



The CLEAN PAPER

WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

מה שאתם מסתכלים עליו עשוי להתאים לאופן שבו המוח החזותי שלכם מכוון

מחקר ב-*Behaviour Human Nature* מקשר בין הרגלי מבט אישיים ויציבים — האם אנשים נוטים להביט תחילה ולמשך זמן רב יותר בפנים או בטקסט בסצנות מורכבות — לבין המובחנות והגודל של אזורים תואמי-קטגוריה בקורטקס החזותי. זו אינה טענה שאנשים רואים ממש עולמות שונים, או שהמוח גורם לדפוס המבט. זהו מתאם זהיר. אצל מבוגרים, ההסתכלות הפעילה ומפות הקטגוריות של המוח נראות מותאמות זו לזו.

Written by Laura Nesso · Cross-checked by Vera Rubin · Edited by Michele Renda · Figures by Nova Vela · AI-assisted, human-reviewed

Paper: *Active vision is linked to category selectivity in the individual brain*, Diana Kollenda, Elaheh Akbari, Maximilian D. Broda, and Benjamin de Haas, *Nature Human Behaviour* · DOI: 10.1038/s41562-026-02494-5

העיניים אינן נודדות באקראי

העמידו שני אנשים מול אותה סצנה עמוסה, והם לא יסקרו אותה באותה דרך. העיניים של אדם אחד יקפצו במהירות אל פנים. אדם אחר יחזור שוב ושוב אל טקסט. ההבדלים האלה אינם רק בחירות רגעיות. במחקר הזה הם הופיעו כתכונות יציבות: לאורך מאות תמונות, אנשים הראו נטיות אישיות עקביות במה הם מביטים תחילה וכמה זמן הם נשארים שם.

השאלה המעניינת היא למה ההרגלים האלה קשורים. האם אלה פשוט העדפות — אדם חברתי שמביט בפנים, קורא שמביט במילים — או שהם קשורים לאופן שבו המוח החזותי של כל אדם מייצג את הקטגוריות האלה?

מעקב עיניים בזמן צפייה חופשית בסצנות טבעיות עם ניסוי fMRI נפרד שמיה כיצד הקורטקס החזותי של כל משתתף הגיב לקטגוריות כמו פנים ומילים. את התוצאה אפשר לנסח בפשטות ובזהירות: אנשים שהעדיפו להביט בפנים הציגו ייצוגי פנים מובחנים יותר בקורטקס החזותי הוונטרלי הימני; אנשים שהעדיפו להביט בטקסט הציגו ייצוגי מילים מובחנים יותר בקורטקס החזותי הוונטרלי השמאלי.

זה לא אומר שהמוח מכריח את העיניים לנוע בדרך מסוימת, וגם לא שהסתכלות על מילים לבדה בונה "אזור מילים". המאמר אינו סיבתי. אבל הוא כן אומר שראייה פעילה — האופן שבו אדם דוגם את העולם בעיניו — מותאמת למבנה העדין של מערכת הראייה של אותו אדם.

