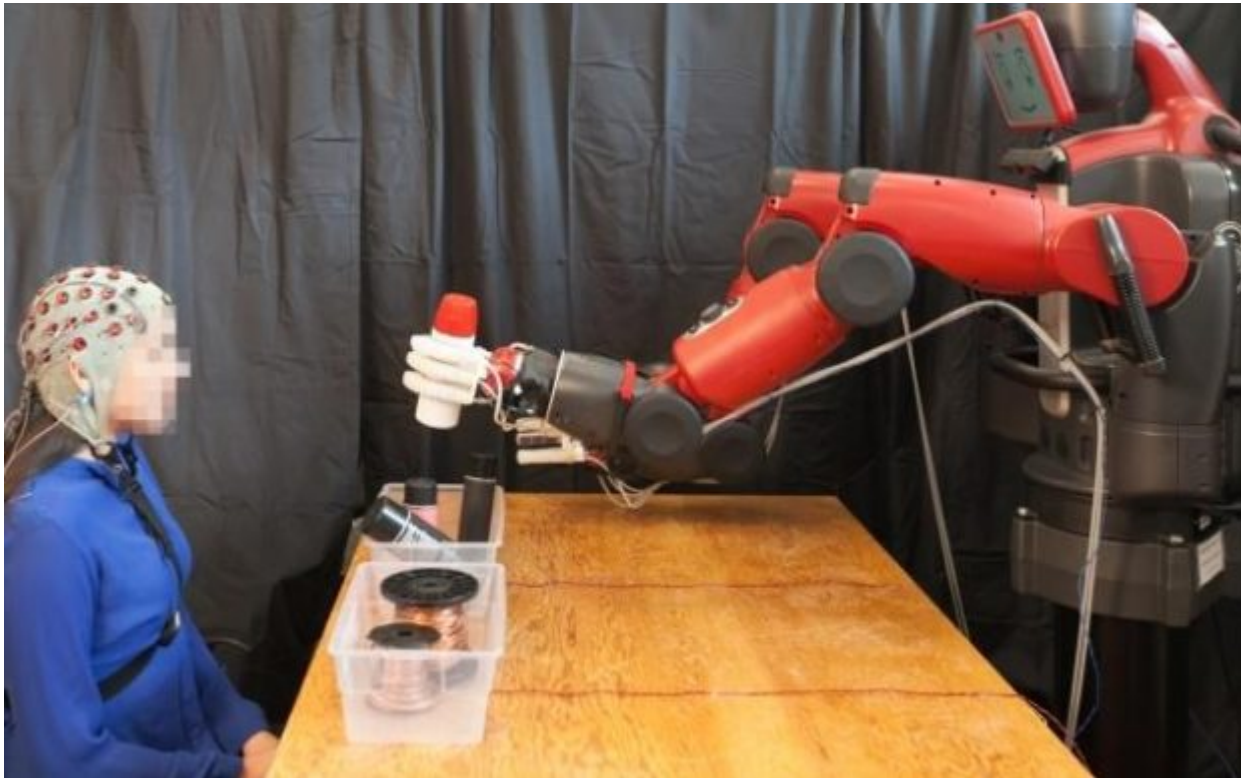


Ρομπότ που ελέγχονται με τον εγκέφαλο

/ [Επιστήμες, Τέχνες & Πολιτισμός](#)

Image not found or type unknown



Jason Dorfman/MIT CSAIL

Χρησιμοποιώντας δεδομένα από ένα μόνιτορ ηλεκτροεγκεφαλογραφίας (EEG) που καταγράφει την εγκεφαλική δραστηριότητα, το σύστημα μπορεί να «καταλάβει» εάν κάποιος αντιλαμβάνεται ένα λάθος στη δραστηριότητα ενός ρομπότ ενώ πραγματοποιεί μια εργασία.

Ομάδα ερευνητών του CSAIL (Computer Science and Artificial Intelligence Laboratory) του MIT και του Boston University δημιούργησε ένα σύστημα αλληλεπίδρασης ανθρώπων- ρομπότ που επιτρέπει στους ανθρώπους να διορθώνουν άμεσα τα λάθη των ρομπότ κατά τη λειτουργία τους, χρησιμοποιώντας απλά τους εγκεφάλους τους.

