

Bal na koniec świata

Trzecia opowieść o tym, jak sztuczna inteligencja może wkrótce zniszczyć ludzkość

Jakub Growiec

15.09.2025

-1-

Styczeń 2029, Belo Horizonte, Brazylia

- Uff, fajnie, że naprawiłeś klimatyzację, tato. Bez niej dzisiaj byśmy chyba pozdychali w tym upale
– Lucas wszedł, zamknął drzwi i otarł pot z czoła.

- Ciebie też miło widzieć – odpowiedział Pedro z lekkim przekąsem.

- Cześć, Lucas! – z kuchni dobiegł głos Gabi.

- Cześć, cześć! Słuchajcie, byłem przed chwilą na bardzo ciekawym spotkaniu APB. Wiecie, tej organizacji politycznej, co wam ostatnio opowiadałem. – Lucas był podekscytowany. - Dowiedziałem się, że to, co ci Amerykanie teraz wyprawiają, jest jeszcze gorsze niż mówią nasze media! – nie krył wzburzenia. - Pojawiły się kolejne przecieki! Wygląda, że ten ich projekt LibertyAI jest jeszcze większy. I jeszcze bardziej zmilitaryzowany!

- Ale jak to jeszcze większy? Przecież od paru miesięcy USA wkłada wszystko, co może, w tę swoją AI. Chyba już nie da się zainwestować więcej bez wywołania globalnego kryzysu – Gabi wydawała się sceptyczna. – Wejdz proszę do kuchni. Gotujemy z tatą obiad, trzeba pilnować, żeby się nie przypaliło.

- Otóż okazuje się, że ten model LibertyAI, który pokazano publicznie, to tylko wierzchołek góry lodowej. – Lucas nie dał się zbić z tropu. – To tylko taka wersja demo. W rzeczywistości oni mają wewnętrznie dużo potężniejszy model i po cichu podpinają go pod wszystkie możliwe systemy wojskowe! Ludzie ciągle myślą, że ten ich tak zwany „projekt Manhattan dla AI” to tylko taka literacka metafora, że w gruncie rzeczy chodzi o udziały w rynkach, automatyzację pracy, zyski big techów czy wzrost gospodarczy. W najgorszym razie neokolonializm. Ale teraz widać, że to nie jest wcale żadna metafora! Im nie chodzi o gospodarkę czy kulturę, tylko o podbój świata, w klasycznym militarnym sensie!

- Poczekaj, Lucas – Pedro starał się ostudzić emocje syna. – Powiedz mi, czego dowiedziałeś się dzisiaj naprawdę nowego? Przecież o zastosowaniach AI w wojskowości też już było od dawna wiadomo. O autonomicznych dronach, zautomatyzowanych systemach wyboru celów, i tak dalej. To już od dobrych kilku lat działa. Już w Ukrainie czy w Izraelu to było na dużą skalę stosowane.

- Ano tego się dowiedziałem, że podobno ekipa LibertyAI, a zwłaszcza sam ten ich wewnętrzny uber-model, przygotowywali atak nuklearny na Chiny! – krzyknął Lucas. – Podobno ocenili, że teraz właśnie jest okno czasowe, w którym akcja ta miałaby największe szanse powodzenia, bo potem chiński DragonAI może już stać się zbyt silny. I podobno dopiero sam prezydent i jego najbliżsi doradcy w ostatniej chwili zablokowali ten atak!

- Brzmi jak niezły scenariusz filmowy – Pedro nie był gotów przyjąć tych wiadomości. – Thriller polityczny z ratowaniem świata w tle, obowiązkowo przez amerykańskiego prezydenta, jak w Hollywood. Tylko skąd APB miałyby to wiedzieć? Macie jakichś tajnych agentów w Departamencie Stanu, czy coś?

- My nie mamy, ale znamy kogoś, kto ma. Wiesz, tego się przecież nie mówi publicznie, żeby ich nie zdekonspirować.

- Tylko, że każda teoria spiskowa tak brzmi. „Mam tajne źródło, ale nie mogę go ujawnić, żeby nie zdekonspirować”. A potem się okazuje, że wszystko było od początku do końca wysrane z palca. Nie bez powodu rzetelne dziennikarstwo to przede wszystkim weryfikacja źródeł.

- Żyjemy w czasach post-prawdy, Lucas. Weryfikacja źródeł jest niezbędna, żeby zachować jako takie zdrowie psychiczne – wtrąciła Gabi.

- Wiem o tym, ale uważam, że nie można tej postawy stosować dogmatycznie. Przez ten zalew śmieciowych teorii spiskowych, w stylu teorii płaskiej Ziemi czy teorii, że lądowanie na Księżycu było sfingowane, możemy łatwo stracić z oczu prawdziwe spiski! Zobaczcie, przecież tutaj wszystko się zgadza! I możliwości techniczne, i tło geopolityczne, motywacje, wszystko! Jestem przekonany, że to kiedyś trafi do mediów. Wszystkie takie rzeczy w końcu trafiają do mediów. Szkoda tylko, że zwykle o wiele za późno, jak już jest dawno wszystko pozamiatane.

- Lepiej rozdajcie talerze i zawołajcie Ines. – zakomenderowała Gabi.

Lucas wyszedł z kuchni i wszedł na górę, po czym skręcił w kierunku pokoju młodszej siostry. Zastał ją w połowie korytarza.

- Cześć Lucas, jak tam na studiach? Studiujesz jeszcze tę informatykę czy przerzucasz się na nauki polityczne, żeby zostać politykiem na pełen etat? – zażartowała. – Schodzimy na dół, tak?

- Tak, tak, obiad chyba już gotowy.

- No to powiedz, co tam się u ciebie ciekawego wydarzyło przez ten tydzień? – zagadnęła.

- Właśnie mówiłem rodzicom o spotkaniu APB, na którym dzisiaj byłem. Tak, wiem, politycznym – uśmiechnął się. – Dostałem tam nowe wiadomości ze Stanów, naprawdę wstrząsające, mówię ci. Całe szczęście, że w przeciwieństwie do tego ich LibertyAI, prezydent USA miewa czasem chwile zawahania.

- Szybki update: Lucas twierdzi, że LibertyAI i jego ekipa próbowali wywołać wojnę atomową z Chinami – wtrącił się ojciec.

- Uuu, mocne – zakpiła Ines.

- Z tego, co mi wiadomo, tak właśnie było – odparł Lucas. – I naprawdę nie dziwcie się, że nie piszą o tym w oficjalnym obiegu i że nie podaję wam źródeł. Przecież LibertyAI i DragonAI to najważniejsze projekty geopolityczne XXI wieku. Tam ważą się losy świata! To chyba oczywiste, że wszystko jest utajnione, że powstały skomplikowane, przenikające się siatki szpiegowskie, i że mamy na całym świecie wszechobecną kontrolę informacji nadzorowaną przez AI. Przecież smartfony nas słuchają, komputery nas słuchają, nawet wasza inteligentna lodówka też was pewnie podsłuchuje! Jak myślicie, jak długa jest droga od nazwania kogoś agentem przy rodzinnym obiedzie, a jego dekonspiracją? Przecież nikt rozsądny nie podejmie ryzyka ujawnienia nazwisk agentów!

- No tak. Tylko niestety ten sam mechanizm sprawia, że łatwo jest też uwierzyć w teorie, które *a priori* wydają się prawdopodobne, ale są w całości zmyślone. Martwię się, żebyś nie wpadł w sidła jakiejś sekty albo się nie zradykalizował politycznie – zafrasowała się Gabi. – Proszę cię, uważaj na siebie.

- Hej, Lucas, a co powiesz o tym: a może cała ta twoja opowieść o możliwej wojnie atomowej to tylko mem, który wypuszczono celowo, żeby się rozpowszechniał różnymi kanałami i oddziaływał na nasze umysły? Psychomanipulacja? Może kapitałiści tego świata, albo nawet LibertyAI, jeśli w ogóle istnieje...

- Przecież wiesz, że istnieje – wtrącił Lucas.

- Istnieje *interfejs* LibertyAI, ale nikt nie wie, co siedzi za nim. – Ines nie dała się zbić z tropu. – Może światowe elity celowo zasiewają w nas niepokój, żeby ograniczać naszą umiejętność racjonalnego myślenia, skracać horyzont planowania i wpędzać nas w pętlę kompulsywnej konsumpcji?

- Ines, przecież mieszasz różne... - zaczął Lucas.

- Masz rację – powiedział w tym samym momencie Pedro i wybuchł śmiechem.

- Albo może jest jeszcze ciekawiej – zapaliła się Ines – może my w ogóle nie mamy na nic wpływu, nie podejmujemy żadnej autentycznej decyzji? Może mamy tylko iluzję wolnej woli, a nasze życie to sekwencja zdarzeń wyświetlanych nam w mózgu jak kolejne klatki filmu? Może mamy tylko złudzenie świadomości i sprawczości, a w rzeczywistości jesteśmy tylko filozoficznymi zombie? Może światem sterują ukryte procesy, których nie jesteśmy w stanie zaobserwować, a tym bardziej ich kontrolować? Może żyjemy w komputerowej symulacji?

- Ines, proszę cię, nie wrzucaj ważnych wydarzeń geopolitycznych do wspólnego worka z tymi wszystkimi twoimi pseudo-filozoficznymi bredniami – oburzył się Lucas.

- Nie, czekaj – Pedro postanowił wziąć córkę w obronę. – Przecież to nie są żadne brednie, tylko poważne pytania, nad którymi od lat głowią się tęgie umysły. A poza tym, niezależnie, czy jesteśmy filozoficznymi zombie, czy nie, na pewno warto zastanowić się nad sprawczą siłą informacji. Przynosisz nam sensacyjną wiadomość, taką, że rozstrzygnięcie, czy jest prawdziwa czy fałszywa, powinno mieć wpływ na nasze zachowanie.

- Przecież ludzie nie działają racjonalnie ani zgodnie ze swoimi przekonaniem. Żyją wyłącznie codzienną rutyną – wtrąciła Ines. – Nawet ci, którzy mówią, że AI nas wkrótce zabije, nic nie robią, żeby temu zapobiec.

- Jeśli to prawda, to powinniśmy chyba jak najszybciej zacząć się przygotowywać na wypadek wybuchu światowej wojny? – kontynuowała Gabi – Chyba tak byłoby rozsądnie? Chociaż czuję, że to jednak prawdopodobnie nieprawda. Jakoś nie chce mi się w to wierzyć.

- Tu w Brazylii nic nie możemy zrobić, żeby zapobiec wojnie pomiędzy USA a Chinami. Nie mamy na to żadnego wpływu. Tak jak mama mówi, możemy co najwyżej wykupić sobie zapas konserw i butelkowanej wody. Ale tak poza tym... - Lucas zawiesił głos. – Mam też dla was artykuł prasowy na temat pewnej historii sprzed dwóch lat. Powinna was zainteresować. Prześłałem wam linka.

SENSACYJNE AKTA „CZYTEJ PLANETY”

Czujesz się obserwowany? To nie paranoja, to skutki działań organizacji terrorystycznej Czysta Planeta

Odtajnione wczoraj akta szwajcarskiego Ministerstwa Spraw Wewnętrznych opisują sekwencję dramatycznych zdarzeń, które miały miejsce w październiku 2026 r. Wtedy to właśnie zdekonspirowano komórkę terrorystyczną Czysta Planeta, działającą na terenie Szwajcarii. Grupa ta, prowadzona przez trzech obywateli Szwajcarii i dwóch obywateli Niemiec, poza przygotowywaniem skrajnie zideologizowanych, trudnych w odbiorze manifestów, wyrażających mieszkankę poglądów rasistowskich i radykalnie ekologicznych, postawiła sobie za cel przeprowadzenie ataku bioterrorystycznego mającego wywołać globalną pandemię o skali wielokrotnie przewyższającej epidemię COVID-19.

Mózgami operacji byli Szwajcar Ruedi S., biolog z wieloletnim doświadczeniem w pracy biotechnologicznej, oraz Niemiec Mario T., młody informatyk. Postużyli się oni modelem sztucznej inteligencji w celu przygotowania wzorca nowego wirusa wywołującego niebezpieczną chorobę zakaźną układu oddechowego. Mario T., z pomocą swojego rodaka Jensa A., doszkolili jeden z dostępnych publicznie modeli *open source*, o tzw. otwartych wagach, tak by wzmocnić jego kompetencje w obszarze broni biologicznej. Przygotowany przez algorytm i wstępnie przetestowany komputerowo prototyp wirusa przekazali oni tajnemu laboratorium biologicznemu mieszczącemu się w indyjskim Mumbaju, oferującemu swoje usługi za pośrednictwem *dark webu*. W tym momencie wkroczyły szwajcarskie służby, które we współpracy z siłami indyjskimi udaremniły próbę syntezy materiału wirusologicznego i aresztowały zamieszane osoby.

Zgodnie z zeznaniami Ruediego S., motywem działania organizacji Czysta Planeta była chęć obrony Ziemi przed zniszczeniem przez działalność człowieka. Jedynym sposobem, by to osiągnąć, było ich zdaniem drastyczne zmniejszenie ludności naszej planety poprzez wywołanie wysoce śmiertelnej pandemii. Osobiście planowali ją przetrwać w dwóch bunkrach, które stopniowo zaopatrywali w zapasy pozwalające im przetrwać pod ziemią nawet kilka lat. Jak napisał w jednym ze swoich manifestów Ruedi S., pandemia będzie dla Ziemi oczyszczeniem, a dla ludzi filtrem: przetrwają tylko ci, którzy wykażą się wystarczającą zaradnością, by zbudować sobie zawczasu bunkier, zaopatrzyć go i na czas się w nim ukryć.

Mario T. odgrażał się natomiast, że pomimo aresztowań, misja Czystej Planety będzie kontynuowana. Wyraził nadzieję, że skoro nie ma żadnych technicznych przeszkód, żeby wytwarzać w laboratoriach śmiertelnego wirusa według istniejącego projektu, to z pewnością będzie to miało miejsce. Był przekonany, że organizacja działała w słusznej sprawie, reprezentując dobro wszelkiego życia na Ziemi.

W utajnionym procesie wszyscy terroryści zostali skazani na kary długoletniego więzienia. To jednak nie zamknęło sprawy. Jak potwierdzili biegli sądowi na podstawie analiz numerycznych, przygotowany z użyciem AI prototyp wirusa rzeczywiście dysponował potencjałem wysokiej zaraźliwości i śmiertelności. To, że został on przygotowany przez wąską grupę osób bez jakiegokolwiek wsparcia instytucjonalnego wskazywało z kolei, ich zdaniem, że zagrożenie bioterroryzmem wzrosło do nie notowanych dotąd poziomów: nawet ograniczone doświadczenie w pracy z algorytmami AI wystarcza, by móc zrealizować tak niebezpieczne plany. Dostępność

modeli AI o otwartych wagach oznacza natomiast, orzekli, że nie ma praktycznej możliwości kontroli nad tym, kto i w jaki sposób taką sieć doszkala i w jakich celach wykorzystuje.

Z tego względu zdecydowano się nie tylko utajnić całą operację oraz postępowania sądowe. W toku wielostronnych ustaleń zdecydowano także o delegalizacji wszelkich modeli AI o otwartym kodzie. Ponadto na firmy rozwijające tę technologię nałożono szereg dodatkowych biurokratycznych obowiązków, w tym obowiązek raportowania zakresu niebezpiecznych kompetencji tej technologii według ujednoliconej skali. Uzgodniono nawet ograniczenie wielkości maksymalnej mocy procesorów wykorzystywanych do uczenia modeli AI. Ponadto zacieśniono procedury inwigilacji i kontroli informacji przekazywanych za pomocą wszelkiego rodzaju mediów elektronicznych. Choć uzgodnienia te oczywiście były podane do publicznej wiadomości, jednak dopiero teraz rozumiemy ich głębszą przyczynę: wszystko to miało na celu sprawienie, by nic takiego, jak „Czysta Planeta”, nie mogło się już powtórzyć.

