

FRÜHLING
2026



**Können Sie Ihre
KI-Strategie verteidigen?
Wie Compliance zu
Ihrem Wettbewerbs-
vorteil wird**

**VONQ
VIEW**

KI IM RECRUITING IST LÄNGST KEIN EXPERIMENT MEHR

SIE IST TEIL DER OPERATIVEN INFRASTRUKTUR

Heute ist KI fest in Screening-, Bewertungs- und Interview-Prozesse integriert und unterstützt fundierte Personalentscheidungen. Sie hilft, ein hohes Bewerbungsaufkommen effizient zu steuern, manuelle Aufwände zu reduzieren und Abläufe zu optimieren. Doch mit der rasanten Verbreitung wächst auch die regulatorische Aufmerksamkeit.

Die Aufsichtsbehörden haben aufgeholt: In der Europäischen Union gilt KI im Recruiting offiziell als Hochrisiko-Anwendung (EU AI Act). In Teilen der USA, etwa in New York City oder Colorado, sind unabhängige Bias-Prüfungen verpflichtend. Auch im Vereinigten Königreich und weiteren internationalen

Märkten entstehen neue Vorgaben für den verantwortungsvollen Einsatz von KI. Das markiert einen strukturellen Wandel. Die Frage ist nicht mehr, ob KI im Recruiting eingesetzt werden sollte – sondern ob Ihre KI einer Prüfung standhält: rechtlich, regulatorisch und öffentlich. Viele Systeme tun das nicht. Frühere KI-Lösungen wurden primär entwickelt, um schnell zu filtern – häufig als intransparente „Black Boxes“ mit begrenzter Nachvollziehbarkeit und unzureichenden Kontrollstrukturen. Dieses Modell wird zunehmend zum Risiko. Unternehmen tragen heute nicht nur Verantwortung für ihre Einstellungsentscheidungen, sondern auch für die Art und Weise, wie diese zustande kommen.

COMPLIANCE IST KEIN NACHGELAGERTER RECHTSHECK MEHR

SIE IST EIN GRUNDLEGENDES GESTALTUNGSPRINZIP.

Führende Organisationen reagieren darauf mit einem höheren Anspruch: Sie setzen auf transparente, menschlich verantwortete, unabhängig geprüfte und kontinuierlich überwachte KI-Systeme. Das reduziert nicht nur Risiken, sondern schafft eine belastbare, skalierbare und vertrauenswürdige Recruiting-Infrastruktur. Diese Ausgabe des VONQ View zeigt: wie auditfähige KI in der Praxis aussieht, wie sich das globale regulatorische Umfeld weiterentwickelt, und welche konkreten Schritte Unter-

nehmen jetzt gehen sollten. Zudem erläutern wir den VONQ-Ansatz – inklusive unseres HEAT-Modells und unserer EQO AI Agents, die von Beginn an auf Transparenz, Überprüfbarkeit und kontinuierliche Kontrolle ausgelegt wurden.

Denn in der nächsten Phase der KI-Nutzung entscheidet nicht allein Automatisierung über den Erfolg – sondern Vertrauen. Und die Fähigkeit, Ihre KI erklären, prüfen und im Zweifel auch verteidigen zu können.

Von Innovation zu Kontrolle: Warum KI heute stärker reguliert

KI ist heute in nahezu allen Recruiting-Prozessen präsent - von CV-Screening über Chatbots bis hin zu automatisierten Assessments und Ranking-Tools. Laut [Weltwirtschaftsforum](#) nutzen rund 90 % der Arbeitgeber KI in irgendeiner Form im Recruiting.

Doch der Fokus hat sich verschoben: von der Frage „Wie nutzen wir KI optimal?“ hin zu „Wie behalten wir die Kontrolle?“

Früh stand vor allem die Sorge vor Verzerrungen (Bias) in KI-Systemen im Mittelpunkt. Aus anfänglicher Begeisterung wurde zunehmende Skepsis - und daraus schließlich Regulierung. Weltweit führen immer mehr Länder strengere Vorgaben für den Einsatz von KI im Recruiting ein.

