Ethik-Richtlinien, Security-Zertifikate, Transparenzberichte, etc.). Diese Anforderungen werden als Kontrollziele in ServiceNow erfasst und den jeweiligen Lieferanten zugeordnet. Das System kann dann verfolgen, ob ein Vendor die geforderten Nachweise erbracht hat. Wird ein Kriterium nicht erfüllt, generiert ServiceNow ein entsprechendes Compliance Issue beim Lieferanten für das TPRM-Team.
- **Integration mit AI Control Tower:** Der AI Control Tower wird zum verbindenden Element zwischen interner KI-Governance und dem Drittanbietermanagement. Im zentralen KI-Inventar kann vermerkt werden, welche KI-Systeme von externen Anbietern stammen. Für solche Einträge lässt sich eine Verknüpfung zum entsprechenden Third Party Risk-Eintrag herstellen. Dadurch fließen Informationen zusammen: Im AI Control Tower sieht das KI-Governance-Team sofort, ob z. B. ein bestimmtes externes KI-Modul mit Einschränkungen belegt ist (etwa „nur für internen Testbetrieb freigegeben“ aufgrund offener Risiken). Gleichzeitig werden im TPRM-Modul automatisch Updates sichtbar, wenn sich bei einem KI-Asset die Risikobewertung ändert. Diese modulübergreifende Transparenz stellt sicher, dass alle KI-Risiken, ob intern oder via Drittanbieter einheitlich gemanagt werden.

Durch den kombinierten Einsatz von AI Control Tower und Third Party Risk Management erreicht das Unternehmen durchgängige Einblicke über den Stand der KI-Compliance in der gesamten Wertschöpfungskette. Kritische KI-Drittanbieter werden früh identifiziert und können gezielt überwacht oder substituiert werden. Gleichzeitig können vertrauenswürdige Partner eingebunden werden, da man ihre Compliance-Nachweise zentral prüft und abspeichert. Insgesamt unterstützt dieser Ansatz dabei, regulatorische Vorgaben und interne Sicherheitsstandards entlang der Lieferkette einzuhalten, ohne auf die Einbindung innovativer KI-Technologien von außen verzichten zu müssen.

FAZIT UND AUSBLICK

Der EU AI Act stellt Unternehmen vor umfangreiche Compliance-Aufgaben. Gleichzeitig bietet er die Chance, **KI systematisch und verantwortungsvoll zu managen**. Mit einer Plattform wie ServiceNow können die vielfältigen Anforderungen, von der Inventarisierung über Risikobewertung und Kontrollimplementierung bis zum dauerhaften Monitoring, in integrierte digitale Prozesse gegossen werden. Wichtig ist, frühzeitig zu starten: Zwar tritt die volle Wirksamkeit der meisten Regelungen des EU AI Act erst 2026 ein, doch Unternehmen sollten die Zeit nutzen. Insbesondere der Aufbau von Inventaren, Prozessen und Kontrollelementen braucht Zeit und Reife. Zudem gelten einzelne Bestimmungen, etwa die Verbote und ersten Transparenzpflichten, schon ab Anfang 2025.

DER GENERAL-PURPOSE AI CODE OF PRACTICE DER EU

Der **General-Purpose AI Code of Practice** der Europäischen Union ist eine freiwillige Selbstverpflichtung für Technologieunternehmen, die generative und allgemein einsetzbare KI-Systeme entwickeln oder betreiben. Ziel des Codes ist es, zentrale Prinzipien und Praktiken zur Sicherstellung einer verantwortungsvollen KI-Nutzung zu etablieren, insbesondere im Hinblick auf Transparenz, Datenschutz, Verantwortlichkeit, Erklärbarkeit und menschliche Aufsicht.

Dieser Verhaltenskodex wurde im Vorfeld des Inkrafttretens des EU AI Act geschaffen, um Unternehmen dabei zu unterstützen, frühzeitig praktikable und einheitliche Standards zur Umsetzung vertrauenswürdiger KI zu implementieren. Die Unterzeichnung des Codes signalisiert, dass sich teilnehmende Unternehmen öffentlich zu diesen Grundsätzen verpflichten und entsprechende Maßnahmen zur Umsetzung ergreifen.