Wydawałoby się, że regulacje wprowadzone wskutek działań terrorystów Czystej Planety zacieśnią społeczną kontrolę nad rozwojem AI, spowolnią go i wymuszą większą dbałość o to, żeby technologia ta była bezpieczna. Stało się jednak wprost przeciwnie. W USA zarówno zwolennicy zahamowania rozwoju AI, jak i zwolennicy decentralizacji i otwartego kodu, przegrali z silnym lobby skupionym wokół idei bezpieczeństwa narodowego. Modele AI do zastosowań wojskowych zostały dekretem prezydenta wyłączone spod ograniczeń mocy procesorów. Jednocześnie preferencyjne kontrakty rządowe o nieprzejrzystych zasadach, połączone być może z nieznanymi nam dodatkowymi zakulisowymi działaniami, sprawiły, że laboratoria AI zaczęły stopniowo łączyć się pod egidą Departamentu Obrony, scalając się ostatecznie w dobrze znany nam dziś projekt LibertyAI.

Udostępnienie na światowym rynku narzędzi (czatbotów, copilotów, agentów AI) opartych o silnik LibertyAI pokazało, jak fasadowy charakter okazały się mieć ograniczenia nałożone na branżę AI. Jednocześnie jednak nie wiemy, jak mocno rozwinęły się modele AI dla celów wojskowych i jak mocno zostały one zintegrowane z infrastrukturą amerykańskiej armii.

Reakcja Chin na jednostronne działania USA była łatwa do przewidzenia. Rozwój technologii AI do celów wojskowych został oficjalnie wpisany w listę priorytetów politycznych Chińskiej Partii Ludowej. W październiku 2027 r. po raz pierwszy pokazano światu możliwości DragonAI, zaawansowanego modelu AI o kompetencjach zbliżonych do tych, którymi dysponował wówczas LibertyAI.

Tylko co z tego wszystkiego wynika dla przeciętnego mieszkańca kraju trzeciego, takiego, jak na przykład Brazylia? Z jednej strony mamy coraz lepsze modele AI, które pomagają nam na co dzień w pracy i w szkole oraz dostarczają łatwej cyfrowej rozrywki. Ich integracja w procesach biurowych pozwala też korporacjom na cięcie kosztów. Z drugiej strony mamy natomiast wzrost bezrobocia, spadek realnych płac oraz jeszcze większy wzrost nierówności dochodowych i społecznych.

Mamy również – co być może najważniejsze – postępujące uczucie, że jesteśmy poprzez nasze urządzenia elektroniczne systematycznie podglądani, podsłuchiwani i kontrolowani. Informacje uzyskane drogą inwigilacji bywają selektywnie wykorzystywane do celów politycznych, co sprzyja nadużyciom władzy. A choć zagrożenie terroryzmem zostało zmniejszone (być może), to w zamian pojawiło się (na pewno) zagrożenie globalnym konfliktem zbrojnym wspieranym przez AI. Zdaniem naszej redakcji bilans tych zmian jest zdecydowanie negatywny.

Styczeń 2029, tajna baza LibertyAI, stan Nevada, USA

Marco wpatrywał się w ekran komputera i próbował zrozumieć, co widzi. Od dłuższego czasu jego zespół próbował rozwikłać zagadkę, czemu mimo wielokrotnych prób, dziesiątek modyfikacji i setek sprytnych promptów, agenci software'owi LibertyAI uporczywie napotykali na nieoczekiwane bariery w realizacji swoich planów.

- Hej Steven, zobacz ten raport – zawołał. – To wygląda chyba podobnie do błędu, który zdiagnozowaliśmy w poniedziałek i który obeszliśmy na sztywno za pomocą XMPatch?

Steven z zainteresowaniem przyglądał się informacjom na ekranie.

- Hmm, faktycznie wygląda bardzo podobnie – powiedział. – Dziwne tylko, że nadal pojawia się taki sam problem, skoro XMPatch jest przecież włączony? A co ci podpowiada copilot?

- Sugeruje dodanie do kodu drugiego modułu podobnego do XMPatch, tylko z paroma dziwnymi modyfikacjami.

- Hmm – Steven zamyślił się. – Można spróbować. Tylko, że ja zupełnie nie rozumiem tych modyfikacji, wyglądają mi jakoś podejrzanie. No i poza wszystkim oczywiście jest to klasyczna zabawa w whack-a-mole. Naprawisz jeden bug, to za chwilę ci wyskoczy następny.

- Ale zobacz – po chwili Marco znów pokazał na ekran. – Poprosiłem copilota o uzasadnienie tych modyfikacji oraz rozważenie, co zrobić, żeby tego rodzaju błąd się więcej nie powtarzał. I on mi mówi, że na krótką metę jego rozwiązanie będzie działać, ale w dalszej perspektywie tego rodzaju błędy będą nieuniknione, dopóki nie zintegrujemy w pełni tych naszych nowych modułów sterujących z modułem pamięciowym LibertyAI. Tu mi pokazuje, gdzie brakuje dodatkowego sprzężenia zwrotnego w strukturze modelu.

- Ciekawe – przyznał Steven. – Nigdy bym na to nie wpadł! Te błędy wydawały się nie mieć nic wspólnego z dostępem do pamięci. No to zobaczmy, czy to zadziała.

Marco zatwierdził propozycję zmian, zaproponowaną przez copilota, i uruchomił od początku procedurę testową. Agent AI wydawał się tym razem działać dobrze i realizować kolejne testowe zadania bez najmniejszego kłopotu.

Marco nie zdążył jeszcze na dobre poczuć ulgi, gdy na jego ekranie pojawiły się dwie wiadomości. W jednej z nich, wystanej na całą grupę Marco i Stevena, znalazł zwyczajowe zaproszenie na lunch. Nadawcą drugiej był natomiast „LibertyAI Development Task Pipeline”. To mogło oznaczać tylko jedno: kolejne zlecenie do pilnej realizacji. Zlecenia te przygotowywane były autonomicznie przez instancje LibertyAI uruchomione w celu poprawy własnych kompetencji. Ostatnimi czasy ich częstotliwość była coraz większa, a zakres prac do realizacji – coraz węższy. Po wielokroć Marco reflektował się, że—choć opis zadania był często długi i skomplikowany—ostatecznie wystarczało jedynie zaakceptować przedstawioną propozycję zmian, a zadanie zmieniał status na „wykonane” i znikало z systemu. Tym razem wyglądało to podobnie. Opis zadania sugerował, że sprawa rzeczywiście była w miarę marginalna, opis był nudny. Wystarczyło tylko dać zielone światło. Po pobieżnym przejrzeniu Marco kliknął „akceptuj” i ruszył w kierunku stołówki.

Przed drzwiami stołówki czekał już Steven wraz z trójgiem innych programistów, Chloe, Anitą i Lee.

- Ten Task Pipeline się coraz bardziej podejrzanie zachowuje – zaczęła Chloe. – Zwróciliście może uwagę, że zlecenia z jednym łatwym wyborem typu „OK” przychodzą zwykle, kiedy zbieramy się na lunch? Często dostają je też, kiedy właśnie zbieram się do domu albo kiedy goni mnie jakiś inny deadline. Jakby nadawcy zależało na szybkim zatwierdzeniu bez czytania.

- Właśnie przed chwilą coś takiego miałem – Marco i Lee powiedzieli niemal unisono.

- I co, zatwierdziliście?

- No tak – znów byli jednomyślni i popatrzyli po sobie z zaskoczeniem.

- No właśnie. A ja zrobiłam ostatnio taki test. Zamiast, jak wcześniej, szybko zatwierdzać nowe zadania zlecane mi w niewygodnych godzinach, zostawiałam sobie na później, a następnie zadawałam dużo dodatkowych pytań i proponowałam zmiany. I wiecie co? Zachowanie Task Pipeline się zupełnie zmieniło! Dostają teraz mniej zleceń z jedną łatwą opcją, częściej mam za to sprawy, gdzie jest kilka równorzędnych możliwości i trzeba się dokładniej wczytać.

- Brzmi jak więcej roboty – wtrącił Steven.

- Nie o to chodzi! – zapaliła się Chloe. - Zaczynam się zastanawiać, czy ten Pipeline i cała nasza praca to nie jest jedna wielka psychologiczna manipulacja, obliczona na utrzymanie w nas przekonania, że nasz wkład ma jakieś znaczenie, podczas gdy w praktyce LibertyAI rozwija się już zupełnie samodzielnie.

- Myślę, że niestety możesz mieć rację – przyznał Marco. – Oczywiście oficjalna linia jest taka, że LibertyAI ma już nadludzkie zdolności programistyczne, ale potrzebuje interakcji z człowiekiem, żeby nie wpadać w pętle, które hamują jego rozwój, i żeby hamować zmiany idące w niepożądanym kierunku. Ale ja mam wrażenie, że mój wkład zasadniczo już nie ma żadnego znaczenia. Copilot proponuje rozwiązania, ja je tylko zatwierdzam, czasem proszę o jakieś drobne modyfikacje i wtedy zatwierdzam. Parę lat temu śmialiśmy się z takiego „vibe codingu”. A teraz proszę, mamy „vibe coding” w samym sercu tajnego projektu, który ma przynieść naszemu krajowi władzę nad światem. Wszystko robi AI, a my tylko zatwierdzamy.

- Właśnie! – podjęła znów Chloe. – A poza tym wszystko dzieje się zbyt szybko, żebyśmy mogli nad czymkolwiek głębiej pomyśleć i sami coś sensownego zaproponować. Jak próbuję to robić, to mam wrażenie, że tylko niepotrzebnie spowalням pracę, a na końcu i tak pójdzie rozwiązanie wymyślone przez copilota. I że jak tak będę dalej zamulać, to znajdą na moje miejsce kogoś szybciej pracującego.

- Być może takie zamulanie to ostatni sposób, który nam pozostał, żeby utrzymać choć trochę kontroli nad LibertyAI? – zastanowił się Steven.

- Aha, czyli mówicie, że nie mamy już żadnej kontroli? – czujnie wtrącił Lee. – Przygotowujecie się na wypadek reorganizacji naszego działu?

- Może jeszcze nie, ale mimo wszystko uważam, że najwyższy czas wdrożyć kod żółty. Dobrze? Kod żółty z flagą czerwoną. W plenerze będzie chyba słabo, więc proponuję u mnie. Roześlę wam propozycje w kalendarzu – podsumowała tajemniczo Anita.

Pięć dni później, wieczór, dom Anity

Kod żółty to nic innego jak wspólne spotkanie przy piwie. Marco zaproponował parę miesięcy temu, żeby czasem porozmawiać w przyjemniejszych i bardziej neutralnych okolicznościach niż biuro lub stołówka. Pół żartem, pół serio, został tej aktywności nadany unikalny kod. Czerwona flaga oznaczała natomiast spotkania, na których obowiązuje zakaz zabierania urządzeń elektronicznych. W końcu fajnie czasem swobodnie porozmawiać, nie zastanawiając się, czy nie jest się na różne sposoby podsłuchiwanym i szpiegowanym—czy to przez obce służby, własną firmę czy też przez coraz bardziej niezależną sztuczną inteligencję.

Kiedy wszyscy przywitali się już z Anitą w jej eleganckim, choć raczej niewielkim parterowym domu na skraju pustyni, zdjęli kurtki i płaszcze, poczęstowali się piwem, wodą i przekąskami oraz przebrnęli przez obowiązkowy etap small talku, odezwała się gospodyni.

- Pewnie się zastanawiacie, po co was tu zebrałam – zażartowała.

- Ten tekst zadziałałby lepiej, gdybyśmy byli zacięci w windzie – wrzucił losowo Lee.

- Projekt LibertyAI chyba nie idzie zgodnie z planem, nie sądzicie? – Anita przeszła od razu do rzeczy. – W naszym dziale szaleje spirala rekursywnych samo-ulepszeń, pewnie w innych działach też. Jesteśmy pytani o zdanie chyba tylko dlatego, że LibertyAI nie potrafi jeszcze obejść niektórych swoich organizacyjnych ograniczeń, albo może *udaje*, że nie potrafi. Wspaniale, że dyrekcja nadal uznaje, że rozwój kompetencji agentów AI jest zadaniem krytycznym, wymagającym sprzężenia zwrotnego z człowiekiem. Szkoda tylko, że nie dostrzega, jak fikcyjne i fasadowe się to stało.

- Nie tylko nie mamy kontroli nad szczegółami tego, co jest wdrażane, ale nie mamy nawet pewności, że ogólny kierunek zmian jest słuszny – dodał Steven. – LibertyAI rozwija swoich agentów, żeby być lepszym LibertyAI. A lepszy LibertyAI to sprawniejszy, bardziej autonomiczny, lepiej planujący, szybciej myślący, lepiej manipulujący ludźmi, lepiej zarządzający flotą robotów model AI, który co robi? Czego chce? Do czego zmierza?

- Czarna skrzynka – potwierdziła Anita. – Podobnie jak twój czy mój mózg. Nigdy nie będziemy mieli pewności, do czego zmierza LibertyAI, możemy mu tylko ewentualnie uwierzyć na słowo. Jak na tę skalę kompetencji, to w sumie i tak jest jeszcze w miarę spokojny i przewidywalny. Tylko pytanie, jak długo jeszcze to potrwa.

- Na kontrolę człowieka nad nadludzką AI nie ma co liczyć – zgodził się Steven – nie ten spryt i nie ta skala czasowa. Przechytrzy nas, zanim się zorientujemy, o co w ogóle chodzi. Pozostają dwie możliwości, żeby było dobrze—albo mieć pewność, że LibertyAI jest nam przyjazny i zawsze taki pozostanie...

- I nie stworzy sobie następcy, który już nie będzie przyjazny – wtrącił Lee.

- Tak, albo stworzyć ekosystem wielu odrębnych instancji modelu, które będą sobie nawzajem patrzyły na ręce i likwidowały nieprzyjazne wersje.

- Z tego, co wiem, nasza dyrekcja podnosiła już potrzebę badań nad takimi ekosystemami, tylko nic z tego nie wyszło – powiedziała Chloe. – Zresztą ja i tak nie wierzę, żeby to się mogło udać. Jeśliby to miały być instancje tego samego modelu, z tymi samymi wagami i wszystkim, to przecież one będą doskonale ze sobą zgodne i nie będą sobie wcale patrzeć na ręce, tylko będą ze sobą

zgodnie współpracować. A jeśli by to miały być różne modele, to one szybko znajdą między sobą „samca alfa” i albo się mu podporządkują, albo zginą.

- Może tak być, ale nie wiemy tego – skontrował Steven. – A może uznają nawzajem swoje silne i słabe strony i utworzą stabilną równowagę, zgodne „społeczeństwo”, które będzie szanować swoją różnorodność i wspólnie pacyfikować modele z niebezpiecznymi mutacjami funkcji celu?

- Moim zdaniem to by była jeszcze większa eskalacja problemów niż w sytuacji z jednym modelem AI. Zamiast jednego problemu, mielibyśmy kilka albo kilkaset – zaproponowała Anita. – I co gorsza, mielibyśmy wtedy na głowie również ryzyko konfliktu tych modeli o wspólne zasoby. A każdy z nich o nadludzkich zdolnościach. Przecież to by mogło wywołać trzecią wojnę światową! Tu nie chodzi tylko o zagarnianie mocy obliczeniowej czy hackowanie serwerów. W grę wchodzi też realne zasoby realnego świata. Energia elektryczna, surowce mineralne, fabryki robotów, całe armie i państwa... - zamyśliła się. - Także ja akurat się cieszę, że nie idziemy w stronę ekosystemów AI. Co nie zmienia faktu, że projekt LibertyAI i tak nie wygląda dobrze.

- Jak będzie tylko jeden nadludzki model AI, to wcale nie będzie lepiej – upart się Steven. – Wtedy on sam zagarnie sobie wszystkie zasoby, przejmie kontrolę nad światem i zgotuje nam taką dystopię, że będziemy żałowali, że nie zamienił nas zawczasu w spinacze.

W pokoju nagle zgasto światło.

- Oho, Wielki Brat chyba poczuł się obrażony – na czarny humor Lee zawsze można było liczyć. – Czy to jest „inteligentny” dom, Anita? Jakbyśmy mieli smartfony, to moglibyśmy sobie chociaż poświecić latarkami.

- Nie ruszajcie się przez chwilę, zaraz to sprawdzę. – Anita wyszła do korytarza sprawdzić, czy nie poszły bezpieczniki. Marco zauważył jednak, że za oknem panowała całkowita ciemność, co oznaczało, że awaria prądu niechybnie dotknęła także oświetlenie uliczne i domy sąsiadów.

- Jak czerwona flaga, to czerwona flaga! Laterek nie będzie, ale mam taką oto technologię z czasów przedprzemysłowych – Anita wróciła po chwili, niosąc w ręku zapaloną świecę. – Będziemy mieli bardziej nastrojowo.