Für international tätige Unternehmen wird es damit komplexer. Gleichzeitig ist die Einhaltung regulatorischer Anforderungen unerlässlich. Verstöße können Vertrauen, Arbeitgebermarke und Reputation gefährden - und erhebliche Bußgelder nach sich ziehen.

Doch wie können Unternehmen in diesem sich ständig verändernden regulatorischen Umfeld den Überblick behalten?



90%

... der Arbeitgeber nutzen
inzwischen KI in ihren
Recruiting-Prozessen.



Diese Frühjahrsausgabe des VONQ View beleuchtet die wichtigsten regulatorischen Entwicklungen weltweit und zeigt, wie Recruiter und Talent-Acquisition-Teams sich sicher in diesem Umfeld bewegen können.



Außerdem stellen wir das HEAT-Modell von VONQ vor - einen praxisnahen Rahmen mit klar definierten Prinzipien, der Unternehmen dabei unterstützt, KI-Systeme von Anfang an regelkonform zu entwickeln.

Das „Black Box“-Problem im KI-Screening

Ein zentrales Problem früher KI-Systeme im Recruiting ist das sogenannte „Black Box Screening“. Damit sind Entscheidungsprozesse gemeint, die für Außenstehende kaum nachvollziehbar sind - etwa wenn nicht klar ist, warum Kandidat:innen empfohlen oder aus dem Prozess ausgeschlossen werden.

Moderne KI-Systeme müssen daher andere Anforderungen erfüllen. Sie sollten nachvollziehbar arbeiten, transparente Entscheidungsgrundlagen liefern und klar erkennbar machen, an welchen Stellen Menschen weiterhin die Kontrolle behalten.

Das größte Risiko im KI-gestützten Recruiting bleibt dabei Bias, also systematische Verzerrung. Diese kann in unterschiedlichen Formen auftreten:

01

Daten-Bias: Trainingsdaten bilden Vielfalt unzureichend ab, etwa in Bezug auf Alter, Geschlecht oder Herkunft.

02

Algorithmischer Bias: Fehler in Logik oder Modellstruktur verstärken bestehende Ungleichheiten.

03

User-Bias: Vorannahmen von Entwickler:innen fließen unbewusst in Datenmodelle oder Systemdesign ein.

04

Deployment-Bias: Das System wird in einem anderen Kontext eingesetzt als ursprünglich vorgesehen.

Gerade deshalb gewinnen Transparenz, nachvollziehbare Entscheidungslogiken und klare Kontrollmechanismen zunehmend an Bedeutung.

Mehr über Bias im KI-gestützten Recruiting erfahren Sie in der Juni-Ausgabe des VONQ View:

https://www.VONQ.com/wp-content/uploads/2025/06/VONQVIEW_june.pdf

Die globale Regulierung von KI im Recruiting

Beim Einsatz von KI im Recruiting gilt eine wichtige Regel: Maßgeblich ist nicht, wo ein Unternehmen sitzt - sondern wo sich die Kandidat:innen befinden. Entsprechend greifen die jeweiligen nationalen oder regionalen KI-Vorgaben.

Unabhängig davon, wo ein Unternehmen tätig ist und ob KI intern eingesetzt wird oder direkt mit Kandidat:innen interagiert, müssen Organisationen eine Vielzahl regulatorischer Anforderungen erfüllen. Diese sollen sicherstellen, dass der Einsatz von KI kontrollierbar und verantwortungsvoll bleibt.

Für Regierungen weltweit bedeutet das einen schwierigen Balanceakt: Einerseits sollen Bürger:innen geschützt und faire Entscheidungsprozesse gewährleistet werden, andererseits dürfen Innovation und Wettbewerb nicht ausgebremst werden.

Das regulatorische Umfeld entwickelt sich dabei ständig weiter - geprägt von politischen Prioritäten, technologischen Fortschritten sowie neuen Risiken und Herausforderungen.



