

Nasze podejście do odpowiedzialnej sztucznej inteligencji (AI) i unijnego aktu o AI (EU AI Act)

Nasz Program Odpowiedzialnej Sztucznej Inteligencji (AI) jest integralną częścią ogólnej strategii ESG Future Processing, zgodną z czterema filarami: Ludzie, Etyczna Doskonałość, Zaangażowanie Społeczne i Zrównoważony Rozwój Środowiska, co zostało ujawnione w naszym raporcie ESG. Uznajemy społeczne, zarządcze i środowiskowe skutki działania AI i będziemy integrować wskaźniki ESG w monitorowaniu naszych rozwiązań opartych na sztucznej inteligencji.





Etyczna i odpowiedzialna AI w Future Processing

Future Processing zobowiązuje się do tworzenia rozwiązań AI, które są etyczne, przejrzyste i w pełni zgodne z obowiązującymi przepisami prawnymi. Jako odpowiedzialny partner technologiczny, aktywnie dbamy o to, by nasze systemy i usługi AI spełniały najwyższe standardy uczciwości i odpowiedzialności, zgodnie z obecnymi i nadchodzącymi przepisami europejskimi - w szczególności unijnym aktem o AI (EU AI Act).

Odpowiedzialny rozwój AI to nie tylko kwestia doskonałości technicznej, ale również zgodności z prawem, etyki i odpowiedzialności społecznej. Nasze podejście do AI opiera się więc na kluczowych ramach regulacyjnych, zasadach etycznych i najlepszych praktykach branżowych.

Będziemy aktywnie angażować interesariuszy (klientów, pracowników, społeczeństwo obywatelskie) w identyfikację istotnych ryzyk i szans związanych ze zrównoważonym rozwojem, obejmujących także etykę AI. Wnioski te zostaną uwzględnione w corocznej analizie istotności (Materiality Assessment) i ujawnione w sposób przejrzysty.

Podejście FP do odpowiedzialnej AI i zgodności z przepisami prawa

Wdrożyliśmy kompleksowe ramy zarządzania oraz procedury wewnętrzne odzwierciedlające obowiązki i zasady określone w unijnym akcie o AI. Nasze podejście obejmuje:

1. Klasyfikacja ryzyka systemów AI

Klasyfikujemy systemy AI zgodnie z kategoriami ryzyka określonymi w EU AI Act (minimalne, ograniczone, wysokie i zakazane). Dla każdego projektu przeprowadzamy szczegółową ocenę ryzyka już na wczesnym etapie rozwoju, co zapewnia, że nie tworzymy systemów zakazanych i stosujemy wymagane zabezpieczenia dla zastosowań wysokiego ryzyka.

2. Etyka i Prywatność w Projektowaniu (Ethics & Privacy by Design)

Kwestie etyczne, zgodność z przepisami oraz ochrona danych są uwzględniane na każdym etapie procesu rozwoju AI — od pierwszych warsztatów z klientem aż po wdrożenie systemu. Stosujemy zasady „Privacy by Design” oraz „Ethics by Design”, by zapewnić odpowiedzialne wykorzystanie technologii AI.

3. Przejrzystość i możliwość wyjaśnienia (Explainability)

Zapewniamy zgodność z wymogami przejrzystości określonymi w AI Act. Użytkownicy są informowani, gdy mają do czynienia z systemami AI, a my dostarczamy jasne i zrozumiałe informacje o sposobie ich działania. Dla systemów wysokiego ryzyka utrzymujemy szczegółową dokumentację oraz wdrażamy mechanizmy umożliwiające wyjaśnienie działania algorytmów, by uniknąć tzw. „czarnej skrzynki”.

4. **Sprawiedliwość i ograniczanie uprzedzeń (bias)**

Aktywnie przeciwdziałamy uprzedzeniom algorytmicznym i dyskryminacji. Nasze procedury obejmują techniki wykrywania uprzedzeń, staranny dobór i weryfikację danych treningowych oraz dokładne testowanie wyników modeli, co pomaga nam tworzyć sprawiedliwe i godne zaufania rozwiązania AI.

5. **Odpowiedzialność i nadzór ludzki**

Dla każdego rozwijanego systemu AI ustanawiamy jasne ramy odpowiedzialności. W naszych rozwiązaniach zapewniamy ludzki nadzór tam, gdzie jest to potrzebne, tak aby kluczowe decyzje nie były podejmowane wyłącznie przez systemy automatyczne.

6. **Odpowiedzialne korzystanie z AI generatywnej**

Dla rozwiązań generatywnych stosujemy szczególne wytyczne etyczne, by zapobiegać tworzeniu szkodliwych, stronicznych lub wprowadzających w błąd treści. Podejmujemy środki przeciwdziałające dezinformacji, deepfake'om i nieuprawnionemu generowaniu treści.