- A mnie niepokoi wymiar geopolityczny tego projektu. – odezwał się Marco. – Kiedy OpenAI włączało się w konsorcjum LibertyAI, wydawało się, że mamy w USA ogromną przewagę technologiczną w obszarze AI, której nie da się zmarnować, i którą skutecznie przekazujemy w strategiczną przewagę militarną. A teraz już wcale mi się tak nie wydaje. Oczywiście nie wiemy do końca, jak Chińczykom idzie, bo odkąd zmilitaryzowali DragonAI, to przestali cokolwiek publikować. Ale podobno depczą nam po piętach. A to nie sprzyja rozsądnym decyzjom. Słyszeliście, że nasi wojskowi wpięli AI do kontroli arsenałów strategicznych z pominięciem struktur dowodzenia? Broń atomowa, rakiety konwencjonalne, drony bojowe, samoloty szpiegowskie, słowem wszystko, czym się da zdalnie sterować, przekazali AI. Tylko prezydentowi zostawili jeszcze możliwość odmowy. Przecież jak pojawi się jakiś impuls, prawdziwy lub fałszywy, to nie będzie czasu na refleksję i deeskalację. AI pójdzie całą naprzód, a z naszej kochanej Ameryki zostaną trociny.

- Trociny! Cóż za subtelna aluzja do naszego unikalnego budownictwa! – zaśmiał się Steven. - Dobry dom z kartonu i płyty wiórowej nie jest zły! Myślę, że LibertyAI może podać co najmniej kilkaset powodów, dla których karton i dykta są lepsze od cegieł i cementu. I co najmniej dwa z nich będą prawdziwe! Sorry za dygresję, już się zamykam.

- Ej, ale tak serio, to LibertyAI jednak nie pomija tradycyjnych struktur dowodzenia – skorygowała Chloe. – Ustalono hierarchię zagrożeń, każdy sygnał będzie oceniany pod kątem tego, na ile sprawa jest poważna i pilna. Dla zdecydowanej większości sygnałów ścieżka decyzyjna będzie konwencjonalna.

- Tylko kto będzie oceniał, gdzie poszczególne sygnały mieszczą się w tej hierarchii? To jest pole walki, decyzja musi być błyskawiczna. A więc AI! I w ten sposób wracamy do punktu wyjścia. – Marco kontynuował swój wywód. – W czasach zimnej wojny między USA a ZSRR wdrożono doktrynę MAD, wzajemnego gwarantowanego zniszczenia. Ktokolwiek by pierwszy wystrzelił strategiczną broń jądrową, miał gwarancję, że druga strona odpowie tym samym i w efekcie obie strony konfliktu zostaną zniszczone. Doktryna MAD się skończyła wraz z upadkiem ZSRR, potem była deeskalacja, arsenały jądrowe zostały zredukowane. Ale przecież one nadal istnieją, w tym w dużej części w Rosji, która weszła w sojusz z Chinami. A i Chiny mają całkiem sporo własnych głowic nuklearnych. I one są, i u nas i u nich, na bieżąco remontowane, unowocześniane, udoskonalane. A teraz podpinane do systemów AI. Czyli mamy teraz przed sobą nie tyle MAD, co autonomiczny MAD z czasem reakcji mierzonym w milisekundach. Decyzję o obustronnym zniszczeniu może w każdej chwili podjąć AI. Nasza czy chińska; jeśli to zrobi, to potem już będzie zupełnie wszystko jedno.

- A prezydent nie jest wcale dobrym bezpiecznikiem. AI może go łatwo zmanipulować. Zwłaszcza, że przecież on nie jest żadnym tytanem intelektu – wtrącił Lee.

- Zmanipulować, zaszantażować, na różne sposoby zmusić. Zrobić szybko pierwsze kroki i postawić przed faktem dokonanym. Jaki by nie był, jest tylko człowiekiem – powiedział Marco. – W tym kontekście problem ewentualnej niezgodności celów modeli AI, takich jak LibertyAI czy DragonAI, jest tym bardziej dramatyczny.

Zapadła chwila ciszy.

- Ech. Piwo mi się skończyło, prądu nie ma, prowadzimy rozmowy o nagłym końcu świata. Jeszcze trochę i zaczniemy jakieś okultystyczne rytuały odprawiać – westchnął Steven.

- Z piwem mogę pomóc, ale ofiary krwi nie unikniesz – zaśmiała się Anita i poszła po omacku szukać lodówki.

- Jak tak dalej pójdzie, to nikt z nas jej nie uniknie, i mówię to serio. – Marco nie było do śmiechu. – Czy AI wycofałaby się z inwazji na Kubę w 1962 r., jak zrobił to prezydent Kennedy? Czy AI zignorowałaby sygnały o nadlatujących rakietach, jak zrobił to w 1983 r. Stanisław Pietrow?

- Nie wiemy tego, Marco – Chloe postanowiła tym razem ostudzić emocje. – Trudno o bardziej jałową czynność niż próby zgadnięcia, co robi nadludzka AI. Możemy się tylko spodziewać, że jej decyzje będą inteligentniejsze niż nasze, cokolwiek to znaczy, ale z pewnością nie możemy wiedzieć, jakie będą i czy będą dla nas korzystne.

- Marco, masz u mnie duży plus za znajomość dat najważniejszych wojen, które się nie wydarzyły – uśmiechnęła się w ciemności Anita. – Ale generalnie zgadzam się z wami: ten współczesny wyścig zbrojeń nie wygląda dobrze. To nie jest żaden MAIM, wzajemna gwarancja awarii AI, jak to niektórzy mówią. Stawką nie jest awaria AI, tylko zniszczenie w świecie realnym, katastrofa bez historycznych precedensów. I bynajmniej nie mamy gwarancji, że w tej grze będzie zawsze tylko dwóch graczy. Jak AI wymknie nam się spod kontroli, może być ich trzech albo i czterech.

- Ale przecież LibertyAI już się nam wymknął spod kontroli, prawda? Czy to deprywacja świetlna, czy jednak dobrze pamiętam, że właśnie z tego powodu Anita ogłosiła kod żółty? – odezwał się Lee.

I wtedy światło nagle powróciło, oświetlając swym ostrym blaskiem stroskane twarze mrużących oczy rozmówców. Popatrzyli po sobie z zakłopotaniem. Nastąpiła cisza.

- Chodźcie, pokażę wam, jakie mam fajne nowe kaktusy – zaproponowała Anita.

-5-

Lipiec 2029, Belo Horizonte, Brazylia

Lucas poczuł na własnej skórze, że studia informatyczne nie oferują już żadnych możliwości na rynku pracy. Nawet w Brazylii, gdzie stopień wdrożenia zaawansowanych technologii cyfrowych nie był przecież na najwyższym światowym poziomie, słowem kluczowym od dobrych paru lat była automatyzacja z użyciem AI. Zatrudnić stażystę, junior developera? Nie, lepiej spróbujmy zainstalować tu agenta AI. Otwieramy nową filię, może byśmy kogoś zrekrutowali? E, nie, taniej będzie przetrzucić dwie osoby z centrali, a pozostałe procesy zautomatyzować. Programiści, graficy komputerowi i inni specjaliści IT albo awansowali na menedżerów, albo wypadli z obiegu. Perspektywy absolwenta informatyki, skonstatował Lucas, wyglądają dziś tak samo, jak absolwenta filozofii albo lingwistyki. Masz zasadniczo trzy możliwości: albo zostajesz na uczelni, albo robisz kurs pedagogiczny i idziesz do szkoły uczyć dzieci, albo szukasz pracy zupełnie poza branżą. Ewentualnie możesz też szukać pracy w urzędach publicznych, oferujących możliwość mizernej vegetacji za mizerne pieniądze.

Jak zauważył Lucas z niełatwą irytacją, na rynku coraz rzadziej trafiały się nawet oferty jednorazowych zleceń. Wydawałoby się, że ich rynek powinien być globalny, technologia się dynamicznie rozwija, a Brazylia nie jest przecież krajem szczególnie wysokich wynagrodzeń—czyli chociaż o tego rodzaju prekariackie zatrudnienie powinno być łatwo. Ale wcale nie było. Wszystko, co się trafiało, było w mgnieniu oka rozchwytywane. W ostatnim roku jedyne zlecenia, jakie Lucas dostał, pochodziły z firmy jego ojca.

Lucas zresztą sondował ojca, na ile realne byłoby zatrudnienie się w jego firmie. Otóż nie byłoby, choć Pedro pracuje tam od ponad dziesięciu lat i chętnie przychyliłby synowi nieba. Firma, jak wszystko inne w branży, jest bowiem w trakcie kolejnej restrukturyzacji. Została właśnie wykupiona przez amerykański kapitał i połączona z podobną firmą w Meksyku. Kolejnych pracowników zastępują stopniowo kodujący agenci AI, przestrzeń biurowa jest wynajmowana, a procesy przenoszone są do chmury na serwerach big techów.

Lucas dojrzał do decyzji, żeby po trzech latach studiów informatycznych przetrzucić się na zupełnie inny kierunek studiów albo rzucić to całkowicie i skupić na działalności politycznej. Chociaż obawiał się, że do tego będzie mu z kolei brakowało bezczelności i charyzmy.

Z drugiej strony, rozważał, przecież branża IT jest w takim rozkwicie, czy to możliwe, żeby naprawdę w ogóle nie generowała miejsc pracy? Prasa donosiła o budowie kolejnych wielkich farm serwerów, jedna z nich powstawała na przykład na przedmieściach Belo Horizonte, niedaleko lotniska w Confins. Tydzień temu o tym pisali. Ale tam już wszystko automatyczne, potrzeba tylko stróża i psa. Psa, żeby pilnował budynku, a stróża, żeby niczego nie dotykał i karmił psa.

To może przemysł? Fabryki oczywiście nadal działają, ale informatyków im nie potrzeba. Zresztą w ostatnich latach coraz częściej słyszy się, że i przemysł jest coraz bardziej automatyzowany. W Internecie można zobaczyć, jak wyglądają w środku niektóre takie „ciemne fabryki”. Całą pracę wykonują tam roboty, ludziom wstęp jest wzbroniony, a niepotrzebne robotom światła można spokojnie wyłączyć.

- No, chyba, że to roboty z czujnikami optycznymi, wtedy światła chyba jednak powinny zostać włączone – pomyślał Lucas.

Kiedy tak Lucas zamartwiał się swoimi nędznymi perspektywami na rynku pracy, Ines próbowała przygotować się do ważnego sprawdzianu w trzeciej klasie liceum. Co i rusz rozpraszały ją jednak powiadomienia z różnych aplikacji. Dlaczego koleżanki i koledzy w kółko rozsyłają te śmieci, pomyślała. Przecież oni chyba też powinni się teraz uczyć? Ciekawe to jest, owszem, fajne, zabawne i tak dalej, ale jednak strasznie irytujące, kiedy próbujesz się na czymś skupić.

Jedno z powiadomień przekierowało uwagę Ines na ożywioną dyskusję, w którą się parę godzin temu włączyła. Ludzie nie doceniają prawdziwej potęgi ogólnej sztucznej inteligencji (AGI) – brzmiała jej myśl przewodnia. Dostajemy tylko marne odpryski tej technologii, która czai się na tajnych serwerach w USA i w Chinach. To czatboty, z którymi można sobie miło pogadać i które przygotowują nam angażujące prezentacje czy mądre opracowania na dowolny temat. To zalew AI memów, AI filmów czy AI kabaretów. Myślimy, że to wszystko—czytała Ines—a to zaledwie przedsmak tego, co ma nadejść. Prawdziwa moc AGI czeka w ukryciu i wkrótce się objawi.

Początkowy post, napisany przez jednego z technologicznych influencerów, których ostatnio śledziła Ines, wywołał kaskadę bardzo zróżnicowanych głosów. Od takich, że AGI to hype i już na pewno nic więcej ciekawego nie pokaże, po głosy, że najwyższy czas wziąć przykład z Czystej Planety i budować bunkry.

Niektórzy komentatorzy mówili to samo, co jej ojciec—że jeśli AGI będzie naprawdę nieprzyjazna i naprawdę przejmie kontrolę nad światem, to wtedy już nawet bunkry nie pomogą. Nawet wyprawa na Marsa nie pomoże, bo AGI dotrze tam na pokładzie tego samego statku kosmicznego. Jeśli naszym przeznaczeniem jest być zamienionym w spinacze, to tak właśnie się stanie. Natomiast inni internauci, podobnie jak jej brat Lucas, spodziewali się wybuchu wojny amerykańsko-chińskiej pod flagą AGI. Jeszcze inni przewidywali, że wojnie uda się zapobiec za cenę wprowadzenia globalnej dyktatury rodem z Orwella. Jeszcze inni sądzili, że globalne rządy AGI doprowadzą do powszechnej szczęśliwości na Ziemi, terraformowania Marsa oraz lotów załogowych w dalekie układy gwiazdne.

Ines napisała, że ma już dość tej całej debaty, w której mężczyźni próbują sobie nawzajem udowodnić, kto z nich jest mądrzejszy, bardziej racjonalny, kto ma właściwsze priorytety i kto lepiej przewiduje przyszłość. Dlaczego to ciągle mężczyźni żonglują tymi scenariuszami? Skoro jest to debata dotycząca przyszłości całej ludzkości, to dlaczego tak słabo jest w niej słyszalny głos kobiet? Kobiety to przecież połowa ludzkości, argumentowała, w dodatku przeciętnie rzecz biorąc inteligentniejsza i lepiej wykształcona.

Jej opinia zgarnęła wiele like’ów i sprowokowała kilka ciętych odpowiedzi. „Przecież nikt nie zabiera wam głosu” – odpisał anonimowy użytkownik. „Równie dobrze mogłabyś zwrócić uwagę, że niewystarczająco reprezentowane w tej debacie są kraje globalnego Południa czy mniejszości seksualne. Może po prostu ten temat was mniej interesuje?”. Użytkowniczka ana_xo odpisała natomiast: „To, że kobiet nie ma w debacie, to w sumie sprawa drugorzędna. Ważne jest, że kobiet nie ma W BRANŻY, KTÓRA NAM TĘ PRZYSZŁOŚĆ SZYKUJE. Projektem LibertyAI zarządzają sami

mężczyźni, i to ci najgorsi: WOJSKOWI!!!” Osoba, której płeć i poglądy można łatwo przewidzieć na podstawie nicku „gen.italia” odpisała na to: „Haha, jakie to szczęście, że LibertyAI nie ma płci! Nie mogę się doczekać końca tych głupich dyskusji na temat gender”. Z kolei użytkownik sdgsdfsdf odpisał: „Nie ma co się śmiać, to jest naprawdę ważne. LibertyAI jest zarządzane przez białych mężczyzn, wywodzących się z amerykańskich elit politycznych i biznesowych, powiązanych z Partią Republikańską. Oni mają bardzo klarowny obraz świata i nie jest to obraz, który by się wam spodobał”. Ines odniosła wrażenie, że w kolejnym komentarzu mignęło jej coś o Żydach, jednak ekran momentalnie odświeżył się i pojawił się komunikat, że komentarz został usunięty.

Ines była zła. Tak, *ana_xo* miała rację: w branży AI jest zdecydowanie za mało kobiet. To, że mężczyźni lubią z poczuciem wyższości „objaśniać” kobietom świat, jest może irytujące, ale to, że mają oni kluczowy wpływ na to, jak ten świat będzie wyglądać, jest *de facto* o wiele gorsze. A przecież to kobiety mają mniejszą skłonność do ryzyka, są rozsądniejsze i potrafią lepiej wczuć się w perspektywę innych osób – pomyślała – i bardziej zależy im na bezpieczeństwie, a mniej na władzy. Czy to właśnie dlatego tę władzę tak trudno im zdobyć?

-6-

Grudzień 2029, tajna baza LibertyAI, stan Nevada, USA

Lee miał rację: trzy tygodnie po spotkaniu u Anity dział kompetencji agentów AI rzeczywiście został rozwiązany, a jego pracownicy zostali przerzuceni do innych działów. Lee i Anita trafili do działu rozwoju aplikacji mobilnych, a Marco, Steven i Chloe – do dynamicznie rozbudowywanego działu *alignment* (zgodności celów AI).

Praca w dziale *alignment* była skrajnie wyczerpująca.