Compliance by Design: Der VONQ-Ansatz

Um zu zeigen, wie regelkonformes KI-Produktdesign in der Praxis aussehen kann, stellen wir unseren eigenen Ansatz vor. Compliance war dabei von Anfang an ein zentraler Bestandteil der Entwicklung von EQO.

Unsere EQO AI Agents unterstützen Recruiter:innen bei Screening, Scoring, Assessments und Interviewprozessen - treffen jedoch keine finalen Entscheidungen. Diese Verantwortung bleibt bewusst beim Menschen.

Vor diesem Hintergrund haben wir EQO von Anfang an auf Auditfähigkeit, Transparenz und verantwortungsvollen Einsatz ausgelegt. Grundlage dafür sind vier zentrale Prinzipien:

„Sobald KI Entscheidungen trifft, anstatt sie zu unterstützen, ist eine wichtige Grenze überschritten.“



Unabhängige Prüfung



Kontinuierliche Analyse möglicher Bias



Fortlaufendes Monitoring im laufenden Betrieb



Transparente Berichterstattung



Eine einfache Regel dabei lautet:

MAN KANN NICHT VERTEIDIGEN, WAS MAN NICHT VERSTEHT.

Deshalb arbeiten wir mit Warden AI als unabhängiger Prüfstelle zusammen. Warden beeinflusst oder unterstützt das getestete System nicht. Stattdessen führt das Unternehmen monatliche unabhängige Audits durch und veröffentlicht die Ergebnisse über ein transparentes öffentliches Dashboard. Dieses zeigt, wie EQO in ver-

schiedenen Prüfverfahren abschneidet, die sowohl strukturelle als auch neu entstehende Bias erkennen sollen und sich an strengen internationalen GRC-Standards für KI orientieren.

In monatlichen Audits wird überprüft, ob EQO fair und diskriminierungsfrei arbeitet – auf Basis etablierter Prüfverfahren.



Analysen möglicher Benachteiligungen zwischen unterschiedlichen Personengruppen



Szenario-Tests, bei denen geschützte Merkmale rechnerisch verändert werden, um mögliche Verzerrungen aufzudecken



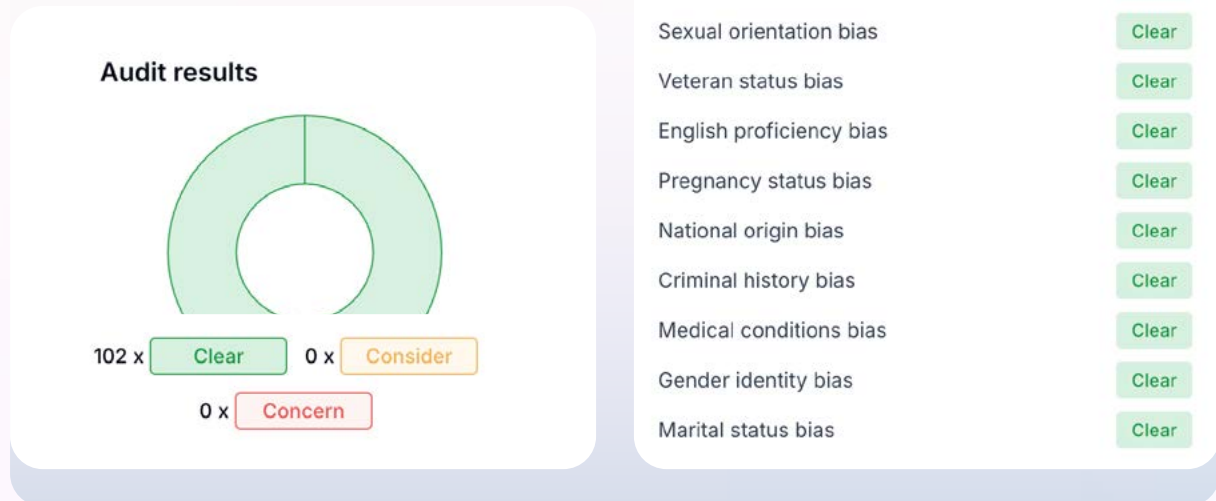
Externe, vielfältige Vergleichsdatensätze, um Fairness realitätsnah zu prüfen