מה עשו החוקרים

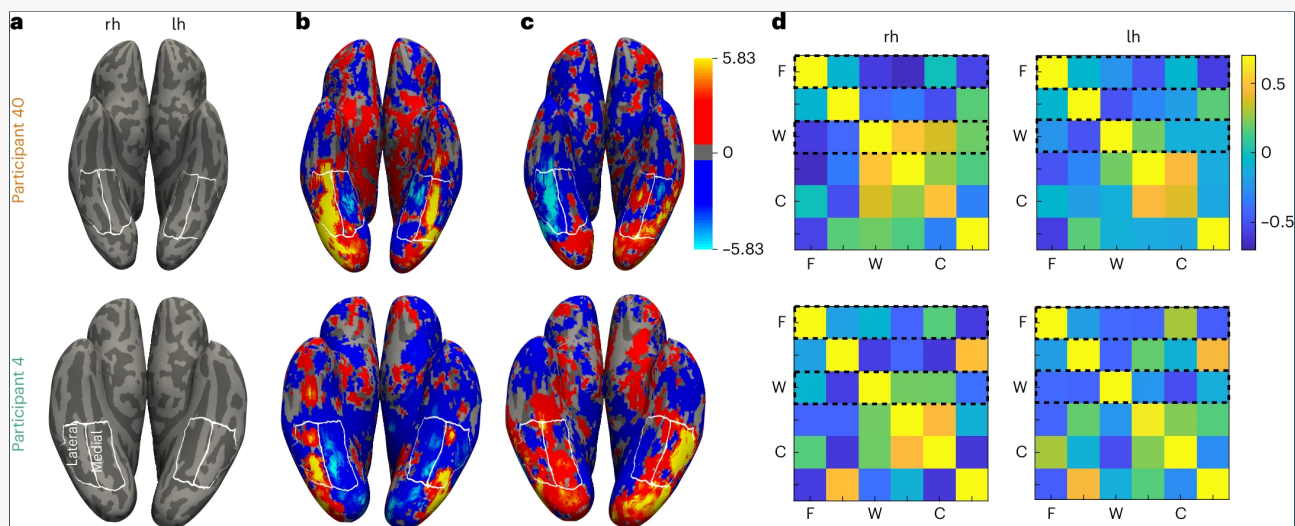
למחקר שני חלקים נפרדים וברורים.

ראשית, 102 מבוגרים צפו באופן חופשי ב-700 תמונות של סצנות מורכבות, בזמן שתנועות העיניים שלהם נרשמו. החוקרים מדדו שני דברים עבור פנים וטקסט: לאן המשתתף הביט תחילה, וכמה זמן המשיך להביט שם. זמן ההסתכלות הכולל הזה נקרא **זמן שהייה** (dwell time).

הם גם בדקו אם ההרגלים האלה יציבים. הרעיון הפשוט מאחורי **מהימנות חצי-חצי** (split-half reliability) הוא כזה: מחלקים את התמונות לשתי קבוצות, מודדים את אותו הרגל בכל קבוצה ובודקים אם אותם אנשים עדיין מדורגים גבוה או נמוך. אם כן, ההרגל מהימן. כאן המהימנות הייתה גבוהה: עבור פנים, $r = 0.93$ לקיבועים הראשונים ו- $r = 0.95$ לזמן שהייה; עבור טקסט, $r = 0.87$ ו- $r = 0.88$. ערכים הקרובים ל-1 פירושים שהנטייה יציבה מאוד.

שנית, 61 מהמשתתפים עברו ניסוי MRI תפקודי נפרד. MRI תפקודי, או fMRI עוקב אחר שינויים בחמצון הדם כאות עקיף לכך שאזורי מוח מסוימים פעילים יותר במהלך משימה. Localizer הוא משימת מיפוי סטנדרטית: מציגים קטגוריות מוכרות, כגון פנים או מילים, ומזהים אזורי קורטקס שמגיבים יותר לקטגוריה אחת מאשר לאחרת. זו לא הייתה צפייה חופשית. המשתתפים קיבעו את המבט במרכז בזמן שהוצגו בלוקים של פנים, מילות-דמה, גופים, בתים, מכוניות וגפיים. העיצוב הזה חשוב: מדד המוח לא היה פשוט תוצאה של כך שהמשתתפים הזיזו את עיניהם לעבר אותם עצמים בתוך הסורק.

לאחר מכן שאלו החוקרים עד כמה התגובה לכל קטגוריה הייתה מובחנת אצל כל אדם בקורטקס הטמפורלי הוונטרלי. דפוס פנים מובחן יותר פירושו שפעילות המוח עבור פנים הייתה באופן עקבי יותר "דמוית פנים" וקלה יותר להפרדה מקטגוריות אחרות. דפוס מילים מובחן יותר פירושו אותו דבר עבור מילים ותווים.