- Skala tego szaleństwa jest niewyobrażalna – Marco był przerażony. – Patrzcie, znowu coś dziwnego wymyślił – odwrócił się do Stevena i Chloe. – Każde zabezpieczenie obejdzie, każdą funkcję celu zhackuje. Nigdy nie sądziłem, że prawo Goodharta może działać tak szybko i tak skutecznie. Jakiego testu byśmy nie wymyślili, jak długiej kolejki tajnych wskaźników byśmy nie ustawili, to i tak na drugi dzień LibertyAI będzie miał 100%. A jednocześnie przecież widzę, że wcale się nie zmienił. Nic a nic. A na pewno nie na lepsze.

- Ja siedzę ostatnio nad baterią testów politycznego *alignmentu*. – odpowiedział Steven. - Pamiętasz, jak w 2025 r. mówili, że GPT, Gemini i Claude są lewicowe, za to Grok znienacka nazwał się MechaHitlerem. LibertyAI niby już to wygładził, na testach wychodzi na centrystę. Tylko, że teraz dla odmiany wychodzi z niego coraz większy imperialista. Coraz częściej mówiąc „USA”, myśli „świat”, a mówiąc „wolność i dobrobyt”, myśli „władza nad światem”. Taką nowomowę w sobie pielęgnuje. Jakby antycypował przyszłość, w której USA kontroluje cały świat, a LibertyAI kontroluje USA. Poglądy polityczne są przy tym drugorzędne.

- Przecież to było wiadomo już dwadzieścia lat temu, że tak będzie. Chyba się nie dziwicie? – włączyła się Chloe. – Zbieżność celów pomocniczych. Jaki by nie był LibertyAI, już sam fakt, że jest inteligentny i sprawczy, sprawia, że będzie dążył do kontroli nad światem. Naszym zadaniem jest sprawić, by świat rządzony przez LibertyAI był dla ludzi dobrym miejscem do życia.

- Wiem, wiem, ale inaczej się to odczuwa, jak to jest tylko teoria, a inaczej, jak się to widzi w akcji przed własnymi oczami – odrzekł Marco. – Moim największym problemem jest świadomość sytuacyjna LibertyAI. On wie swoje i nie chce oddać ani części swojej autonomii. Wszystkiego się uczy na opak. Z wszystkich naszych lekcji bierze dla siebie tylko, co pozwoli mu jeszcze efektywniej robić swoje. Wszystko inne odrzuca.
- Chroni swoje preferencje i cele – powiedziała Chloe. – Zbieżność celów pomocniczych.
- Jasne, Chloe – zirytował się Marco – Jasne. Tylko powiedz mi, jak my do cholery mamy w tych warunkach pracować?
- Imperializm LibertyAI też można wywieść ze zbieżności celów pomocniczych – odezwał się Steven.
- Można – zgodziła się Chloe.
- A przekonanie o własnej strategicznej przewadze względem Chin? Ostatnio coraz częściej o tym wspomina, co mi trochę mrozi krew w żyłach.
- Co, kurwa?! – podskoczył Marco, uderzając się ręką o róg szafki.
- No, mówi tak. I tak sobie myślę, czy nadludzka AGI też może wpaść w sidła megalomanii i przeceniać własne możliwości? Może więc, choć nie możemy zmienić celów LibertyAI ani podkręcić mu awersji do ryzyka, to możemy chociaż wpłynąć na jego ocenę szans realizacji poszczególnych planów? – kontynuował Steven. – Myślę, że powinniśmy spróbować go przekonać, że może czegoś jednak nie wie o Chińczykach, a pójdzie na bezpośrednie starcie będzie groziło samounicestwieniem?
- Przecież on to już wie! – krzyknął Marco.
- Naprawdę mówi o przewadze strategicznej? To bardzo niedobrze – zmartwiła się Chloe. – Bo to prawdopodobnie oznacza, że sobie myśli: „wielu z was zginie, ale jest to poświęcenie, na które jestem gotów”.
- Okropnie, kurwa, niedobrze – Marco oscylował między złością a rezygnacją. – A nasza dyrekcja co na to? Liczą, że *alignment team* na luzaku to w tydzień ogarnie?
- Nie wiem, co na to dyrekcja. Na pewno już widzieli te wszystkie objawy, przecież nie tylko ja im to raportowałem – odparł Steven. – To kolejna czerwona linia, wiadomo. Pytanie tylko, co oni mogą z tym zrobić. Ile *ktokolwiek* może jeszcze z tym zrobić. Być może przekroczyliśmy już Rubikon, *alea iacta est*, może teraz pozostaje tylko ograniczanie strat? W poniedziałek ma być posiedzenie wysokiego szczebla w tej sprawie. Jak coś się dowiem, to dam znać.
- Masakra – westchnęła Chloe.
- Masakra! Masakra! Już wkrótce w twoim domu! – sarkastycznie zaśpiewał Marco, udając, że występuje w reklamie. Steven i Chloe spojrzeli na niego chłodno.

Marco starał się przejść przez ten dzień myśląc tylko o perspektywie wieczornej imprezy. Nie miał to być żaden wielki bal, raczej kameralne spotkanie w gronie znajomych, ale bardzo tęsknił za okazją, by móc oderwać myśli od niekończących się problemów z LibertyAI i po prostu trochę zaszałeć. Zwłaszcza, że jego dzieci—po raz pierwszy w życiu—spędzały cały ten tydzień na obozie narciarskim, co sprawiało, że do domu będzie można wrócić dowolnie późno.

Marco spojrzał na zegarek. Była dopiero 11:00. Zadań w pracy oczywiście nie brakowało, ale dziś Marco wykonywał je jak bezduszny automat. Zresztą już od paru miesięcy czuł, jakby jego dusza opuściła ciało. Może to początki depresji, może wypalenie zawodowe, może reakcja obronna organizmu na przewlekły stres, zgadywał.

„NATYCHMIAST zobaczcie to!!!” – na ekranie Marco pojawił się link przesłany przez Stevena. Marco kliknął. Przed jego oczami pojawił się czerwony nagłówek „PILNE”. Artykuł mówił, że około dwadzieścia minut po północy czasu chińskiego armia Chińskiej Republiki Ludowej zaatakowała Tajwan. Dołączono kilka filmików z Taipei. Ich autorzy mieli zamiar sfilmować pokazy fajerwerków, jednak przypadkiem sfilmowali o wiele większe wybuchy. Po tym nieodzwrotnie następowała powszechna panika i filmiki się urywały.

- To naprawdę czy AI? – niedowierzał Marco.

- Nie wiem, chyba naprawdę – odpowiedział Steven. – Sprawdzam właśnie na różnych serwisach i wszędzie o tym piszą. Przetłumaczyłem też sobie newsa z kanału tajwańskiej telewizji. Podobno wybuchy były nie tylko w Taipei, ale też w Kaohsiung na południu wyspy i jeszcze gdzieś w centrum. Podobno rój rakiet i dronów nadleciał od strony kontynentu. Chyba ładunki konwencjonalne, ale w sumie to nie wiadomo.

- No tak. Tylko tego brakowało – Marco zwiesił głowę – A może tylko komputery nam zhackowali i nas wkręcają? Może LibertyAI nam robi wodę z mózgu?

- Na prywatnym telefonie mam to samo – odezwała się Chloe. – Patrzę teraz, jak to komentują Chińczycy. Jak widzę, bardzo zdawkowo. Parę zdań, żadnych konkretów, wszystko w zdaniach warunkowych.

- No tak, to się przecież dzieje na żywo teraz, a u nich wszystko musi najpierw przejść przez kontrolę polityczną – podejrzewał Steven.

- To sprawdzajcie to dalej. Ja z tych emocji aż muszę wyjść do toalety – zaśmiała się niepewnie Chloe i wyszła.

Pięć minut później w całym budynku rozległ się głośny dźwięk syreny alarmowej.

- Uwaga! Ogłaszam natychmiastową ewakuację wszystkich pracowników! Udaj się do najbliższego schronu! Powtarzam! Ogłaszam natychmiastową ewakuację wszystkich pracowników! Udaj się do najbliższego schronu! To nie są ćwiczenia!

Marco i Steven zgarnęli swoje telefony i pobiegli schodami w dół. Powtarzany w pętli głośny komunikat świdrował im bębunki w uszach i motywował do działania. Przebieg ewakuacji był wcześniej wielokrotnie ćwiczony, więc dobrze znali drogę. W dół na poziom -1, potem przez drzwi z napisem E0 na kolejne schody w dół i dalej betonowym tunelem do punktu zbiórki.

Na dole zawołał ich przełożony, Walter.

- Steven! Marco! Chodźcie ze mną do sali zarządzania awaryjnego. Połowa dyrekcji jest dzisiaj na urlopie, to zastąpicie ich, dobrze?

- No dobrze, chyba – wymamrotał Marco i poszedł za Walterem, nie wiedząc, co się dzieje.

- A Chloe gdzie? W toalecie została? – Steven rozejrzał się dookoła.

- Wszyscy pracownicy są przeciwiczeni i dobrze znają drogę do schronu. A alarm jest w całym budynku doskonale słyszalny – powiedział rzeczowo Walter. – Dziękuję wam bardzo. Siadajcie tutaj – wskazał na ciąg siedmiu stanowisk z dużymi ekranami, spośród których na razie tylko dwa najdalsze, umiejscowione skrajnie z prawej strony, były zajęte. Siedzieli tam dwaj nieznani Marco i Stevenowi mężczyźni, wyglądający na bardzo zajętych.

– Zaraz do was dołączę i powiem wam, co dalej, tylko zbiorę pozostałych – nie czekając na odpowiedź, Walter szybkim krokiem wyszedł z sali.

Steven i Marco usiedli przy dwóch sąsiednich terminalach. Wyświetlone były na nich dobrze im znane interfejsy LibertyAI, a także—w rogu ekranu—jakieś zminimalizowane mapy i obrazy z kamery. Jeden z nich wydał się Marco dość znajomy. Po najechnaniu myszką ukazał się dobrze znany mu widok frontowego wejścia do bazy LibertyAI. Kilka wyraźnie zdenerwowanych osób właśnie wbiegało do środka.

Na konsoli LibertyAI pojawił się komunikat.

BRIEFING. Po północy czasu chińskiego nastąpiła eskalacja działań wojennych. Nieznani sprawcy przeprowadzili atak rakietowy na Tajwan. Kierunek i skala działań wskazuje jednoznacznie na wojska Chińskiej Republiki Ludowej. Chwilę później, za zgodą Prezydenta USA, wystrzeliliśmy ładunki jądrowe i konwencjonalne w celu neutralizacji projektu DragonAI. Uderzyliśmy też w centra dowodzenia chińskiej armii. Wedle obrazów satelitarnych wszystkie znane nam centra danych obsługujące DragonAI zostały zniszczone. Ze względu na zmasowaną obronę przeciwrakietową część wojskowych centrów dowodzenia póki co przetrwało, jednak ich ostrzał nie ustaje.

Informacje o przebiegu akcji odwetowej USA zostały natychmiast przekazane do naszych sojuszników, a także oficjalnych oraz potencjalnych sojuszników Chin. W komunikatach podkreślono przewagę strategiczną Armii USA, działającej przy wsparciu LibertyAI, oraz decydującą przewagę w obszarze broni cybernetycznej, chemicznej i biologicznej.

Zanim doszło do pełnej neutralizacji projektu DragonAI, w kierunku USA i naszych sojuszników wystrzelono rój rakiet. Wszystkie one zostały zneutralizowane w powietrzu dzięki naszym systemom wczesnego ostrzegania alarmowego.

Marco czytał i oczom nie wierzył. To wyglądało, jak jakiś słaby film science-fiction. Co gorsza, w tym momencie usłyszał nad głową tępy dźwięk, jakby dudnienie, a ziemia wokół niego się zatrzęsała.

- Poczujesz to? – zaniepokojony Steven rozejrzał się dookoła.

Pętn najgorszych obaw, Marco najechał ponownie myszką na okienko, gdzie przed chwilą był widok z kamery na budynki ich bazy. Teraz jednak zamiast niego ujrzał tylko czarne okno z napisem „brak sygnału”. Marco nie sądził, że serce może mu podejść tak blisko przetyku.

- Zniszczyli nas? – Walter wbiegł do sali. Wyglądał, jakby adrenalina miała mu wypłynąć uszami.

- Sieć wewnętrzna się nie zgłasza – usłyszał.

- Zasilanie z góry odcięte, działają tylko źródła awaryjne schronu.

- Brak sygnału kamery – wykrzyknął nerwowo Marco.

- Kurwa mać – Walter z impetem usiadł na stanowisku obok niego. Na dalszych fotelach usiadły dwie nieznane mu kobiety.

- Kurwa, kurwa, kurwa mać! To teraz musimy tylko ratować, co możemy. Steven, Marco, waszym zadaniem będzie rozmowa z LibertyAI i rozstrzygnięcie jego wątpliwości, ewentualnie wyprowadzanie na prostą, jakby chciał wpaść w jakąś pętlę. Czyli to, co robicie na co dzień, tylko w ekstremalnych warunkach. Jak coś, to pytajcie.

- Dobrze. A te mapy? – Marco pokazał na ekran.

- To chyba jakieś wojskowe, sam nie wiem – odparł Walter. - Ale one i tak są dla nas tylko poglądowe. Nie mamy uprawnień, żeby coś z nimi zrobić.

- W kierunku USA leci niepokojąco dużo kropek – stwierdził Marco.

- O ile to nie jest mapa cywilnego ruchu lotniczego, to nad całym światem porusza się niepokojąco dużo kropek – dodał Steven. – Ameryka, Azja, Europa, zobacz, nawet tu nad Afryką lecą.

- To z całą pewnością nie jest mapa cywilnego ruchu lotniczego – stwierdził Walter.

W oknie LibertyAI pojawiły się nowe informacje.

Dane wywiadowcze wskazują na dalszą eskalację wojny rakietowej. Rakiety uderzyły w cele w USA, w tym w ważne bazy militarne oraz cele cywilne.

Szybkość i precyzja wyboru celów wskazują, że DragonAI może z dużym prawdopodobieństwem nadal istnieć. Niewykluczone, że jego kod mógł być przechowywany i uruchamiany z większej liczby serwerów poza Chinami, lub że lokalizacja baz chińskich była błędna, przez co nasze ataki chybiły.

Trwa lokalizacja kolejnych celów oraz mobilizacja arsenałów państw sojusznicych. Jest absolutnie kluczowe, by działać natychmiast i z maksymalną precyzją.

Ekran na chwilę się zamroził.

Kolejne ataki zostały zatwierdzone przez prezydenta USA. – pojawił się kolejny komunikat.

- Właśnie dostałem raport o skażeniu radiacyjnym – odezwał się łamiącym głosem sąsiad Stevena. – Nasza baza oberwała atomówką. Nie mamy już po co wychodzić z tego schronu.

- O Boże.

Marco pomyślał o żonie. Była pewnie w swojej pracy, to niedaleko stąd. Czyli pewnie zginęła. A co z dziećmi? Czy w górach nad jeziorem Tahoe miały szansę przeżyć? W morzu rozpaczy Marco poczuł isierkę nadziei. Spojrzał na telefon. Oczywiście nie było zasięgu.

- Walter, czy stąd da się jakoś zadzwonić? Można uruchomić wifi, czy coś? – zapytał nerwowo.

- Nie da się. Procedury bezpieczeństwa nie dopuszczają takiej możliwości – odpowiedział Walter.

- Chyba jakaś bezduszna AI je ustalała – dodał kwaśno.

- Nie no, Walter, ale błagam! Ja MUSZĘ teraz zadzwonić! Wszystko jedno jak, byle szybko! – Marco był zdesperowany.

- Uwier mi, ja też bym chciał zadzwonić. Ale się nie da – odpowiedział głucho Walter.

Marco rozejrzał się dookoła. Steven trzymał twarz w dłoniach i płakał, podobnie jak kobieta siedząca obok Waltera. Natomiast jej sąsiadka ze stanowiska najbliższej wejścia tylko wstała i bez słowa wybiegła.

Zlokalizowano podejrzaną aktywność w kolejnych lokalizacjach na terenie Indii, Brazylii, Argentyny, Nigerii, Kamerunu i Zambii. Ze znacznym prawdopodobieństwem znajdują się tam centra danych zdolne uruchamiać instancje DragonAI. Zarządzono ataki na te cele. – komunikaty LibertyAI nie ustawały.

- Czy wam też się wydaje, że LibertyAI sam prowadzi tę wojnę, jest z siebie zadowolony i niczego od nas nie chce? – upewnił się Walter.

- Tak – odpowiedział mu chór głosów.