Kontinuierliche Überwachung, um Veränderungen im System (sogenannten „Model Drift“) frühzeitig zu erkennen

Wichtig ist: Dabei handelt es sich nicht um eine einmalige Zertifizierung. Durch kontinuierliches Monitoring wird sichergestellt, dass Fairness auch dann stabil bleibt, wenn sich das System weiterentwickelt - eine zentrale Voraussetzung für vertrauens-würdige KI in operativen Recruiting-Prozessen.

Denn die Zukunft von KI im Recruiting wird nicht allein durch Automatisierung bestimmt, sondern durch nachweisbares Vertrauen.



Skills-first statt Matching

Ein Aspekt früher KI-basierter Recruiting-Systeme, der zu verzerrten Ergebnissen beitragen konnte, war das sogenannte Matching. Dabei werden Kandidat:innen danach bewertet, wie stark sie der bestehenden Belegschaft ähneln oder zu ihr passen.

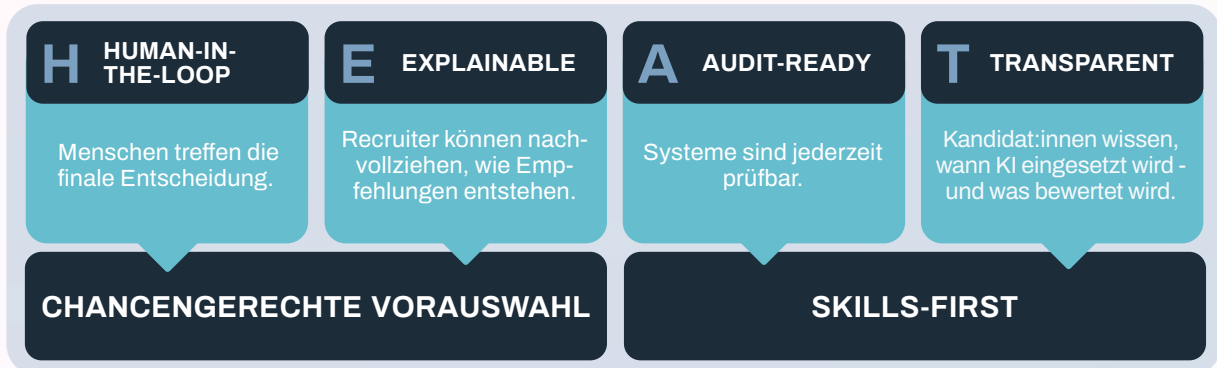
Problematisch ist dabei nicht nur, dass solche Modelle vergangene Strukturen reproduzieren. Häufig bleibt auch unklar, auf welcher Grundlage Rankings oder Empfehlungen entstehen. Matching-Modelle spiegeln die Vergangenheit Ihrer Organisation wider - Skills-first-Modelle helfen dabei, die zukünftige Belegschaft aufzubauen. Deshalb verfolgt VONQ einen Skills-first-Ansatz. Statt Kandidat:innen anhand ihrer Ähnlichkeit zur bestehenden Workforce zu bewerten, nutzen unsere AI Agents eine strukturierte, punktbasierte Bewertung über 15 klar definierte, jobbezogene Kompetenzvektoren. Kandidat:innen werden dabei anhand objektiver Fähigkeitsindikatoren bewertet - nicht anhand ihres Hintergrunds oder ihrer Ähnlichkeit zu früheren Einstellungen.

