

Czy Twój samochód-robot powinien Cię zabić, by ratować innych?

18 maja 2014

Czy w krytycznej sytuacji automatyczny samochód powinien poświęcić życie kierowcy dla wyższego dobra? Pytanie trudne, a przecież automatyczne samochody są coraz bliżej.

Różne problemy etyczne związane z robotami są roztrząsane od lat, ale nigdy nie wydawały się one tak ważne jak dziś. Automatyczne samochody Google zaczynają już jeździć po miastach. Powstaje prawo regulujące korzystanie z nich. Lada dzień będziemy powierzać robotom swoje bezpieczeństwo i życie. Czy nasz robot powinien je chronić za wszelką cenę?

Mniej więcej takie pytanie zadał sobie Patrick Lin, specjalista od etyki z California Polytechnic State University, na łamach serwisu Wired.

AUTO MA ZABIĆ KIEROWCĘ?

Lin zaproponował następujący eksperyment myślowy. W wyniku jakiejś nieszczęśliwej awarii Twoje automatyczne auto traci część kontroli. Są przed nim dwa samochody i jest pewne, że auto uderzy w jeden z nich. Czy powinno uderzyć w samochód wiozący całą rodzinę, czy w mały pojazd, w którym jest tylko kierowca?

Możemy nieco zmodyfikować eksperyment. Załóżmy, że automatyczne auto może zboczyć z drogi i zjechać w przepaść, zabijając tylko kierowcę i oszczędzając np. auto z całą rodziną. Na tej podstawie można sformułować ogólne pytanie. Czy robot powinien zabić właściciela, by oszczędzić więcej ludzkich istnień? Czy robot w ogóle powinien dokonywać ocen i decyzji, które byłyby podyktowane wyższym dobrem?

Odpowiedź na to pytanie nie jest łatwa. Z jednej strony możemy powiedzieć, że każdy człowiek ma prawo ratować własne życie. W panice większość z nas będzie ratować siebie lub swoje dzieci. Uznajemy to za naturalne.

Z drugiej strony można powiedzieć, że na twórcach robotów spoczywa pewna ogólna odpowiedzialność za ludzkość. Czy systemy kierujące ruchem nie powinny być projektowane tak, aby w razie problemów uratować jak najwięcej istnień?

DYSKRYMINACJA?

Z „trzeciej strony” tworzenie systemów nastawionych na wyższe dobro mogłoby dawać nieuzasadnione korzyści określonym osobom. Załóżmy, że auta są tak programowane, aby nigdy nie uderzyć w duży rodzinny samochód albo karetkę pogotowia. Załóżmy, że kierowca tego rodzinnego samochodu jest człowiekiem, który również popełnia błąd. Czy nie dojdzie wówczas do tego, że rozwiązanie „dla wyższego dobra” stanie się powodem czyjejś zguby? Czy nie dojdzie do dyskryminacji pewnych ludzi (np. motocyklistów), którzy z zasady nie jeżdżą z całą rodziną?

MORALNE SZCZĘŚCIE?

Oczywiście Patrick Lin, jako specjalista od etyki, przemyślał ten problem o wiele głębiej. W swoim artykule w Wired odwołał się on do pojęcia „moralnego szczęścia”.

Czasem bywa tak, że podejmujemy jakąś decyzję nawet ze złymi intencjami, ale ostateczny jej efekt okazuje się chwalebny. Bywa też odwrotnie. Nierzadko działamy po omacku, a od przypadku zależy późniejsza ocena słuszności naszej decyzji. Wychodzi na to, że czasem przypadek ma większe znaczenie niż dbanie o wyższe dobro.

Roboty mogłyby w pewnych sytuacjach bazować na przypadku, tzn. mogłyby losować rozwiązanie problemu. Takie podejście ma jedno moralne usprawiedliwienie: kierowcy-ludzie też często bazują na przypadku, więc roboty działałaby podobnie. Szybko jednak

pojawi się wątpliwość. Czy roboty naprawdę powinny naśladować ludzi? Czy nie możemy zrobić czegoś lepszego od nas samych?

Takie rozważania można ciągnąć bez końca. Nawet jeśli nie znajdziemy jasnej odpowiedzi to wcale nie oznacza, że te rozważania są bezsensowne. Nie chodzi w nich tylko o udzielenie ostatecznej odpowiedzi „tak” lub „nie”. Takie eksperymenty pozwalają nam wstrzymać się przed wydawaniem szybkich, bezrefleksyjnych osądów, które z reguły do niczego dobrego nie prowadzą.

Autor: Marcin Maj

Źródło: [Dziennik Internautów](#)