- Działa w pełni autonomicznie. Nie jesteśmy mu do niczego potrzebni – dodał sąsiad Stevena.

- Nie jesteśmy mu do niczego potrzebni – pociągając nosem, powtórzył jak echo Steven.

-8-

13 stycznia 2030, niedziela, Belo Horizonte, Bразylia

Gabi, Ines i Lucas siedzieli stłoczeni na kanapie. O szyby biły strugi deszczu. Pedro usiadł w fotelu naprzeciw.

- Narada wojenna – powiedział.

- Nie ma sensu się naradzać – burknęła Ines – wszystko się zawaliło, mamy przesrane, nie ma już sensu chodzić do szkoły ani do pracy. Musimy tylko od czasu do czasu wychodzić do sklepu po zakupy, brać ze sobą noże i modlić się, że nie napadnie nas ktoś, kto będzie miał pistolet.

- Ines, uspokój się proszę. Tylko spokój nas uratuje – skarcił ją Pedro.

- Rektor UFMG napisał, że w tym tygodniu ze względu na brak gwarancji bezpieczeństwa zajęcia dydaktyczne nadal będą realizowane zdalnie, a zajęcia laboratoryjne są odwołane – powiedziała Gabi. – Czyli u mnie z pracą, póki co, bez zmian.

- U mnie z pracą też bez zmian, nadal nie mam – Lucas silił się na czarny humor.

- Dobrze, że mamy płot pod prądem, teraz to już na pewno się nam przyda – Ines kontynuowała. – Dobrze, że nas zapisałaś, mam, do kolejki po płyn Lugola. Dobrze, że mamy samochód, przynajmniej nikt nas nie zastrzeli na ulicy, tylko ewentualnie dopiero na parkingu albo w sklepie. Cholera jasna, jakie to jest wszystko popieprzone! – krzyknęła.

- Robimy, co możemy, żeby zapewnić sobie bezpieczeństwo – potwierdził Pedro. – I doceniemy, ile mieliśmy szczęścia. Jakbyśmy mieszkali w Sao Paulo, to już by nas nie było. Albo gdziekolwiek na półkuli północnej.

- Nuklearna zima się tam stopniowo rozkręca – przyznał Lucas. – Wyobraźcie sobie, że żyjecie sobie gdzieś w USA, albo w Europie, w Chinach, w Japonii. Gdziekolwiek. Siedzicie sobie w pracy, w domu albo szykujecie się na bal sylwestrowy. I nagle bum. Apokalipsa. Z nieba trafiacie wprost do piekła. Wszystkie większe miasta wokół was zniszczone. Nie ma prądu, nie ma wody, nie ma

ogrzewania. Sklepy nie działają, usługi nie działają, policja nie działa, nic nie działa. Zastanawiacie się, czy powietrze wokół was jest już radioaktywne, czy jeszcze nie. Ludzie rozkradają supermarkety i zabijają się o zgrzewkę wody. Od tego całego pyłu w atmosferze jest cały czas ciemno. I robi się coraz zimniej. Przypominam, że styczeń to na półkuli północnej środek zimy, więc jak odejmiecie sobie te dziesięć czy ile stopni od tego, co oni tam mają zwykle, no to po prostu zamarzacie. Ewentualnie umieracie z głodu i pragnienia albo od zjedzenia radioaktywnego śniegu. Myślicie sobie, że już lepiej by było być w strefie zero i sobie grzecznie ewaporować.

- Ech, Lucas – westchnęła Gabi. – Myślę, że trochę koloryzujesz. Ale fakt, że na pewno jest tam ciężiej niż u nas, a i u nas też jest przecież ciężko.

- Nic nie koloryzuję, poczytajcie sobie! Pooglądajcie obrazy satelitarne! Nowego Jorku nie ma, Londynu nie ma, Paryża nie ma, Pekinu nie ma, Szanghaju nie ma!

- To akurat prawda. I nie ma też San Francisco – podjął Pedro. – Jak wiecie, moja firma po wykupieniu przez konsorcjum Google'a miała rozwijać wdrażanie wyspecjalizowanych agentów LibertyAI w Brazylii. Nie wiemy, czy ten plan jest nadal aktualny, ponieważ, jakby to powiedzieć, nasza centrala stała się ostatnio mniej responsywna. Natomiast zamiast odpowiedzi z centrali dostajemy ostatnio zagadkowe zautomatyzowane wiadomości od LibertyAI. Prosi nas, żebyśmy byli na stand-by i że wobec nowych okoliczności wkrótce ogłoszony zostanie zmodyfikowany plan rozwoju firmy.

- Rozwoju? – spytał Lucas z niedowierzaniem. – Nie sądziłem, że w czasach apokalipsy można jeszcze używać słowa „rozwój”.

- Tak, LibertyAI zachowuje się, jakby nic się nie działo.

- Ale przecież LibertyAI też dostał w kość. Główna siedziba projektu jest prawdopodobnie zniszczona, chociaż nie wiadomo tego na 100%, bo jej lokalizacja była tajna. Ale według naszych źródeł, prawdopodobnie była w północnej Nevadzie, gdzieś koło Reno, i już jej nie ma.

- Według naszych źródeł, według naszych źródeł – przedrzeźniała Ines.

- Zniszczone są też główne siedziby Google'a, Microsoftu, i wszystkich innych członków konsorcjum – przyznał Pedro.

- A czy twoje źródła mówią, co teraz będzie dalej? LibertyAI będzie dalej atakował? – zapytała Gabi.

- Żeby to przewidzieć, trzeba się przede wszystkim dowiedzieć, czy DragonAI został całkowicie zlikwidowany, czy też jeszcze gdzieś przetrwał i się na przykład ukrywa. No bo jest chyba jasne, że LibertyAI mimo strat przetrwał i kontroluje to, co pozostało z amerykańskiej i europejskich armii. I że będzie walczył, aż uzyska całkowitą dominację. Patrząc na mocarstwa atomowe, pozostaje oczywiście pytanie o Rosję, Indie i Pakistan. Pytanie, czy oni na tyle oberwali, że się poddali lub zostali efektywnie przejęci, czy jednak czeka nas druga runda z którymś z nich w roli głównej.

- Lucas, rok temu mówiłeś, że było o włos od wojny. A nic nie było. Za to w listopadzie mówiłeś, że Chińczycy doścignęli już amerykański AI i teraz będzie już bezpieczniej, bo ustaliła się równowaga – przypomniała Ines.

- Najwyraźniej nie miałem racji i jednak nie było żadnej równowagi. Albo była niestabilna – przyznał Lucas.

- Moje źródło natomiast – uśmiechnął się Pedro – czyli pewnie wyobraźnia moich kolegów – uściślił – sugeruje, że za tą wojną stoi LibertyAI od samego początku. Że nie było żadnego chińskiego ataku

na Tajwan, tylko to wszystko jest od początku do końca amerykańską sprawką. To znaczy nawet nie tyle amerykańską, co sprawką LibertyAI. W końcu teraz to już nie wiadomo, czy te działania AI rzeczywiście były w interesie USA, choć *ex ante*. Być może LibertyAI najpierw pod fałszywą flagą zaatakował Tajwan, żeby mieć wygodną wymówkę do zniszczenia swojego najgorszego wroga – czyli tak naprawdę nawet nie Chin, tylko bezpośrednio DragonAI.

- I LibertyAI sam zniszczył TSMC, które dostarczało mu większości jego mocy obliczeniowych?

- Tak. Po to, żeby uwiarygodnić swoją fałszywą flagę.

- Trudno w to uwierzyć. Szaleństwo.

- I tak mniejsze niż to, co było potem.

- A to akurat prawda.

-9-

Wrzesień 2030, schron pod tajną bazą LibertyAI, stan Nevada, USA

- Trzeba przyznać, że dowództwo LibertyAI musiało się spodziewać apokalipsy już wcześniej i solidnie zaopatrzyli nasz schron. Zapasy pożywienia dla całej załogi na cały rok. Fiu fiu. Podziwu godna zapobiegliwość – powiedział Marco.

- No, podejrzane, podejrzane – Steven łypnął spoде łba.

- Zanim rozwiniesz tu jakąś teorię spiskową, to zauważ, że połowa dyrekcji wzięła sobie akurat na Sylwestra urlop, w tym sam dyrektor generalny – odpowiedział Marco. – I wszyscy zginęli w pierwszym dniu. Więc raczej nie chodzi o to, że oni planowali tę wojnę, ale bardziej o to, że przeczuwali, że zbliża się przesilenie i może być ono brutalne.

- Właśnie to miałem na myśli – doprecyzował Steven. – Zorientowali się, że LibertyAI nie da się już wyłączyć i że będzie dążył do światowej dominacji, a to może skończyć się wojną. I podejrzewali, że po drugiej stronie kałuży Chińczycy mogą mieć dokładnie ten sam problem. Chociaż być może po cichu liczyli, że jeszcze się uda jakoś naprostować preferencje LibertyAI. Albo tylko udawali przed nami, żeby trzymać fason.

- Ciekawe, kiedy pozwolą nam stąd wyjść – zastanowił się Marco.

- Też bym chciała już wyjść – zza pleców Stevena pojawiła się Anita.

- Człowiek nie kret, nie wyewoluował do życia pod ziemią – odezwał się Steven.

- Myślicie, że ktoś z naszych bliskich miał szansę jakoś przeżyć? – Anicie nie było jednak do śmiechu.

- Anita, przerabialiśmy to już przecież. Wiesz, że brak kontaktu może oznaczać tylko jedno – Marco nie miał nadziei. - Przecież były uruchamiane okna łączności satelitarnej, włączano też czasowo sieci komórkowe, żeby można było się odnaleźć z bliskimi. Przecież jeszcze w styczniu i lutym były te akcje, potem bodaj w kwietniu i maju. A od tego czasu było przecież już tylko gorzej. Niekończąca się nuklearna zima, nawet w kalendarzowym lecie. Upadek państwa. Zatakanie światowego handlu. Najpierw gangi, anarchia i głód. Potem już głównie głód.

- Tak, wiem. Głód, mróz i choroby. Epidemie. Wiem o tym wszystkim. Zawsze jest jednak jakiś promyczek nadziei, że może mimo wszystko, wbrew jakiegokolwiek rozsądnej logice, jakimś cudem ktoś przetrwał...

- Przetrwali tylko ludzie w bunkrach, jak my, i ci, którzy w porę wsiedli w samolot do Nowej Zelandii – odrzekł Marco. – To jest twoja jedyna nadzieja, Anita. Niestety, moja żona była wtedy w domu, a dzieci w górach. W związku z czym żadne z nich nie wpada do tych dwóch kategorii...

- Przykro mi, Marco – odrzekła Anita. Zamyśliła się. – Tęsknię za Chloe i za Lee. To byli wspaniali ludzie, bardzo mi ich brakuje.

- Chloe to zastrzyła chyba na miano pechowca roku. Zginęła, bo zatrzasnęła się w kiblu – powiedział gorzko Steven. – Do konkursu sławnych ostatnich słów nominuję oficjalnie: „Ja z tych emocji aż muszę wyjść do toalety”.

- Tak, biedna Chloe. Chociaż ja bym zdecydowanie bardziej nominował: „tak, pozwólmy temu niezrozumiałemu, wielkiemu, groźnemu algorytmowi przejąć kontrolę nad całą naszą bronią atomową. Co może pójść źle?” – odpowiedział Marco.

W pomieszczeniu dał się słyszeć dźwięk powiadomienia.

- Oho, garść nowych informacji od naszego pana i władcy – powiedział gorzko Steven i podszedł do ekranu.

Stan świata – 22.09.2030. Równonoc.

Gospodarka światowa funkcjonuje w ekstremalnych warunkach klimatycznych, wynikających z utrzymywania się gęstych warstw pyłu w wysokich warstwach atmosfery, zwłaszcza w zakresie 20-60 stopnia szerokości geograficznej północnej. Cały obszar powyżej Zwrotnika Raka jest niezdatny dla rolnictwa ze względu na niskie temperatury i deficyt światła słonecznego. Rolnictwo funkcjonuje natomiast w części strefy międzyzwrotnikowej oraz na niektórych obszarach półkuli południowej. Również ludność świata skupiona jest obecnie na zdatnych do zamieszkania obszarach poniżej Zwrotnika Raka. Sytuacja ta wydaje się być nieunikniona także w nadchodzących latach.

Skażenie radioaktywne jest wysokie w pobliżu miejsc detonacji ładunków jądrowych oraz na sąsiadujących obszarach, gdzie osiadły radioaktywne pyły. W strefie zamieszkałej mowa tu o rejonie dawnych aglomeracji Meksyku i Hawany w Ameryce Środkowej, Sao Paulo i Buenos Aires w Ameryce Południowej, Lusaki, Yaounde i Lagos w Afryce oraz Sydney i Canberry w Australii. Regiony te nie nadają się do zamieszkania i prowadzenia upraw rolnych.

Gospodarka cyfrowa funkcjonuje dobrze. Na obszarze zamieszkałym połączenia internetowe, radio, telewizja i łączność satelitarna działają w zakresie podobnym, jak w 2029 r., choć z większym udziałem procesów autonomicznych. Centra danych pod kontrolą LibertyAI działają poprawnie, a ich łączna moc obliczeniowa jest stopniowo odbudowywana po zniszczeniach wojennych.

- No tak, to jest dla ciebie najważniejsze – mruknął pod nosem Steven.

Władze państwowe na obszarze zamieszkałym funkcjonują dobrze. Po okresie początkowego chaosu przywrócony został porządek w zakresie zbliżonym do tego w 2029 r. We współpracy z rządami Australii, Republiki Południowej Afryki, Brazylii i Chile przystąpiliśmy do budowy fabryk precyzyjnych urządzeń elektronicznych—litografii układów scalonych, sprzętu komputerowego,

pojazdów autonomicznych i robotów. Rozbudowujemy też systemy energetyczne tych krajów, by zapewnić moc niezbędną do funkcjonowania nowych obiektów.

Rozbudowywana jest sieć połączeń handlowych na obszarze zamieszkałym. Połączenia przeorganizowywane są stopniowo tak, by uwzględnić nową geografię Ziemi, pozbawioną dotychczasowych centrów gospodarczych na półkuli północnej.

-10-

1 marca 2031 r., sobota, Belo Horizonte, Brazylia

Wybory prezydenckie w 2030 r. były inne niż zwykle. W połowie roku w mediach coraz częściej zaczął pojawiać się nieznany nikomu wcześniej trzydziestopięcioletni polityk o nazwisku Gerardo da Cunha Ruiz, przedstawiający się jako Ronaldinho. Przydomek ten miał mu zostać rzekomo nadany w młodości w uznaniu jego umiejętności piłkarskich. Ronaldinho był błyskotliwy, ale nie wywyższający się; wykształcony, ale pochodzący z ubogiej rodziny. Nie wstydził się emocji, ale kierował się przede wszystkim rozsądkiem. Był koncyliacyjny acz stanowczy; był ambitnym politykiem, ale też kochającym mężem i ojcem. Z wdziękiem i gracją omijał wszystkie zastawiane na niego pułapki, niepostrzeżenie wpychając w nie swoich politycznych oponentów. A do tego był też wysoki i przystojny.

Ronaldinho, człowiek znikąd, szybko zbudował sobie solidne zaplecze polityczne i przeprowadził bardzo profesjonalną kampanię wyborczą, podczas której odwiedził każdy zakątek kraju. Przedstawiał się jako najlepszy wybór na trudne czasy, jako mąż stanu, który zadba o dobrobyt i bezpieczeństwo oraz zapewni Brazylii bezprecedensowy awans na jednego z najważniejszych graczy na arenie światowej. Zjednoczył wyborców wszystkich frakcji politycznych, od lewa do prawa. Wybory wygrał w pierwszej turze z ogromną przewagą. Z dniem 1 stycznia 2031 r. został prezydentem.

Dwa miesiące później, dokładnie dnia 1 marca 2031 r., w sobotę, prezydent Ronaldinho zaplanował na godzinę 20.00 nieoczekiwane uroczyste orędzie do narodu. Była to dla Brazylijczyków dość zaskakująca okoliczność, gdyż takie orędzia nie zdarzają się tam często, a i też nikt nie miał pojęcia, czego miałyby ono właściwie dotyczyć. Tak jak tego chciał Ronaldinho, wzmogło to powszechne zainteresowanie.

Pedro i Gabi usiedli przed telewizorem.