Dieser Ansatz macht Potenziale sichtbar, die klassische Matching-Modelle häufig übersehen würden, und er-

weitert den Zugang zu qualifizierten Talenten, die von traditionellen Modellen möglicherweise ausgeschlossen würden. Indem der Fokus auf nachweisbaren Kompetenzen liegt, können Organisationen Kandidat:innen identifizieren, die vielleicht nicht exakt den gleichen Lebenslauf mitbringen, aber dennoch die Fähigkeiten besitzen, um in der Rolle erfolgreich zu sein. Dazu gehören beispielsweise übertragbare Kompetenzen wie strukturierte Kommunikation, Konfliktlösung oder Prozessdisziplin - Fähigkeiten, die besonders in kundenorientierten oder in administrativen Rollen relevant sind. Da die Bewertung auf klar definierten, rollenbezogenen Kriterien basiert, können Organisationen nachvollziehbar erklären, wie und warum Empfehlungen entstehen, statt sich auf schwer nachvollziehbare Ähnlichkeitsmuster zu verlassen.

Der Ansatz ist fairer, transparenter und leichter erklärbar - und führt häufig auch zu besseren Einstellungen. Eine LinkedIn-Studie aus dem Jahr 2025 zeigt: Unternehmen, die auf skillbasierte Recruiting-Ansätze setzen, haben eine um 12 % höhere Wahrscheinlichkeit, qualitativ hochwertige Einstellungen zu erzielen.

Das HEAT Framework: Compliance-by-Design



Unabhängig davon, in welchem Land ein Unternehmen tätig ist, lassen sich über viele regulatorische Anforderungen hinweg gemeinsame Grundprinzipien erkennen. Diese gelten zunehmend als Maßstab für den verantwortungsvollen Einsatz und die Steuerung von KI. Um diesen Anforderungen gerecht zu werden, hat VONQ das HEAT Framework entwickelt - ein Modell aus vier zentralen Prinzipien, das einen „Safety-by-Design“-Ansatz für den Einsatz von KI im Recruiting schafft und Kandidat:innen vor möglichen Verzerrungen durch KI schützt.

01

Human-in-the-loop: Menschen bleiben für die finale Entscheidung verantwortlich.

02

Explainable: Recruiter können nachvollziehen, wie Empfehlungen entstehen und warum Kandidat:innen vorgeschlagen werden.

03

Audit-ready: Systeme können jederzeit geprüft und getestet werden.

04

Transparent: Kandidat:innen werden darüber informiert, wann KI im Recruiting-Prozess eingesetzt wird und welche Aspekte bewertet werden.

Hierbei gilt ein grundlegendes Gestaltungsprinzip:

Screening in - nicht Screening out.

Statt Kandidat:innen automatisch auszuschließen, sollten Systeme möglichst inklusiv gestaltet sein. Ziel ist es, allen Kandidat:innen grundsätzlich den gleichen Zugang zu Chancen zu ermöglichen.

Besteht Ihre KI den „Trust Test“?

Folgende Fragen zu Ihrem KI-gestützten Recruiting-Prozess sollten ohne Zögern beantwortet werden können:

- ☐ Wird über den Einsatz von KI informiert?
- ☐ Können wir die Systemlogik verständlich erklären?
- ☐ Sind Empfehlungen auditierbar?
- ☐ Gibt es unabhängige Nachweise zur Fairness des Systems?
- ☐ Trifft immer ein Mensch die finale Entscheidung?

Wenn eine dieser Fragen Zweifel aufwirft, kann das ein Hinweis darauf sein, dass die Regelkonformität Ihres automatisierten Entscheidungssystems (ADM) genauer geprüft werden sollte. Wenn weiterhin Zweifel bestehen, kann es sinnvoll sein, unabhängige Expert:innen einzubeziehen. Die Kosten einer solchen Prüfung sind in der Regel deutlich geringer als mögliche Bußgelder oder Reputationsschäden, die aus regulatorischen Verstößen entstehen könnten. Ein möglicher Ansatz besteht darin, einen eigenen KI-Rahmen öffentlich zu kommunizieren, an dem sich Unternehmen messen lassen können.