Χρησιμοποιώντας δεδομένα από ένα μόνιτορ ηλεκτροεγκεφαλογραφίας (EEG) που καταγράφει την εγκεφαλική δραστηριότητα, το σύστημα μπορεί να «καταλάβει» εάν κάποιος αντιλαμβάνεται ένα λάθος στη δραστηριότητα ενός ρομπότ ενώ πραγματοποιεί μια εργασία. Οι αλγόριθμοι machine learning επιτρέπουν στο σύστημα να ταξινομεί τα εγκεφαλικά κύματα σε κλίμακες 10-30 μιλισεκόντ.

Αν και το σύστημα αυτή τη στιγμή διαχειρίζεται απλές περιπτώσεις δυαδικών επιλογών, οι ερευνητές υποστηρίζουν πως η έρευνα υποδεικνύει ότι κάποια στιγμή θα είναι δυνατός ο έλεγχος των ρομπότ με πολύ πιο ενστικτώδεις τρόπους.

«Φανταστείτε να είστε ικανοί να πείτε άμεσα σε ένα ρομπότ να προβεί σε μια συγκεκριμένη ενέργεια, χωρίς να χρειάζεται να πληκτρολογήσετε εντολή, να πατήσετε ένα πλήκτρο ή ακόμα και να πείτε μια λέξη» λέει η διευθύντρια του CSAIL, Ντανιέλα Ρας. «Μία ομαλή προσέγγιση σαν αυτή θα βελτιώνει τις δυνατότητές μας στην επίβλεψη εργοστασιακών ρομπότ, αυτόνομων οχημάτων και άλλων τεχνολογιών που δεν έχουν εφεύρει ακόμα».

Στο πλαίσιο της συγκεκριμένης έρευνας, η ομάδα χρησιμοποίησε ένα ανθρωποειδές ρομπότ με το όνομα «Baxter» από τη Rethink Robotics, εταιρεία του πρώην διευθυντή και συνιδρυτή της iRobot, Ρότνεϊ Μπρουκς.

Προηγούμενες έρευνες πάνω σε ελεγχόμενα μέσω EEG ρομπότ προϋπέθεταν οι άνθρωποι να «σκέφτονται» με προκαθορισμένο τρόπο, τον οποίο οι υπολογιστές μπορούν να αναγνωρίσουν- για παράδειγμα, ένας χειριστής θα έπρεπε να κοιτά σε μια από δύο φωτεινές οθόνες η καθεμία από τις οποίες αντιπροσωπεύει άλλη εντολή για το ρομπότ. Η ομάδα της Ρας ήθελε να κάνει την εμπειρία πιο φυσική, επικεντρώνοντας σε εγκεφαλικά σήματα υπό την ονομασία ErrP (error-related potentials) που παράγονται όποτε ο εγκέφαλος ανιχνεύει ένα λάθος. Καθώς το ρομπότ υποδεικνύει σε ποια επιλογή σκοπεύει να προβεί, το σύστημα χρησιμοποιεί ErrP για να διαπιστώσει εάν ο άνθρωπος συμφωνεί ή όχι με την επιλογή αυτή.

«Καθώς βλέπεις το ρομπότ, το μόνο που έχεις να κάνεις είναι να συμφωνήσεις ή να διαφωνήσεις διανοητικά με αυτό που κάνεις» λέει η Ρας. «Δεν χρειάζεται να εκπαιδευτείς να σκέφτεσαι με συγκεκριμένο τρόπο- η μηχανή προσαρμόζεται σε εσένα, όχι το αντίθετο».

Τα σήματα ErrP είναι πολύ αχνά, κάτι που σημαίνει ότι το σύστημα πρέπει να είναι ρυθμισμένο στην εντέλεια, έτσι ώστε να μπορεί τόσο να ταξινομεί το σήμα όσο και να το ενσωματώνει στο «κύκλωμα» για τον άνθρωπο χειριστή. Πέρα από την παρακολούθηση των αρχικών ErrP, η ομάδα επίσης επιδίωξε τον εντοπισμό «δευτερευόντων» λαθών, που γίνονται όταν το σύστημα δεν αντιλαμβάνεται την αρχική διόρθωση από πλευράς του ανθρώπου.

Αν και το σύστημα δεν μπορεί ακόμα να αναγνωρίζει δευτερεύοντα λάθη σε πραγματικό χρόνο, αναμένεται κάποια στιγμή να είναι δυνατόν να φτάσει σε επίπεδο ακριβείας άνω του 90%.

Πηγή: naftemporiki.gr