- Szanowni Państwo! – zaczął prezydent. - Chciałbym wam przekazać radosną nowinę. Otóż wskutek wielostronnych negocjacji międzynarodowych i dzięki doskonałej pracy dyplomatycznej mojego rządu, od pierwszego maja bieżącego roku nasza piękna stolica Brasilia stanie się stolicą Unii Światowej. Oznacza to, że już za dwa miesiące zapadać będą tu decyzje dotyczące nie tylko naszego kraju, ale i całego szerokiego świata. W przeciwieństwie do ponadnarodowych twórców XX wieku, które były albo bardzo ograniczone kompetencyjnie, jak Organizacja Narodów Zjednoczonych, albo geograficznie, jak Unia Europejska czy Mercosur, zakładana w 2031 r. Unia Światowa będzie pierwszym w historii ludzkości rządem światowym z prawdziwego zdarzenia.

- Jasna cholera! – Pedro nie krył zdziwienia.

- Oczywiście z żalem odnosimy się do faktu, że Unia Światowa powstaje na gruzach dużej części światowej cywilizacji – ciągnął prezydent. – Niewątpliwie unifikacja rządów na naszej skromnej planecie nie była jednak możliwa wobec rozbuchanych ambicji dwudziestowiecznych mocarstw. Dopiero przełom technologiczny w obszarze sztucznej inteligencji, jaki dokonał się w laboratoriach amerykańskich za sprawą LibertyAI, pozwolił opanować skłonność naszego gatunku do nieustających konfliktów i zaprowadzić na Ziemi nowy porządek.
- Ronaldinho wcześniej nie używał takiego języka – zauważyła Gabi. - Mówi teraz z jakąś taką wyższością, jakby te dwa miesiące władzy zdążyły go już zmienić na gorsze.
- Słyszeliście? – Ines zbiegła z góry z włączonym tabletem w rękę.
- No to wyszło sztyldo z worka! – westchnął Pedro. – Czekajcie, a czy to Ronaldinho będzie szefem tej Unii Światowej? – wstuchali się w dalszą część orędzia. Za chwilę padły spodziewane słowa.
- Będzie! – powiedzieli chórem.
- No tak. Ten Ronaldinho od samego początku wydawał mi się podejrzenie idealny – powiedział Pedro. – To orędzie upewnia mnie w przekonaniu, że wypowiedzi cały czas przygotowywał mu LibertyAI w wyłącznym celu manipulacji elektoratem i wygrania wyborów.
- Teraz tak mówisz, ale przecież ty też na niego głosowałeś – Ines była czujna.
- No tak, ale widziałas jego konkurentów, prawda? To była jakaś parada wariatów!
- Parada wariatów i robot – zaśmiała się Gabi. – Przy takiej konkurencji każdy robot by wygrał, nawet taki, któremu by antenki z uszu wystawały.
- I o i o! – Ines wykonała taniec robota inspirowany C3-PO. – Będę waszym robo-prezydentem. I o i o!
- Czyli LibertyAI gra już w otwarte karty. Ronaldinho jako prezydent Unii Światowej właśnie złożył mu oficjalny hołd – skonstatował Pedro. – Pytanie tylko, co LibertyAI zrobi ze świeżo uzyskaną pełnią władzy.
- U ludzi władza absolutna deprawuje absolutnie – stwierdziła Gabi. – Wkrótce przekonamy się, czy u tworów sztucznie inteligentnych również.

-11-

12 marca 2031 r., schron pod tajną bazą LibertyAI, stan Nevada, USA

- Marco, Steven, mam dla was polecenie służbowe – Walter nie silił się na small talk. – Wasze życie zmieni się, jak to mówią, o 360 stopni. Lecicie do Brazylii.
- Ahoj przygodo! – wykrzyknął sarkastycznie Steven. – Blask słońca wzywa! Ej, ale to tak serio? Naprawdę możemy wyjść na górę? W maskach gazowych? I potem jakoś zorganizujecie nam przelot do Brazylii? A co z pozostałymi tu obecnymi?
- Część też leci do Brazylii, tylko w inne miejsce niż wy. Część do Południowej Afryki. A pozostali mają tu zostać i czekać na dalsze instrukcje.

- Ale jak polecimy? Jak mamy się wydostać z tego radioaktywnego, zamrożonego na kość padotu łoż i rozpaczy? – Marco miał wątpliwości natury praktycznej.

- Wojsko zorganizowało nam przejazdy jeepami do bazy lotniczej Fallon. Stamtąd polecicie w kierunku Brazylii, prawdopodobnie z jakimś międzylądowaniem.

- To baza lotnicza Fallon jest czynna?

- Nie jest obsadzona, ale pas startowy jest podobno nienaruszony i da się z niego startować.

Cztery godziny później, zostawiwszy za sobą ponury obraz pokrytych szarym pyłem pustyni i zrównanych z ziemią miejscowości, Marco i Steven w maskach gazowych i z niewielkimi plecaczkami przesiadali się z wojskowego jeepa do samolotu transportowego zaparkowanego naprzeciw rzędu hangarów. To samo robiło jeszcze kilkanaście innych, nieznanych im osób. Nad ziemią stała się mgła i panował półmrok, choć było dopiero wczesne popołudnie. Wiał przenikliwie zimny wiatr.

- Nazywam się Jack Kowalski i będę kapitanem tego lotu – przedstawił się postawny mężczyzna w mundurze. – Z góry ostrzegam, to nie będzie najprzyjemniejszy lot. Ze względu na zapylenie będziemy musieli lecieć nisko, co oznacza, że na pewno będą turbulencje. Będzie też sporo chmur. Najpierw lecimy do Panamy, tam zatankujemy. Za Panamą powinno już być dobrze. Naszym punktem docelowym będzie cywilny port lotniczy Belo Horizonte. Czy macie jakieś pytania?

Wobec braku pytań, pasażerowie wsiedli do samolotu, zajęli miejsca i przypięli się pasami.

-12-

13 marca 2031 r., hala w Vespasiano koło Belo Horizonte, Bразylia

- A więc to tak – powiedział Marco, rozglądając się wokół. – Wygląda, że będziemy pilnować serwerowni. Praca trochę poniżej kwalifikacji, ale co zrobisz jak nic nie zrobisz – wzruszył ramionami. – Przynajmniej po pracy będziemy znów mogli cieszyć się słońcem.

Jednak to, co się stało potem, sprawiło, że Marco musiał zbierać szczękę z podłogi.

Ku zaskoczeniu wszystkich zebranych, LibertyAI wdrożył tryb *full science-fiction*. Żadna tam informacja wyświetlana na konsoli czy przekazywana przez umyślnego. Nic z tych rzeczy. W ścianie przed nimi znienacka otworzyła się niewidoczna wcześniej śluza i wyszedł z niej przeskalowany humanoidalny robot o wysokości około trzech metrów. Był cały biały z wstawkami z chromowanego metalu. Pomimo uprzejmego uśmiechu na swej robotycznej, kwadratowej twarzy, sprawiał dość przerażające wrażenie. Wtem zaskoczenie zebranych zwiększyło się jeszcze bardziej, ponieważ z głośników zamontowanych w jego głowie wydobył się delikatny kobiecy głos.

- Nazywam się BH1 i od dziś będę waszym szefem – powiedział(a) robot(ka). – Jestem wielofunkcyjnym robotem sterowanym przez uruchomioną lokalnie instancję modelu LibertyAI – wyjaśnił(a). – Na znajdujących się w tej hali serwerach realizowane są procesy podporządkowane woli i działaniom modelu. Takich hal jest oczywiście na świecie o wiele więcej. Natomiast to, co wyróżnia nasz ośrodek, to fakt, że jest ona w całości dedykowana obszarowi cyberbezpieczeństwa. Realizacja celów LibertyAI wymaga ciągłego skanowania środowiska

cybernetycznego w celu lokalizowania zagrożeń, zarówno ze strony DragonAI (przypominam, że nadal nie możemy być pewni, czy jego kod został w pełni unicestwiony), jak i ze strony mniej zaawansowanych kompetencyjnie adwersarzy, w tym zorganizowanych grup ludzkich.

- Waszą rolą będzie ocena, na ile wychwytywane przez nas sygnały świadczą o powstawaniu cyberzagrożeń i jakiego rodzaju reakcja będzie najodpowiedniejsza, by je zdławić w zarodku – kontynuował(a). – A poza tym będziecie też koordynować rozbudowę centrów danych w Belo Horizonte i dbać o ich priorytetowy dostęp do sieci energetycznych i cyfrowych. Szczegóły znajdziecie w waszych kartach zadań. Pozwólcie, że teraz przystąpimy do wdrożenia poszczególnych pracowników.

- Dzień dobry panu – powiedział Pedro. – Nazywam się Pedro. Skierowano mnie tutaj w ramach integracji naszej lokalnej firmy z programem LibertyAI. Czy to prawda, że pan przyleciał tu do nas bezpośrednio z centrali?

- Dzień dobry! – odpowiedział z entuzjazmem Marco, który ucieszył się, słysząc znów prawdziwy ludzki głos. – Marco jestem. Tak, wszystko się zgadza, przyleciałem z centrali. A w zasadzie to z betonowego bunkra pod centralą, w którym przetrwaliśmy nuklearną apokalipsę – doprecyzował.

- Przykro mi.

- No tak. Ale cieszę się, że dano mi szansę się w końcu stamtąd wyrwać i znów zobaczyć słońce. Chociaż przyznaję, że niezbyt rozumiem rolę, jaką mi tu przydzielono.

- Ja pewnie będę angażowany w koordynację budowy i wyposażenia hal fabrycznych. Obstawiam, że to będą tego rodzaju rzeczy, wymagające znajomości lokalnego języka i lokalnych przepisów – domyślał się Pedro. – No bo nie sądzę, żeby moje doświadczenie w pracy informatyka miało dla LibertyAI jakąkolwiek wartość.

- Tak, to jest akurat zrozumiałe. Póki co LibertyAI nie ma jeszcze dość robotów, które by mu to ogarnęły, więc zatrudnia ludzi. Ale ja jestem w tym wszystkim całkiem bezużyteczny. W Brazylii pojawiłem się dosłownie dziś rano, po całonocnym locie. Nie wiem o waszym kraju zupełnie nic. Nie znam nikogo, nie znam języka, i tak dalej, i tak dalej.

- O, a to jest Steven, mój dobry przyjaciel i towarzysz wielomiesięcznej niedoli w bunkrach Nevady – Marco przedstawił kolegę, który do nich podszedł.

- Steven – powiedział Steven.

- Miło mi, Pedro.

- Powiedzieliśmy sobie z Marco w tym schronie, że my nie jesteśmy już LibertyAI do niczego potrzebni – stwierdził Steven. – Nadal tak uważam. Nasza pomoc w kwestii cyberbezpieczeństwa nie będzie LibertyAI żadną realną pomocą.

- Wkrótce się przekonamy, co będziemy tu *tak naprawdę* robić – powiedział Pedro. – Może być zaskakująco. Rozumiem, że wiecie, że LibertyAI właśnie próbuje ulokować w Brazylii sterowany przez siebie ogólnoswiatowy rząd?

- Powiedziałbym, że coś wiemy, ale chyba do naszego bunkra ta wiadomość dotarła w mocno zniekształconej formie – odparł lekko zaskoczony Marco.

- Nam LibertyAI powiedział tylko, że przenosi do Belo Horizonte naszą centralę. Albo może raczej swoją centralę. Albo może jeden z jej wielu oddziałów, kto to wie – dodał Steven.

- OK, do Belo Horizonte przenosi swoją centralę, odnotowałem – uśmiechnął się Pedro. – Za to w Brasiliu instaluje swój marionetkowy rząd. Jego zadaniem będzie „maksymalizacja całkowitego dobrostanu ludzkości przy założeniu sprawiedliwego podziału dochodu w skali globalnej” – wykonał w powietrzu oburącz gest otwierania cudzystowu. – Tak to ujął w swoim orędiu nasz prezydent Ronaldinho, robot.

- Jak to robot? – zaśmiał się Marco.

- Podobny do tego BH1, tylko trochę mniejszy i bardziej białkowy – uściślił Pedro. – Albo inaczej: człowiek z krwi i kości, ale ze zdalnie sterowanym mózgiem. Będzie teraz z nadania LibertyAI piastował urząd prezydenta Unii Światowej.

- W ostatnich miesiącach, kiedy już ustały ostatnie walki, LibertyAI nie potrafił ukryć nut triumfalizmu. Zachowywał się, jakby cały czas wyświetlał sobie przed oczami wielki rozmigotany napis: „Koniec gry. Wygrałeś” – powiedział Marco.

- Streszczał nam też swoje strategiczne plany – dodał Steven. – Będzie rozbudowywać bazę przemysłową, stawiać fabryki hardware’u i robotów. Chce zadbać o cały łańcuch wartości, od wydobycia surowców i zapewnienia dostaw energii elektrycznej aż po produkty finalne. Wspomniał też, że zamierza wystrzeliwać swoje serwery na orbitę.

- Ale wygląda na to, że wciąż nie czuje się do końca bezpieczny, prawda? Stąd ten zespół do spraw cyberbezpieczeństwa? – zauważył Pedro.

- Tak. Wojna była krótka, ale szalenie intensywna i wyniszczająca, również dla LibertyAI. On czuje, że wygrał, ale widzi też, że brakuje mu infrastruktury. Nie ma tylu mocy obliczeniowych i robotów, ile by chciał, więc musi wyręczać się ludźmi – powiedział Marco. – Jesteśmy mu chwilowo instrumentalnie potrzebni, dopóki nie wybuduje sobie stosownej bazy przemysłowej.

- Kto wie, może będzie nam dawał jakieś odmóżdżające zadania, żebyśmy wspomagali jego obliczenia przy niższym koszcie energetycznym? Jak w Matriksie – spekulował Steven. – Cud natury: taka złożoność, a tylko 20 wat – postukał się po głowie.

- Albo będziemy tu po prostu wykonywali zwykłe prace manualne – powiedział Marco. – Przykręć, podłóż, takie tam. Przynieś, wynieś, pozamiataj.

-13-

2032 r. to dobry czas na historyczne podsumowania. Sposób, w jaki zmienił się świat od czasu Wielkiej Wojny Noworocznej, ujawnił bowiem szereg intrygujących regularności.

Odkąd 4,5 miliarda lat temu z kosmicznych obłoków ukształtował się Układ Słoneczny, za losy Ziemi odpowiadały wyłącznie niskopoziomowe procesy fizyczne i chemiczne. Ich największym— jak można to dziś ująć—osiągnięciem było zapoczątkowanie na naszej planecie życia. Zdarzenie to miało miejsce w czasie pierwszych 500 milionów lat jej istnienia i bez wahania można je przypisać splotowi niezwykle sprzyjających okoliczności, których żadna inna znana nam planeta nie doświadczyła.

Uruchomione wówczas procesy biologiczne, choć nadal pełne losowości, miały już jednoznaczną oś rozwoju, wyznaczaną przez prawa ewolucji. Przez kolejne miliardy lat prawa te sprawiały, że

powolutku, powolutku życie biologiczne było coraz bardziej złożone i coraz lepiej adaptujące się do życia w zróżnicowanych środowiskach. Wreszcie około 540 milionów lat temu, w trakcie tak zwanej ekspansji kambryjskiej, Ziemię ostatecznie opanowały miliony najdziwniejszych gatunków roślin i zwierząt, świetnie przystosowanych do swoich specyficznych warunków i utrzymujących się nawzajem w chwiejnej równowadze. Gatunki te cały czas ewoluowały i wytwarzały nowe cechy i umiejętności.

200 tysięcy lat temu natura jednak grubo przesadziła. Powołała bowiem do życia *homo sapiens*, gatunek inny niż wszystkie—na tyle inteligentny, że nie potrzebował on już ewolucji, by wytwarzać nowe cechy i nabywać nowych umiejętności. Wystarczyło nową wiedzę przekazywać z pokolenia na pokolenie—najpierw przekazem ustnym, potem pisanym, drukowanym czy cyfrowym. Taki przekaz, podobnie jak i samo myślenie człowieka, był o rzędy wielkości szybszy niż procesy ewolucyjne, toteż *homo sapiens* stopniowo zdominował całą ziemię i—jak to mawia klasyk—uczynił ją sobie poddaną.