ServiceNow gehört zu den Unternehmen, die diesen EU AI Code of Practice unterzeichnet haben. Für Unternehmen, die ServiceNow einsetzen, bedeutet dies, dass die Plattform bereits wesentliche Anforderungen und Prinzipien der verantwortungsvollen KI-Nutzung berücksichtigt und somit eine belastbare Grundlage bietet, um künftige regulatorische Anforderungen, insbesondere die des EU AI Act, einfacher erfüllen zu können.

ÜBER DIE AUTOREN



KENTARO ELLERT

Associate Director
+49 160 910 512 87
kentaro.ellert@protiviti.de

Kentaro Ellert unterstützt Unternehmen bei der Einführung von Governance Modellen und Compliance Management Systemen zur nachhaltig erfolgreichen Steuerung Künstlicher Intelligenz. Als Associate Director bei Protiviti im Bereich Technology Risk & Resilience unterstützt er Unternehmen dabei, die Anforderungen des EU AI Act in der Praxis umzusetzen und diese automatisiert zu überwachen. Als Practice Lead für die ServiceNow Risk Management Lösungen leitet er Implementierungsprojekte, um regulatorische Anforderungen workflow-gestützt mit ServiceNow zu realisieren.



ANDREJ GREINDL

Managing Director
+49 172 698 30 53
andrej.greindl@protiviti.de

Andrej Greindl unterstützt Unternehmen dabei, technologische Risiken zu beherrschen, regulatorische Anforderungen zu erfüllen und nachhaltige Resilienz aufzubauen. Als Managing Director bei Protiviti verantwortet er übergreifend die Bereiche Technology Risk & Resilience sowie Technology Audit. Sein Fokus liegt auf der Gestaltung und Prüfung von Informationssicherheits-, IT-Risikomanagement und IT-Compliance-Strukturen – sowohl aus beratender als auch aus prüferischer und implementierender Perspektive. Dazu zählen unter anderem Projekte zur Umsetzung des EU AI Act sowie die Prüfung der Anforderungen durch die Interne Revision.



SEBASTIAN MAYER

Managing Director
+49 162 276 58 55
sebastian.mayer@protiviti.de

Sebastian Mayer unterstützt Unternehmen dabei, regulatorische Anforderungen, technologische Komplexität und digitale Transformation in Einklang zu bringen. Als Managing Director bei Protiviti verantwortet er die Bereiche ServiceNow sowie Technology Strategy & Architecture. Sein Schwerpunkt liegt auf der Nutzung von ServiceNow als strategische Plattform für Resilienz, Automatisierung und integrierte Steuerung. Zudem setzt er sich mit der Frage auseinander, wie Künstliche Intelligenz die digitale Transformation in Organisationen nachhaltig gestalten kann.

ÜBER PROTIVITI

Protiviti (www.protiviti.com) ist ein weltweit tätiges Beratungsunternehmen, das fundiertes Fachwissen, objektive Erkenntnisse, einen maßgeschneiderten Ansatz und beispiellose Zusammenarbeit bietet, um Führungskräfte selbstbewusst in die Zukunft blicken zu lassen. Protiviti und seine unabhängigen und regional ansässigen Mitgliedsfirmen unterstützen Kunden mit Beratungsleistungen und Managed Solutions in Bezug auf Finanzen, Technologien, Operatives, Digitales, HR, Risiko und Interne Revision – mithilfe eines Netzwerkes von mehr als 90 Niederlassungen in über 25 Ländern.

Protiviti wurde zum zehnten Mal in Folge in die Liste der **Fortune 100 Best Companies to Work For®** aufgenommen und hat mehr als 80 Prozent der Fortune-100- und fast 80 Prozent der Fortune-500-Unternehmen betreut. Protiviti arbeitet ebenfalls mit Regierungsbehörden sowie kleineren und wachsenden Unternehmen zusammen, einschließlich solchen, die den Börsengang planen. Protiviti ist darüber hinaus eine hundertprozentige Tochter der **Robert Half Inc.** (NYSE: RHI).

© 2025 Protiviti Inc. Protiviti ist nicht als Wirtschaftsprüfungsgesellschaft zugelassen oder registriert und gibt keine Einschätzung zu Finanzberichten oder anderen Bestätigungsleistungen ab.

www.protiviti.de



© 2025 Protiviti GmbH