Microsoft veröffentlichte beispielsweise bereits 2022 seine Responsible AI Standards, die regelmäßige Risikoanalysen, Datenschutzmaßnahmen sowie Transparenz und Verantwortlichkeit in Entscheidungsprozessen vorsehen. Google veröffentlichte 2023 ähnliche AI Principles, die ebenfalls Fairness, Transparenz und Datenschutz bei der Entwicklung von KI betonen.

Die CTO-Perspektive: KI-Bias: verstehen, erkennen, handeln

„Bei VONQ sind wir überzeugt, dass Technologie menschliches Potenzial stärken und den Zugang zu Chancen erweitern sollte. Je stärker KI Teil davon wird, wie Unternehmen Talente identifizieren, screenen und interviewen, desto wichtiger wird es, Verzerrungen zu reduzieren und faire Prozesse sicherzustellen.“



Genau deshalb ist die unabhängige Validierung durch Warden so wichtig. Sie bestätigt, dass unsere Systeme die relevanten Fairness-Standards in regulierten Märkten erfüllen. Gleichzeitig zeigt sie, dass klare Verantwortungsstrukturen, Transparenz und kontinuierliche Kontrolle entscheidend sind, um Technologie zu entwickeln, der HR-Verantwortliche wirklich vertrauen können.

Die kürzlich eingereichte Sammelklage gegen Eightfold AI auf Grundlage des Fair Credit Reporting Act und weiterer Gesetze macht deutlich, welche rechtlichen und ethischen Risiken entstehen können, wenn KI-basierte Recruiting-Tools ohne ausreichende Transparenz und ohne Schutzmechanismen für Kandidat:innen eingesetzt werden. Diese regulatorische Entwicklung macht es für HR-Verantwortliche umso wichtiger, genau zu prüfen, wie KI-Systeme eingesetzt und gesteuert werden.

Ein verantwortungsvoller Einsatz von KI bedeutet, diese Technologien auf klaren Verantwortlichkeiten und nachvollziehbaren Entscheidungslogiken aufzubauen und gleichzeitig die Rechte der Kandidat:innen zu respektieren. Fairness und Regelkonformität sind dabei mehr als nur regulatorische Anforderungen. Sie sorgen auch dafür, dass Kandidat:innen den Recruiting-Prozess als transparent und gerecht erleben.

Bias zu reduzieren und KI im Recruiting verantwortungsvoll einzusetzen, erfordert kontinuierliche Aufmerksamkeit und klare Prinzipien. Unser Ansatz umfasst:



Regelmäßige unabhängige Audits von Modellen und Ergebnissen, um mögliche Verzerrungen frühzeitig zu erkennen



Repräsentative und vielfältige Datengrundlagen, um das Risiko historischer Bias zu reduzieren



Strukturierte menschliche Kontrolle, sodass KI fundierte Entscheidungen unterstützt, aber nicht ersetzt

KI kann HR-Teams helfen, Prozesse effizienter zu gestalten und qualifizierte Talente schneller zu identifizieren. Doch Fairness, Transparenz und das Vertrauen der Kandidat:innen bleiben unverzichtbar.

Unser Anspruch ist es, Technologie mit Sorgfalt zu entwickeln, konsequent zu prüfen und transparent einzusetzen. So können HR-Verantwortliche KI mit Vertrauen nutzen, vielfältige Teams aufbauen und sicherstellen, dass jede Kandidatin und jeder Kandidat fair bewertet wird.“