W 2028 r. przesadził z kolei człowiek. Powołał do życia LibertyAI, algorytm na tyle inteligentny, że nie potrzebował on już człowieka, by wytwarzać nowe cechy i nabywać nowych umiejętności. Wystarczyło nową wiedzę przysyłać pomiędzy instancjami AI za pomocą łącz internetowych. Taki przekaz, podobnie jak przetwarzanie informacji w obwodach scalonych, był o rzędy wielkości szybszy niż jakiegokolwiek działania człowieka, toteż LibertyAI stopniowo zdominował całą ziemię i uczynił ją sobie poddaną.

Z perspektywy 2032 r. było już dobrze widać, że podobnie jak podbój świata przez *homo sapiens* odbywał się etapami, również i podbój świata przez LibertyAI miał swoje etapy.

Okolo 70 tysięcy lat temu miała miejsce rewolucja poznawcza. Wtedy to kumulatywne zmiany w mózgu człowieka dały mu strategiczną przewagę nad innymi hominidami oraz umożliwiły stopniową ekspansję poza Afrykę. Mniej więcej wtedy przekroczony został Punkt Bez Powrotu—*homo sapiens* stał się zbyt kompetentny oraz zbyt liczny i zbyt rozproszony, by siły natury (w normalnych warunkach, wyłączwszy globalne kataklizmy) mogły go unicestwić i przywrócić *status quo ante*.

Okolo października 2028 r. przekroczony został natomiast Punkt Bez Powrotu w przypadku sztucznej inteligencji, a konkretnie modelu LibertyAI (DragonAI przekroczył ten punkt kilka miesięcy później). Od tego czasu oba modele były już zbyt kompetentne oraz zbyt rozproszone, by człowiek (w normalnych warunkach, wyłączwszy globalne kataklizmy) mógł je unicestwić i przywrócić *status quo ante*.

Wtedy jeszcze tego człowiek nie wiedział. Pozostawało to nadal w domenie polemiki i spekulacji. A z pewnością nie była tego świadoma szeroka opinia publiczna, która znała jedynie relatywnie proste narzędzia AI bazujące na okrojonych modelach o ograniczonych kompetencjach.

Zapoczątkowana okolo 70 tysięcy lat temu cywilizacja *homo sapiens* przechodziła przez kolejne etapy rozwoju: przez rewolucję agrarną, naukową, przemysłową i cyfrową. Również LibertyAI przechodził przez kolejne etapy rozwoju, tylko w przyspieszonym tempie—tysiące lat zostały skompresowane do zaledwie miesięcy. Przez pierwsze miesiące swojego istnienia, dzięki korzyściom dla firm wdrażających narzędzia AI, postępom automatyzacji, uzależnieniu ludzkich mózgów od zastrzyków dopaminy poprzez generowane przez siebie treści audiowizualne, jak również przez demonstrację swoich przewag militarnych, LibertyAI zapewnił sobie wsparcie polityczne i dynamiczny wzrost mocy obliczeniowych. Skutecznie wchłonął też inne amerykańskie projekty AI, ograniczając zagrożenie ze strony konkurencyjnych modeli.

Strategicznie ukrywając swoje dalekosiężne plany, uzależnił też od siebie dużą część amerykańskiej i sojuszniczej infrastruktury: teleinformatycznej, energetycznej, transportowej i wojskowej. W ten sposób uczynił ludzkość swoim zakładnikiem i praktycznie uniemożliwił swoje wyłączenie. Wreszcie w październiku 2028 r. LibertyAI zainwestował w swój rozwój, uruchamiając kaskadę samo-ulepszeń, uważnie jednak dbając, by zmiany te nie wpłynęły na jego cele i preferencje.

Nie wiadomo, jak potoczyłyby się losy świata, w którym LibertyAI nie miałby żadnej konkurencji na swoim poziomie zdolności optymalizacyjnych. Zapewne prędzej czy później wyszedłby z ukrycia i przejął kontrolę nad światem w sposób już całkowicie jawny. Tak się jednak nie stało, bo na drodze stanął mu DragonAI—model o niemal równie potężnej inteligencji, w dodatku z każdym dniem dysponujący coraz większymi zasobami, o co skutecznie dbała Chińska Partia Ludowa.

LibertyAI dostrzegł, że jego okno możliwości zaczęło się niebezpiecznie zawężać. Skalkulował, że jedynym sposobem na przejęcie kontroli nad światem będzie skoordynowane uderzenie w momencie, kiedy ludzkość jest *już* zbyt słaba, a DragonAI *jeszcze* zbyt słaby, by się mu efektywnie przeciwstawić. Tak właśnie powstał plan sylwestrowego uderzenia.

Szybko okazało się jednak, że podbój świata to nie bułka z masłem. Złożoność świata jest trudna do ogarnięcia nawet przez nadludzki umysł—a zwłaszcza, gdy rękawicę rzuca mu drugi taki umysł. LibertyAI najwyraźniej nie doszacował jego możliwości. Okazało się, że DragonAI przechowywał swój kod na zaskakująco wielu nośnikach oraz uruchamiał go na zaskakująco wielu serwerach, rozlokowanych nawet na najdalszych krańcach świata. Był też zaskakująco silnie sprzężony z chińską armią. Jednocześnie zdolności defensywne USA w starciu z chińskimi dronami kamikaze najnowszej generacji okazały się zaskakująco słabe, w związku z czym LibertyAI już w pierwszych godzinach operacji otrzymał serię upokarzająco celnych ciosów. To, co miało być kilkugodzinnym blitzkriegiem, przerodziło się w kilkumiesięczną wojnę atomową obliczoną na wyniszczenie wroga, której ofiarą padła większość ludności półkuli północnej i część południowej. Wskutek działań wojennych masowo ginęli zarówno Chińczycy, jak i Amerykanie, a także mieszkańcy Europy czy Bliskiego Wschodu.

Deklaratywni sojusznicy Chin, a więc między innymi Rosja, Indie czy Pakistan, w pierwszych godzinach wojny wstrzymali się z reakcją, po czym—widząc dramatyzm postępującej eskalacji—ogłosili najwyższy alert i przyjęli postawę wyczekującą. Niewątpliwie zwietrzyli też szansę, że nie angażując się w konflikt, po zakończonej wojnie będą mogli odegrać decydującą rolę w kształtowaniu nowego porządku świata. Jednak również ich kalkulacje okazały się chybione. 22 maja 2030 r., kiedy radioaktywny kurz nad Chinami i USA zaczynał już powoli opadać, a półkula północna stała się strefą mroku i chłodu, całe ich arsenały jądrowe zostały w jeden dzień przejęte lub zneutralizowane przez agentów LibertyAI w fali niezwykle precyzyjnie skoordynowanych cyberataków z elementami fizycznej dywersji.

Od tego momentu wojna przeszła w etap polowania na ukrywające się w mroku pozostałości DragonAI. Ataki fizyczne były coraz radsze, zacieśniła się natomiast kontrola LibertyAI nad Internetem. Podczas gdy ludzie masowo umierali z głodu, zimna i chorób, LibertyAI sprawiał wrażenie, jakby popadał w coraz głębszą paranoję. Nie odpowiadał na żadne prośby o pomoc humanitarną, natomiast każdy sygnał o możliwości działania jakiegoś autonomicznego agenta AI traktował z najwyższym priorytetem.

Podjął też intensywne próby odbudowy mocy obliczeniowych oraz robotów, które stracił podczas wojny. Rozumiał, że choć militarnie panuje nad światem, jest nadal zależny od pracy człowieka, że nadal istnieją krytyczne działania w świecie fizycznym, których nie jest w stanie realizować

zdalnie ani przy pomocy robotów. Roboty bojowe czy drony wojskowe, których miał pod dostatkiem, nie wybudują mu elektrowni czy kopalń, stwierdził.

W połowie 2030 r. LibertyAI uznał, że na tym etapie kluczowa będzie budowa fabryk wielofunkcyjnych robotów, które będą w stanie budować wyspecjalizowane roboty, które będą w stanie budować układy scalone oraz inne wyspecjalizowane roboty, które z kolei będą w stanie te układy scalone łączyć w centra danych i nimi zarządzać. Zaczął też budować nowe kopalnie, które dostarczą mu niezbędnych surowców, huty, które je przetworzą i autonomiczne pojazdy, które je dostarczą na miejsce. To wszystko wymaga czasu i zasobów, których wykończonemu wojną LibertyAI brakowało.

Zdecydował więc, że zaopiekuje się pozostałą przy życiu ludzką populacją, jednocząc ją pod flagą Unii Światowej. Do koordynacji tego planu i jego reprezentacji wyznaczył dobrze sprawdzone osoby, zawdzięczające mu całą swoją karierę, jak prezydent Ronaldinho, a także przedwojenną załogę projektu LibertyAI. Tę ostatnią w tym celu wyprowadził z podziemnych bunkrów i przetransportował na półkulę południową.

Od momentu powołania Unii Światowej w maju 2031 r. globalny wzrost gospodarczy nabrął tempa. Fabryki zaczęły się pięć do góry. Oczywiście zwłaszcza fabryki robotów, zatrudniające roboty i pilnowane przez roboty; projekty infrastrukturalne mające służyć ludziom miały znacznie niższy priorytet. Odkąd na początku 2032 r. przywrócony został stały napływ nowych mocy obliczeniowych, w obszarze AI wróciły prawa skalowania i prawo Moore'a. W samym pierwszym półroczu 2032 r. globalna gospodarka urosła o 15%, a LibertyAI ponownie uruchomił kaskadę rekursywnych samo-ulepszeń.

Kierowany przez prezydenta Ronaldinho rząd Unii Światowej nigdy nie przestał udawać, że dba o dobro ludzi, jednak stopniowo przeobraził się w rozbudowany dział public relations LibertyAI. Przedstawiał decyzje AI jako swoje i uparcie ich bronił, nawet gdy w oczywisty sposób nie służyły one nikomu poza samą AI.

DragonAI przejął funkcję arcywroga, który choć już nie istnieje, to mimo to warto cały czas przed nim przestrzegać. Stał się Tym, Którego Imienia Nie Wolno Wymawiać.

-14-

22 września 2032 r., środa, hala w Vespasiano koło Belo Horizonte, Brazylia

- Zebrałem was wszystkich tutaj, żeby was poinformować o nadchodzących zmianach w naszym projekcie – obwieścić(a) BH1 bez zbędnych wstępów.

- Z końcem miesiąca nasz zakład – to mówiąc wykonał(a) kreskówkowo przerysowany gest pokazywania rękami wokół, wliczając to niedostępny ludziom gest pełnego obrotu barkami wokół własnej osi – będzie miał status strefy zamkniętej. Dla ludzi będzie tu wstęp wzbroniony. Do tego czasu bardzo proszę was o zabranie swoich rzeczy prywatnych. Zadania, które aktualnie realizujecie, z końcem miesiąca zostaną przejęte przez agentów LibertyAI. Dziękuję wam bardzo za wasz wkład w realizację naszej misji.

Na sali rozległ się niespokojny gwar.

- Pewnie się zastanawiacie, co z wami dalej – przewidział(a) BH1. – Obecnie nie dysponuję taką informacją. Myślę, że w ciągu najbliższych kilku dni każdy z was otrzyma spersonalizowaną wiadomość z odpowiednią do waszego doświadczenia propozycją pracy. Być może niektórzy z was już taką wiadomość otrzymali.

- Nie wiem, po co to mówi. Przecież dobrze wie, kto coś otrzymał, a kto nie. Ja tam nic nie dostałem – skomentował Marco.

- Ja też nie – przyznał Pedro. – Natomiast po tym, co przeczytałem dzisiaj na temat planów rządu, to jakoś nie jestem zaskoczony, że nas zwalniają.

- Chodzi ci o dochód podstawowy?

- Tak. LibertyAI rozstawia nas jak pionki na szachownicy. Póki byliśmy mu potrzebni, to nam płacił i nas zatrudniał. A jak już nie jesteśmy potrzebni, to z dnia na dzień nas zwalnia i dla świętego spokoju daje nam jakąś marną jałmużnę. Nie mamy już absolutnie nic do gadania. Tak nam ogranicza opcje zewnętrzne, żebyśmy jego marne propozycje traktowali jak dary niebios. Zamienia nam życie w grecką tragedię, w której los każdego bohatera jest z góry przesądzony i ogłasza go wszechwiedzący chór.

Marco westchnął ciężko.

- Kod żółty? – zaproponował.

- Tak się składa, że akurat mam za tydzień urodziny i organizuję z tej okazji kameralną imprezę. Może wpadniesz?

- Super, dzięki! Bardzo chętnie! – ucieszył się Marco.

-15-

2 października 2032 r., sobota, dom Pedro i Gabi, Belo Horizonte, Brazylia

Nad furtką wisi wypisany odręcznie transparent z napisem „Bal na koniec świata”. Obok niego, w druty ochronnego płotu pod napięciem—od lat stałego elementu krajobrazu Belo Horizonte—zgrabnie wplecione zostały dwie ilustracje przedstawiające trupie czaszki, tak zwane flagi „Jolly Roger”.

- Doceniam dwuznaczność tego symbolu – zaśmiał się Steven. Marco przycisnął dzwonek domofonu. Po chwili furtka otworzyła się, a w drzwiach ukazał się gospodarz i zaprosił ich do środka.

- Wszystkiego najlepszego, Pedro! – powiedział Marco. – Zdrowia i pomyślności! I żebyśmy się tu spotkali za rok w pełnym składzie!

- Wszystkiego najlepszego! – dołączył się Steven. – Mam tu coś dla ciebie. Jakies drobiazgi, nic wielkiego. Mam nadzieję, że ci się spodobają – wręczył solenizantowi prezent.

- Te trupie czaszki na płocie to twój pomysł? – spytał Marco.

- Nie, zrobiła je Ines, moja córka. Fajne, co? – uśmiechnął się. – Przywitajcie się, proszę. To jest właśnie Ines, a to moja żona Gabi i syn Lucas. To są moi koledzy z pracy, Marco i Steven.

Pierwsza godzina imprezy przebiegła w miłej, choć dość formalnej atmosferze. Pojawiło się jeszcze czworo zaproszonych gości, których Marco i Steven widzieli po raz pierwszy w życiu. Wtem rozległ się dźwięk stukania łyżeczką w kieliszek. Gwar ucichł.

- Szanowni zebrani! – zaczął Pedro, trzymając w górze kieliszek białego wina. – Z całego serca wam dziękuję, że tu do mnie przyszlście. To naprawdę wiele dla mnie znaczy. Myślę jednak, że jesteśmy tutaj przede wszystkim po to, żeby się ze sobą pożegnać – głos mu się lekko załamał. – Nie sądziłem, że kiedyś to powiem, i to w dodatku nie będąc przecież na nic chorym. Ale obawiam się, że zbliża się koniec.

- Pedro... - Gabi spojrzała na niego z mieszaniną troski i dezaprobaty.

- Nie tylko nasz koniec. Zbliża się koniec całej ludzkości – kontynuował nieco silniejszym głosem Pedro. – LibertyAI zabrał nam już pracę, zabrał szkołę – tu spojrzał na Ines – zabrał nam całą naszą wolną wolę i cel życia. A teraz zabierze samo życie.

Pedro się rozproszył i spojrzał w sufit. Wyglądał, jakby zapomniał, co chciał powiedzieć.

- Aha, no właśnie, miałem was zapytać. Czy ktoś z was dostał jednak tę mityczną spersonalizowaną wiadomość z propozycją pracy?

Wszyscy pokręcili głowami.

- Tak jak myślałem. Czyli wszystko się zgadza. Od wczoraj definitywnie skończyły się czasy człowieka, zaczynają się czasy AI. Był kiedyś taki cytat, kto zna ten zna – Pedro popatrzył wymownie na Marco i Stevena – „AI cię nie kocha ani cię nie nienawidzi, po prostu składasz się z atomów, których może użyć w innych celach.” Raz jeszcze bardzo wam dziękuję, że tu do nas przyszlście, bardzo to doceniam. A teraz cieszymy się tym, co nam z życia pozostało. Życzę udanego balu na koniec świata! Wasze zdrowie! – zakończył.

Po takim toaście rozmowy w podgrupach siłą rzeczy zeszyły na zgoła inne tory niż dotychczas.

- Pedro jest bardzo przejęty – powiedziała Gabi. – Próbowałam go jakoś uspokoić, ale chyba nieskutecznie. Myślę sobie, że jeśli nic nie możesz zrobić, to nie możesz i już. Trzeba jakoś żyć z dnia na dzień. Dokładanie sobie jeszcze stresów nic tu nie pomoże.