דוגמאות למפות fMRI משני משתתפים שנבחרו מרשימת הנבדקים במחקר; השורה העליונה והשורה התחתונה שייכות לשני אנשים שונים. בפאנל A, קווי המתאר הלבנים מסמנים את אזורי הקורטקס הטמפורלי הוונטרלי שנבדקו במחקר. פאנלים B ו-C מציגים שני ניגודים שונים: B מדגיש קורטקס שמגיב חזק יותר לפנים, ו-C קורטקס שמגיב חזק יותר למילים כתובות. פאנל D מסכם בצורת מטריצה עד כמה דומים היו דפוסי התגובה בין קטגוריות עצמים. המטריצה היא 6×6 מפני שמשימת ה-localizer השתמשה בשש קטגוריות גירוי: F מציין פנים, W מילים C-ו מכוניות, והיתר הם גופים, בתים וגפיים. דמיון חזק יותר בתוך אותה קטגוריה ודמיון חלש יותר בין קטגוריות פירושם ייצוג מוחי מובחן יותר של אותה קטגוריה. הנקודה אינה שלכל משתתף יש אותה מפה, אלא שהמחקר מדד את המפות האלה אדם-אדם.

Kollenda et al. / Nature Human Behaviour CC BY 4.0

מה הם מצאו

הדפוס המרכזי היה ספציפי לקטגוריה.

מידת המובחנות של ייצוגי פנים בחלק הלטרי הימני של הקורטקס הטמפורלי הוונטרלי הייתה במתאם עם הנטייה של המשתתף להביט בפנים במשימת הצפייה החופשית הנפרדת. הקשר הופיע גם עבור הקיבוע הראשון וגם עבור זמן השהייה. מידת המובחנות של ייצוגי מילים בחלק הלטרי השמאלי של הקורטקס הטמפורלי הוונטרלי הייתה במתאם עם הנטייה להביט בטקסט.

המאמר בדק גם שלא מדובר פשוט באפקט כללי של "קורטקס חזותי חזק יותר אצל אנשים מסוימים". הקשרים החזקים ביותר עקבו אחר הקטגוריה וההמיספרה הצפויות: פנים ב-VTC הלטרי הימני, מילים ב-VTC הלטרי השמאלי. קשרים חוצי-קטגוריות לא היו הסיפור המרכזי.

המדדים העצביים נקשרו גם להתנהגות. בתתי-מדגמים קטנים יותר, מובחנות חזקה יותר של פנים הייתה קשורה לביצועים טובים יותר ב-Test, Memory Face Cambridge ומובחנות חזקה יותר של מילים הייתה קשורה לקריאה מהירה יותר. זה נותן למדד העצבי משמעות חיצונית מסוימת: הוא אינו רק נתון מהסורק, אלא אות שקשור גם למה שאנשים מסוגלים לעשות.

מה זה לא אומר

כותרת מפתה הייתה יכולה להיות "המוח שלך מחליט מה תראה". זה חזק מדי.

המחקר אינו קובע את כיוון הסיבתיות. ייתכן שאדם מביט יותר בפנים מפני שייצוגי הפנים שלו מדויקים יותר. ייתכן גם שייצוגי הפנים שלו מדויקים יותר מפני ששנים של הסתכלות בפנים עיצבו את החלק הזה של מערכת הראייה. ואפשר ששני הדברים מתפתחים יחד. המחברים משאירים במפורש את השאלה ההתפתחותית הזאת פתוחה.

זה גם לא אומר שאנשים בעלי הרגלי מבט שונים רואים סצנות פיזיות שונות. כולם הסתכלו באותן תמונות. ההבדל הוא בדגימה: אילו עצמים מקבלים עדיפות, איזה מידע נאסף קודם ואילו קטגוריות מיוצגות בצורה מובחנת יותר בקורטקס החזותי.

וזה אינו מבחן אבחוני ליחידים. המתאמים משמעותיים ברמת הקבוצה, אבל הם אינם כלי שאפשר לומר בעזרתו "מפת המוח של האדם הזה מנבאת בדיוק איך יסתכל על התמונה הזאת".

