

**„Zastosowanie autonomicznej AI w procesie pracy a odpowiedzialność cywilna
pracodawcy wobec osób trzecich”**

Kierownik projektu: dr Iwona Gredka-Ligarska

Streszczenie popularnonaukowe

Celem projektu badawczego jest ustalenie jaki model odpowiedzialności cywilnej powinien obowiązywać w prawodawstwie Unii Europejskiej, gdy autonomiczna AI (sztuczna inteligencja) wykorzystywana przez pracodawcę w procesie pracy wyrządzi szkodę osobie trzeciej. Obecnie odpowiedzialność cywilna za powstające szkody jest rozczłonkowana i dostosowana do znanego dotychczas podziału na ludzi i rzeczy. AI nie jest ani rzeczą, ani też nigdy nie stanie się człowiekiem. Wiele wskazuje natomiast na to, że AI w procesie pracy zastąpi nie tylko rzeczy (maszyny; urządzenia), ale również ludzi (pracowników). Z pewnością nie nastąpi to od razu i nie będzie dotyczyć wszystkich zawodów, ale w tym kontekście dotychczasowe dopasowanie zasad i modeli odpowiedzialności cywilnej do podziału na ludzi i na rzeczy może się nie sprawdzić. W prawodawstwie Unii Europejskiej dotychczas nie zdecydowano się na nadanie AI osobowości prawnej, połączonej z samodzielnym ponoszeniem przez AI odpowiedzialności cywilnej za wyrządzane szkody. Nie można jednak z góry wykluczyć konieczności opracowania całkowicie nowego, nieznanego dotychczas modelu odpowiedzialności autonomicznej AI. Być może zasadne będzie także zastąpienie różnych rodzajów odpowiedzialności cywilnej, jednym tylko rodzajem odpowiedzialności, związanym z autonomiczną AI. Hipoteza badawcza opiera się na założeniu, że odpowiedzialność musi być dostosowana do działania w pełni autonomicznej AI, która w określonym obszarze podejmować będzie samodzielne decyzje, bez udziału czynnika ludzkiego. Szczegółowym celem badawczym jest udzielenie odpowiedzi na pytanie, kto powinien odpowiadać za szkody wyrządzone osobie trzeciej przez autonomiczną AI, wykorzystywaną przez pracodawcę w procesie pracy i na jakiej zasadzie należy oprzeć tę odpowiedzialność.

Niektóre systemy AI funkcjonują tylko w świecie wirtualnym (np. asystenci głosowi, oprogramowanie do analizy obrazu, wyszukiwarki, systemy rozpoznawania mowy i twarzy), podczas gdy inne systemy AI są wbudowane w sprzęt (np. zaawansowane roboty, samochody autonomiczne, drony). Ważny jest również podział na tzw. słabą AI oraz na silną AI. Słaba AI jest określana jako zadaniowa lub ilościowa. Natomiast silna AI obejmuje sieci modelowane na wzór ludzkiego mózgu. Chociaż do chwili obecnej nie stworzono jeszcze takiego systemu, który można byłoby uznać za silną AI, naukowcy podkreślają, że jest to tylko kwestią czasu.

Uzasadnieniem dla podjęcia przedstawionego problemu badawczego jest okoliczność, że w UE nie zdecydowano się na nadanie AI osobowości prawnej, połączonej z samodzielnym ponoszeniem przez AI odpowiedzialności cywilnej za wyrządzane szkody. AI zakwalifikowano póki co – nie jako cyfrową osobę prawną, lecz – jako produkt i skoncentrowano się na odpowiedzialności za produkt oraz na odpowiedzialności operatora systemu AI. Rodzi się wobec tego pytanie czy prawodawstwo UE nie jest zbyt zachowawcze i czy pozwoli na budowanie pełnego zaufania do AI oraz szerokiego wykorzystywania AI w procesie pracy. Unia Europejska ma ambicje aby być w awangardzie w zakresie rozwoju w

pełni autonomicznej AI. Cel taki chce osiągnąć m.in. dzięki prawodawstwu, które będzie kompatybilne z rozwojem najnowszych technologii AI. Należy w związku z tym zbadać jakie rozwiązania w zakresie odpowiedzialności AI są zgłaszane w odniesieniu do prawodawstwa światowych liderów w zakresie technologii AI takich jak USA, Chiny oraz Singapur. Propozycje doktryny dotyczące obcych systemów prawnych mogą być pomocne dla ustalenia czy kierunek obrany przez UE (odpowiedzialność za AI = odpowiedzialność za produkt) jest słuszny. A jeśli nie, jakie rozwiązanie można w tym zakresie zaproponować.

Wyniki przeprowadzonych badań pozwolą odpowiedzieć na pytanie: czy możliwe jest optymalne uregulowanie odpowiedzialności AI przy użyciu jednego ze znanych już modeli odpowiedzialności, a jeśli tak to który model nadaje się do tego najlepiej. Jeśli jednak wniosek płynący z badań byłby taki, że znane modele odpowiedzialności są niekompatybilne w stosunku do odnotowanych problemów związanych z AI, to w ostatnim etapie pracy badawczej sformułowana będzie propozycja nowego modelu odpowiedzialności autonomicznej AI.