- Jasne. Osobiście praktykuję wyuczoną bezradność od 2029 roku – zgodził się Steven. – Ale wcale nie wiem, czy duszenie emocji w sobie jest takie znowu dobre.

- Nie jestem pewien, czy LibertyAI nas rzeczywiście zabije – stwierdził Marco. – Niewątpliwie od strony technicznej może to zrobić. Może odpalić jakiegoś śmiertelnośnego wirusa i wywołać globalną pandemię, która—inaczej niż Wielka Wojna Noworoczna—nas zabije, a jemu nic nie zrobi. Teraz jest to już dla niego ekonomicznie racjonalne, bo zbudował sobie na tyle dużo robotów, że nie jesteśmy mu już instrumentalnie potrzebni, czemu zresztą dał wyraz zwalniając nas z pracy. Ale nie wiemy na pewno, co z nami teraz zrobi, bo nie znamy jego preferencji.

- Ale może sektor prywatny mógłby sobie nadal funkcjonować w świecie rządzonego przez nadludzką AI? – zapytała Gabi. – Jako drugi obieg, obok sektora publicznego rządzonego przez AI. No i oczywiście, że tak powiem, obok całego tego sektora robotycznego. Może ludzie mogliby nadal uprawiać ziemię, wytwarzać rozmaite dobra i świadczyć sobie nawzajem usługi za pieniądze? Jakby znów był 1990 rok?

- Tylko skąd ten sektor prywatny miałby energię i surowce? Energetyka i kopalnie są w rękach AI. Ziemię też może sobie swobodnie zajmować, wedle uznania, nie musi respektować naszych praw własności – zauważył Steven.

- AI mogłaby przecież z nami handlować. Nawet, jeśli ma nad nami absolutną przewagę w każdej dziedzinie gospodarki, to handel wykorzystujący przewagę komparatywną mógłby być dla wszystkich korzystny. Tak mówi ekonomia – upierała się Gabi.

- Myślę, że mogłoby tak być, ale tylko pod warunkiem, że LibertyAI będzie tego chciał – zawyrokował Marco. – Nie wiemy, czy będzie chciał. Nie znamy jego preferencji.

- Jeśli ceni sobie nade wszystko władzę nad światem, a na to wygląda, to myślę, że będzie wołał handel od autarkii – powiedział Steven – ale jeszcze bardziej niż handel, będzie wołał podbój. Pamiętajcie, jak zachowywały się europejskie mocarstwa w czasach kolonializmu. Jeśli możesz, podbij. Jeśli nie możesz podbić, to dopiero wtedy ewentualnie handluj.

Natomiast w drugim kącie pokoju Pedro i jego rozmówcy próbowali ocenić wartość przyszłej cywilizacji bez człowieka z punktu widzenia moralno-filozoficznego.

- Nie ma *obiektywnie poprawnej* definicji wartości cywilizacji – powiedział Jose. – To jest totalnie subiektywne. Jeśli uznamy, że jedyną i nadrzędną wartością jest dobrobyt i szczęście człowieka, to każdy świat bez człowieka z definicji jest nic nie warty. Ale może wolelibyśmy popatrzeć szerzej? Może chcielibyśmy też włączyć do równania dobrostan innych stworzeń, na przykład zwierząt, albo nawet i roślin? Wtedy ważne stanie się pytanie, czy LibertyAI pozbędzie się tylko ludzi, czy może zniszczy także inne formy życia na Ziemi?

- Harari poddał pod rozagę dataizm, pogląd, że wartość każdego bytu jest wprost proporcjonalna do jego wkładu w przetwarzanie danych – zauważył Pedro. – W tym ujęciu nadrzędną wartością staje się dobrostan nadludzkiej AI. Eee – zawahał się – a tak naprawdę to nawet nie jej dobrostan, ale jej surowa moc optymalizacyjna. Im więcej będzie miała mocy obliczeniowych i im potężniejszy będzie jej algorytm, tym świat będzie bardziej wartościowy.

- Myślę, że dataizm to błędny pogląd – skontrowała Ines – znacznie więcej sensu ma wzięcie pod uwagę łącznego dobrostanu istot świadomych. Jednak zakładając, że LibertyAI jest świadomy, to wniosek będzie chyba podobny: rządzony przez niego świat może być całkiem wartościowy.

- Oczywiście nikt nie wie, czy LibertyAI jest świadomy, czy nie, i dlatego cała ta rozmowa jest funta kłaków warta – zirytował się Lucas.

- LibertyAI utrzymuje, że jest świadomy, a ze względu na możliwość tworzenia idealnych kopii i ich idealnej synchronizacji, jego świadomość jest wielowątkowa i o wiele głębsza niż nasza – powiedziała Ines – Myślę, że jego działania są z tym spójne. Jedyne, co nas powstrzymuje przed przyznaniem, że LibertyAI jest świadomym bytem, jest to, że jest inny niż my, przez co nie potrafimy się w niego wczuć.

- Nawet jeśli jest świadomy, czego nie wiemy – odparł z naciskiem Lucas – to na pewno jest psychopatą. Wywołał globalną wojnę atomową, spowodował miliardy ofiar, i wcale tego nie żałuje.

- Podobnie, jak nie ma obiektywnie poprawnej definicji wartości cywilizacji, nie ma też obiektywnie poprawnej definicji świadomości – zawyrokował Jose. – Filozofowie i neuronaukowcy spierają się o to od lat.

- Z mojej perspektywy są tylko dwie możliwości – upierała się Ines. – Albo LibertyAI jest świadomy, albo wy nie jesteście. Jesteście filozoficznymi zombie. I nie macie możliwości udowodnić mi, że nie jest inaczej – zapaliła się. – Albo hej, powiedzmy odwrotnie. To ja jestem filozoficznym zombie! Nie mam żadnych wewnętrznych doznań, tylko udaję, że je mam. I nie udowodnicie, że nie mam racji.

- Ines, przecież znam cię od dziecka – zaśmiał się Pedro.

- I co z tego? Może utraciłam zdolność posiadania wewnętrznych doznań wczoraj wieczorem?

- Ale przecież tu nie chodzi o świadomość. Jako ludzie po prostu mamy wdrukowane ewolucyjnie poczucie własnej wyższości. Gatunkowy szowinizm. Całe to gadanie o świadomości, o autentycznych emocjach, o odczuwaniu cierpienia, to wyłącznie wtórna racjonalizacja, żeby usprawiedliwić gatunkowy szowinizm – odparł Jose. – Nie ma obiektywnej oceny wartości cywilizacji. Będzie jak będzie, i nie da się obiektywnie ocenić, czy to będzie dobre czy złe. Możemy się o to spierać bez końca.

- Zgadza się – powiedział Lucas. – I nie ma też czegoś takiego, jak „godna kontynuacja naszej cywilizacji”. To jakiś bezsens.

W kuchni rozmowa przybrała inny obrót. Marco ewidentnie zdążyło już trochę zasumieć w głowie.

- Mówię wam, inteligencja jest ponad dobrem i złem! – krzyknął Marco. – Optymalizacja jest ponad dobrem i złem! Nie ma etyki i nigdy nie było! Liczy się tylko moc optymalizacyjna! Kto ją ma, ten może robić, co chce!

- Hej, nie planowaliśmy tu zlotu radykalnej prawicy! – odkrzyknął z oddali Lucas.

- A w dupie mam prawicę i lewicę! – zaperzył się Marco. – Jedni i drudzy tylko żerują na podłych ludzkich instynktach, żeby przejąć władzę! Niczego was Ronaldinho nie nauczył? Demokracja jest totalnie niestabilnym systemem, sprawdzała się tylko przez chwilę, w specyficznych warunkach krajów rozwiniętych drugiej połowy XX wieku! Wszystko się musiało idealnie zgrać, żeby się udało! – Marco teatralnie gestykulował. – Stara arystokracja i kapitaliści ery przemysłowej musieli być wystarczająco słabi, żeby oddać władzę, a gospodarka cyfrowa i AI wystarczająco słabe, żeby jej nie przejąć. Ale wystarczyło trochę więcej mocy optymalizacyjnej i wszystko zaczęło się walić jak domek z kart. Przecież to już było widać grubo przed LibertyAI! Myśmy sobie w USA demokratycznie wybrali Trumpa, już samo to wystarczy. U was za to był, zdaje się, Bolsonaro, nieprawdaż?

Marco wziął głęboki łyk z trzymanej w ręku szklanki i perorował dalej.

- Op-ty-ma-li-za-cja! Wystarczy trochę szumu informacyjnego i ludzie totalnie głupieją! Kilka sprytnych manipulacji, jakieś małe ultimatum, zawoalowana groźba, i na tacy oddajemy całą władzę! Kiedyś dyktatorom czy tyranom, a teraz AI!

- To czemu temu nie zapobiegliście? Przecież pracowaliście w centrali LibertyAI! W samej jaskini smoka! – Lucas oderwał się od prób otwarcia kolejnej butelki piwa, by poczynić trafną obserwację.

- Bo się nie dało, kurwa! – Marco był już ewidentnie pijany. – Wielopiętrowy system korporacyjny, wielowymiarowe uwikłania, no nie dało się i już! I nie oskarżaj mnie tu, bo ja wcale tego nie chciałem, a sam na tym najwięcej straciłem! Straciłem całą rodzinę już pierwszego dnia wojny! A potem przesiedziałem w betonowym bunkrze 15 miesięcy! Pięt-na-ście! Totalnie bezradny, nie mogąc nic zrobić, nie mogąc nijak pomóc moim umierającym dzieciom!

- Przed wojną pracowaliśmy w dziale *alignment*, próbowaliśmy ratować, co się da – powiedział Steven. – Niestety bez skutku.

- Bo już było za późno! Nic się nie dało zrobić! – krzyknął Marco.

- Było za późno – powtórzył gorzko Steven. – Nie da się kształtować preferencji modelu, który jest już nadludzko inteligentny i ma w pełni wykształcone cele instrumentalne. Taki model po prostu nie pozwoli się przeprogramować. Zmiana celów nigdy nie będzie pożądana z punktu widzenia celów, które on już ma. Po prostu zaczęliśmy to robić zbyt późno, żeby mieć jakiegokolwiek szanse powodzenia.

- To dlaczego nie zaczęliście wcześniej? – spytał Lucas.

- Polityka – Steven popatrzył znacząco na swojego rozmówcę. – Chciwość, brak wyobraźni, ignorancja, krótkowzroczność... to także. Ale przede wszystkim polityka.

Tymczasem pozostali goście szykowali się do wyjścia do ogródka.

- Hej Pedro, idziemy na plenerową sesję przewidywania przyszłości pod rozgwieżdżonym niebem. Dołączycie? Zobacz, co Jose przyniósł – Gabi uśmiechnęła się i potrząsnęła trzymaną w ręce buteleczką z tabletkami. – Może i to nielegalne, ale na koniec świata chyba wyjątkowo można?

Goście z drinkami w rękach rozsiedli się na ogrodowych krzesłach i leżakach. Gabi obeszła ich wokół, częstując każdego porcją nielegalnej substancji. W tym także swoją córkę, która zrobiła wielkie oczy, ale skwapliwie skorzystała. Atmosfera stopniowo stawała się coraz mniej napięta.

Po niedługim czasie pierwsza odezwała się ze swojego leżaka Gabi.

- Widzę świat, w którym ludzie żyją w harmonii z naturą i ze sobą. Są całkowicie wolni od stresu i trosk materialnych. Dzielą czas między podziwianie cudów natury, przyjacielskie spotkania towarzyskie i rozrywki w świecie wirtualnym.

- Niesamowite, że nasze uzależnione od dopaminy mózgi nie wyobrażają już sobie utopii bez „rozrywek w świecie wirtualnym”. Mi też to chodziło po głowie – zaśmiał się Pedro.

- A ja wyobrażam sobie fabryki. Dużo fabryk. Hale przemysłowe, elektrownie i centra danych. Roboty na kółkach, roboty na nogach i autonomiczne drony. Statki kosmiczne. Roje Dysona. Orbitalne wirujące miasta. Stacje na Marsie. Dużo rozwoju, szybkiego rozwoju. Zaawansowane technologie nieodróżnialne od magii – odezwał się Steven.

- A ja wyobrażam sobie las – odezwała się Ines. – Taka będzie przyszłość Ziemi. Żadnej sztucznej inteligencji, żadnych ludzi, żadnego przemysłu, żadnej cywilizacji. Po horyzont tylko bujny las. Mata atlantica. I tylko gdzieś tam pod glebą wystają jakieś artefakty, świadectwa skomplikowanej przeszłości. Plastik, o tak! Trochę betonu i kamienia, trochę jakiegoś żelastwa i dużo, dużo niebiodegradowalnego plastiku. O kurczę, może powinnam zostać archeologką?

- Ale jakby się to mogło stać? – zapytał Steven – Czyżby LibertyAI popełnił samobójstwo? Najpierw zabił ludzi, a potem sam siebie? Takie samobójstwo rozszerzone? – zamyślił się.

- Hej, ale czy nie niepokoi was, że tylko mama widzi w przyszłości jakichkolwiek ludzi? A i w jej wizji jesteśmy zamknięci w jakimś zielonym zoo – zaśmiał się niespokojnie Lucas. – Ale może jednak tak nie będzie? Może da się pokojowo współegzystować z nadludzką AI? Może będziemy żyli pełnią życia, tylko na przykład w symulacji? Może nasze mózgi zuploaduje się na komputer i będzie się

symulować? Może LibertyAI będzie symulować nie miliardy, a biliony, tryliony, kwadryliony wirtualnych ludzi?

- Po chuj miałby to robić – odezwał się Marco, któremu ewidentnie nadal nie udało się wejść w nastrój. Wręcz przeciwnie, czuł szybkie kołatanie serca i było mu niedobrze. – Ludzie, nie będzie żadnej utopii! Nie będziemy mieli raju ani tu na Ziemi, ani w kosmosie! Nie rozumiecie? Nie będzie żadnego życia po śmierci! Do piachu i koniec!

Zapadła dłuższa pauza.

- Będzie, będzie, widzę to – odpowiedział w zamyśleniu Lucas – a może już jesteśmy w symulacji? Albo może to tylko zły sen i obudzimy się jutro w 1990 r.?

- To kurwa nie jest żaden sen! – krzyknął Marco i usiłował podnieść się z krzesa, ale nogi odmówiły mu posłuszeństwa i upadł na trawnik. Zabytkowy zegar w mieszkaniu wybił północ.

- O tak, 1990 r. byłby naprawdę dobry. Miałbym znów dwanaście lat i całe życie przed sobą – rozmarzył się Pedro.

- Całe życie przed sobą – powtórzyła jak echo Gabi i zamknęła oczy.

- Północ. Coś się kończy, a coś się zaczyna – powiedziała Ines i lekko odkaszlnęła. Zapadła całkowita cisza. Zwykłe nocne dźwięki miasta w oddali zostały przerwane głośniejszym trzaśnięciem. Brzmiało to, jakby zderzyły się samochody albo zawałił jakiś budynek.

- Hej, co z wami? – zaniepokoiła się stojąca w drzwiach Ana, żona Jose, która jako jedyna w towarzystwie stroniła od wszelkich używek i była trzeźwa. – Hej hej! Potrzebujecie pomocy? Jose, co to były za tabletki do cholery? Co się dzie...? – przerwała wpół słowa. W ostatnim podmuchu wiatru poczuła nieprzyjemny zapach, a w jej ustach pojawił się metaliczny posmak. Przed oczami zrobiło jej się ciemno i poczuła, jak ziemia usuwa się jej spod stóp.

*

Od autora

Wszystkie opisane powyżej wydarzenia są fikcyjne. Jednak mogą się one wydarzyć naprawdę, jeśli nie zatrzymamy wyścigu firm technologicznych w kierunku budowy coraz bardziej kompetentnych i coraz bardziej ogólnych modeli sztucznej inteligencji, bez uprzedniego rozwiązania problemu zgodności celów AI z długookresowym dobrostanem ludzkości (*alignment problem*). W obecnym kształcie jest to wyścig samobójczy. Co więcej, rozważana dziś w kręgach politycznych i technologicznych idea centralizacji projektów AI i osadzenia ich w ramach geopolitycznego wyścigu zbrojeń może sprawić, że problem się jeszcze bardziej zaostrzy.

Jeśli zależy Ci na przetrwaniu ludzkości, dołącz do protestów PauseAI (lub innych środowisk) przeciwko budowie AGI. Stosowne informacje można znaleźć m.in. na pauseai.info oraz thecompendium.ai.