7. **Dokumentacja regulacyjna i gotowość audytowa**

Prowadzimy rejestr wszystkich rozwijanych systemów AI — zarówno na potrzeby wewnętrzne, jak i dla klientów — by zapewnić przejrzystość i gotowość do audytów. Pozwala nam to wykazać zgodność z wymogami dokumentacyjnymi i rejestracyjnymi AI Act, szczególnie w przypadku systemów wysokiego ryzyka.

8. **Podnoszenie świadomości i szkolenia pracowników**

Inwestujemy w programy szkoleń, by rozwijać kompetencje AI wśród naszych pracowników. Nasze zespoły regularnie uczestniczą w szkoleniach dotyczących etyki AI, wymogów regulacyjnych oraz najlepszych praktyk, dzięki czemu nasza wiedza rozwija się razem z otoczeniem prawnym.

9. **Bezpieczeństwo i zarządzanie ryzykiem**

Stosujemy rygorystyczne testy i procedury zarządzania ryzykiem, aby identyfikować i ograniczać podatności na zagrożenia w systemach AI. Obejmuje to ochronę przed atakami typu „adversarial” oraz niewłaściwym wykorzystaniem systemów.

10. **Monitorowanie przepisów i zarządzanie zgodnością**

Aktywnie śledzimy zmiany regulacyjne dotyczące AI. Nasz zespół ds. zarządzania zapewnia, że nasze polityki, procedury i rozwiązania są na bieżąco aktualizowane zgodnie z nowymi wymogami prawnymi i standardami etycznymi.

11. **Prawa człowieka**

Nasze podejście do odpowiedzialnej AI zostanie zintegrowane z procesem należytej staranności w zakresie praw człowieka (Human Rights Due Diligence), zgodnie z Wytycznymi ONZ (UNGPs) oraz wskaźnikiem GRI 410-1. Będziemy analizować ryzyka związane z prywatnością, uprzedzeniami, dyskryminacją i nadzorem w kontekście rozwoju AI.



12. Wpływ AI na emisje GHG i środowisko

Future Processing jest świadome wpływu technologii AI na środowisko, w szczególności w zakresie zużycia energii i emisji dwutlenku węgla wynikających z trenowania modeli AI i działania centrów danych. Będziemy mierzyć i raportować te emisje w ramach naszego inwentarza GHG, zgodnie z GHG Protocol.

Ocena wpływu AI na prawa człowieka – dla systemów AI o dużym wpływie

W przypadkach, gdy Future Processing angażuje się w rozwój lub wdrożenie systemów AI na dużą skalę lub o znacznym wpływie, przeprowadzimy formalną **Ocenę Wpływu AI (AIIA – AI Impact Assessment)**. Proces ten obejmie kompleksową analizę potencjalnych skutków działania systemu AI na podstawowe prawa człowieka, zgodnie z wymogami EU AI Act oraz rekomendacjami OECD AI Principles.

Ocena wpływu AI będzie uwzględniać m.in.:

- Potencjalny wpływ na prywatność, wolność wypowiedzi i niedyskryminację,
- Ryzyka dla autonomii człowieka, godności i równości,
- Możliwe skutki społeczne i ekonomiczne, w tym wpływ na zatrudnienie i grupy wrażliwe,
- Środki zaradcze w celu ograniczenia negatywnych skutków.

Wyniki oceny zostaną udokumentowane, udostępnione zainteresowanym stronom i zintegrowane z naszymi procesami zarządzania ryzykiem i ładu korporacyjnego. W stosownych przypadkach zaangażujemy także zainteresowane strony lub ich przedstawicieli, by uwzględnić ich prawa i obawy.

Przygotowanie na przyszłe obowiązki wynikające z AI Act

Aktywnie przygotowujemy się do spełnienia przyszłych wymogów wynikających z kolejnych etapów wdrażania unijnego aktu o AI, w tym:

- Zgodność z przepisami dotyczącymi ogólnego przeznaczenia modeli AI (General-Purpose AI Models, GPAI), które wejdą w życie w sierpniu 2025 r.,
- Zapewnienie, że wszystkie systemy AI wysokiego ryzyka spełniają rygorystyczne wymogi dotyczące dokumentacji technicznej, oceny zgodności oraz monitorowania po wdrożeniu,
- Ciągła ocena i aktualizacja procesów rozwoju AI w celu odzwierciedlenia najnowszych zmian regulacyjnych i standardów etycznych. latest regulatory updates and ethical standards.