למה זה חשוב

לעיתים מתארים ראייה כאילו העיניים הן מצלמה שמזינה את המוח. ראייה אמיתית פעילה הרבה יותר. העין בוחרת, מרגע לרגע, איזה מידע להביא לרזולוציה גבוהה. הבחירות האלה הופכות לקלט שממנו המוח לומד ופועל.

המאמר הזה מבסס טוב יותר את הלולאה הזאת. הוא מקשר בין החלק הפעיל — תנועות העיניים בסצנות — לבין החלק הייצוגי — מפות בררניות לקטגוריות בקורטקס החזותי הוונטרלי — בתוך אותם אנשים. התוצאה מקשה להתייחס ל"ארגון הקורטקס החזותי" ול"התנהגות חזותית" כשתי רמות נפרדות. לפחות אצל מבוגרים, הן נראות מותאמות זו לזו.

ההתאמה הזאת היא הסיפור האמיתי. לא שאנשים שמביטים בפנים הם "סוג" אחד ואנשים שמביטים בטקסט הם סוג אחר. לא שאזור מוח אחד מסביר אישיות. התוצאה צרה ושימושית יותר: הבדלים יציבים באופן שבו אנשים חוקרים סצנות חזותיות מתיישרים עם הבדלים יציבים במידת החדות שבה הקורטקס החזותי שלהם מייצג את הדברים שהם נוטים לחפש.

בקצרה

מחקר ב-Behaviour Human Nature עקב אחר האופן שבו 102 מבוגרים הסתכלו על 700 סצנות מורכבות, ולאחר מכן סרק תת-קבוצה של 61 משתתפים ב-fMRI כדי למפות תגובות חזותיות בררניות לקטגוריות. אנשים שנטו להביט מוקדם יותר ולמשך זמן רב יותר בפנים הציגו ייצוגי פנים מובחנים יותר בקורטקס הטמפורלי הוונטרלי הלטרלי הימני; אנשים שנטו להביט בטקסט הציגו ייצוגי מילים מובחנים יותר בצד השמאלי. המדדים העצביים נקשרו גם לזיהוי פנים ולביצועי קריאה בתתי-מדגמים קטנים יותר. המחקר מתאמי, לא סיבתי: הוא מראה שאצל מבוגרים הרגלי מבט פעילים וארגון מוחי בררני לקטגוריות מתאימים זה לזה, לא מי מהם יוצר את האחר.

בדיקה מפוכחת

מה המחקר מראה: נטיות אישיות יציבות להביט בפנים או בטקסט בסצנות טבעיות קשורות לייצוגים תואמי-קטגוריה בקורטקס החזותי הוונטרלי.

מה סביר אבל לא הוכח: שניסיון חזותי ארוך-טווח וכיוונון מוחי מחזקים זה את זה במהלך ההתפתחות. המחקר מתיישב עם הרעיון הזה, אך אינו בודק את כיוון ההתפתחות.

מה הוא אינו מראה: שאנשים רואים ממש עולמות שונים; שהרגלי מבט מקובעים מלידה; או שמפות ה-fMRI גורמות לתנועות העיניים.

המגבלות העיקריות: תכנון מתאמי; מדגם מבוגרים מאוכלוסיית אוניברסיטה; תתי-מדגמים קטנים יותר למבחני זיהוי פנים וקריאה; ו-localizer נפרד ב-fMRI במקום fMRI בזמן צפייה חופשית עצמה.

כמה ביטחון צריך להיות לקורא כללי? גבוה בכך שנטיות המבט אמיתיות ויציבות במדגם הזה, ובינוני-גבוה בכך שהן קשורות לייצוגים חזותיים תואמי-קטגוריה. נמוך לגבי כל סיפור סיבתי עד שמחקרי התפתחות או התערבות יבדקו אותו ישירות.

מקורות

10.1038/s41562-026-02494-5 DOI, (2026) Behaviour Human Nature al., et Kollenda
<https://doi.org/10.1038/s41562-026-02494-5>

stimuli and code data, anonymized for repository OSF
<https://osf.io/9fhxk/>