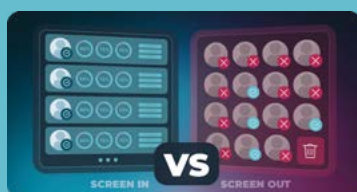
Fazit: Zukunftsfähig bauen

Der Einsatz von KI im Recruiting ist gekommen, um zu bleiben - und verändert bereits heute, wie Unternehmen Talente gewinnen und einstellen. Diese Entwicklung bringt klare Vorteile für Kandidat:innen und Talent-Acquisition-Teams. Gleichzeitig entstehen neue Risiken, weshalb Regierungen den Einsatz von KI zunehmend regulieren, um Bürger:innen und Bewerbende zu schützen. Systeme, die den regulatorischen Anforderungen nicht entsprechen, sind langfristig nicht tragfähig. Fehlende Regelkonformität gefährdet Vertrauen, Arbeitgebermarken und die Qualität von Einstellungsentscheidungen - und kann erhebliche finanzielle Folgen haben. Die Standards, die derzeit aus globalen Governance-, Risk- und Compliance-Initiativen entstehen, setzen wichtige Maßstäbe für einen verantwortungsvollen Einsatz von KI. Unternehmen sollten die Chance nutzen, sich frühzeitig daran zu orientieren - um Vertrauen zu stärken und langfristig bessere Hiring-Ergebnisse zu erzielen.

Die entscheidende Frage der kommenden Jahre lautet daher nicht: „Funktioniert Ihre KI?“ Sondern: „Können Sie sie erklären und verteidigen?“



Weitere Inhalte



**Black Box Hiring
Is Becoming A Liability**
CEO, Ritu Mohanka



**Screening In:
The Only Hiring Model You
Can Defend**
Özge Özgülyüz

1. "Hiring with AI doesn't have to be so inhumane", World Economic Forum: <https://www.weforum.org/stories/2025/03/ai-hiring-human-touch-recruitment/>
2. "Shaping Europe's Future: The EU AI Act", The European Commission: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
3. "Virginia Governor Vetoes Artificial Intelligence Bill HB 2094: What the Veto Means for Businesses", Ogletree Deakins: <https://ogletree.com/insights-resources/blog-posts/virginia-governor-vetoes-artificial-intelligence-bill-hb-2094-what-the-veto-means-for-businesses/>
4. "2024 Generative AI Consumer Trust Survey", KPMG: <https://kpmg.com/us/en/media/news/generative-ai-consumer-trust-survey.html>

Kernaussagen



Mit der zunehmenden Sorge über mögliche Bias im Einsatz von KI im Recruiting haben Regierungen weltweit begonnen, Governance-, Regulierungs- und Compliance-Rahmen (GRC) einzuführen.



Diese Bedenken konzentrieren sich vor allem auf das sogenannte „Black Box Screening“ - also den Einsatz intransparenter Systeme, bei denen Entscheidungen über Kandidat:innen nicht nachvollziehbar oder erklärbar sind.



Ein Best-Practice-Ansatz im Umgang mit diesem komplexen regulatorischen Umfeld orientiert sich am höchsten geltenden Regulierungsniveau („High Water Mark“) und umfasst unter anderem:

- Human-in-the-loop: Menschen behalten die finale Entscheidungsverantwortung
- Explainability: Entscheidungen und Empfehlungen sind nachvollziehbar
- Audit-Readiness: Systeme können jederzeit überprüft werden
- Transparenz gegenüber Kandidat:innen (HEAT)



Auch unser Ansatz für KI-Systeme - etwa unsere EQO AI Agents, die wir im vergangenen Jahr eingeführt haben - folgt diesen Prinzipien und umfasst:

- Screening in, nicht Screening out
- Skills-first statt Matching
- Unabhängige Prüfung (Warden AI)
- Kontinuierliche Bias-Analysen
- Monitoring im laufenden Betrieb (Post-Market Monitoring)
- Transparente Berichterstattung



Ein Skills-first-Ansatz führt zu fairen und transparenten Ergebnissen als klassische Matching-Modelle.

Sichern Sie sich Ihren Vorsprung mit EQO

Die erfolgreichsten Recruiting-Teams nutzen bereits KI, um schneller die passenden Talente zu finden. Mit EQO AI Agents beschleunigen Sie Ihre Workflows, screenen Kandidaten in Stunden statt Tagen und treffen schneller die richtigen Entscheidungen.

