

דוח ניתוח זיהוי AI

בדיקה היוריסטית של טבעיות הכתיבה וסיכון לזיהוי AI. לעבודה שאינה באנגלית אין מודל זיהוי AI אמין, ולכן הסיכון מוערך לפי מדדים היוריסטיים מקומיים בלבד.

הופק בתאריך: 11 July, 2026 | נושא: Use of artificial-intelligence tools and learning motivation among undergraduate students: a quantitative survey study

סיכום פעולה ללקוח

- סיכון כולל לזיהוי AI: לא נקבע על ידי המודל — שפת העבודה אינה אנגלית, ולכן חלים רק המדדים ההיוריסטיים.
- רמת ודאות: נמוכה — מבוססת על מדדים היוריסטיים מקומיים בלבד.
- סעיפים לבדיקה: אף סעיף אינו דורש טיפול מבני — רק הצעות הניסוח האופציונליות שלמטה.
- פעולות מומלצות:
- גוון את אופן הפתיחה של המשפטים.

ניסוחים לתיקון

בעיה	תיקון מוצע
משפטים רבים נפתחים ב-"עבור" (8 פעמים)	גוון את אופן הפתיחה של המשפטים.
משפטים רבים נפתחים ב-"לפיכך" (6 פעמים)	גוון את אופן הפתיחה של המשפטים.
משפטים רבים נפתחים ב-"מאמר" (5 פעמים)	גוון את אופן הפתיחה של המשפטים.
משפטים רבים נפתחים ב-"הממצאים" (5 פעמים)	גוון את אופן הפתיחה של המשפטים.
משפטים רבים נפתחים ב-"המאמר" (5 פעמים)	גוון את אופן הפתיחה של המשפטים.

בעיה	תיקון מוצע
משפטים רבים נפתחים ב-"מבחינה" 4) (פעמים)	גונו את אופן הפתיחה של המשפטים.

תקציר מנהלים

דוח זה מנתח את המסמך לאיתור סימני תוכן שנוצר ב-AI, ומשווה ציונים לפני ואחרי שלב הליטוש (polish) האקדמי.

מדד	לפני ליטוש	אחרי ליטוש	שינוי
ציון זיהוי AI	A/N	A/N	לא זמין
Burstiness	100%	100%	↓ 0pp
Perplexity Proxy	82%	82%	↓ 0pp

⚠ ציון זיהוי ה-AI מבוסס-המודל מוצג רק לעבודות באנגלית. המזהה (chatgpt-/Hello-SimpleAI) אומן על טקסט אנגלי והיה מחזיר ציון לא-אמין עבור שפת העבודה הזו, ולכן לא הופעל. מדדי ה-heuristic המקומיים (Perplexity Proxy, Burstiness) המוצגים מעלה תקפים לכל השפות.

הבנת המדדים

כלי זיהוי AI אינם חסיני טעות. מזהים שונים משתמשים בשיטות שונות ועלולים להפיק תוצאות סותרות. דוח זה מספק מחווני זיהוי המבוססים על ניתוח סטטיסטי ועל מודל זיהוי ייעודי, אך יש לפרש את התוצאות כהכוונה, לא כאמת מוחלטת.

ציון זיהוי AI (מודל chatgpt-detector-roberta)

מה זה מודד: ציון זה מופק על ידי מודל ייעודי לזיהוי AI (chatgpt-detector-/Hello-SimpleAI) (roberta) — טרנספורמר RoBERTa שעבר fine-tuning ספציפית לסיווג טקסט כאנושי או כמוצר-AI. המודל מנתח דפוסים לשוניים, התפלגויות בחירת מילים, ועקביות הקשרית כדי להפיק ציון הסתברותי.

איך לקרוא את זה: 0% = ודאי שנכתב על ידי אדם, 100% = ודאי שנוצר על ידי AI — ככל שהציון גבוה יותר, כך גדלה הסבירות שהטקסט יסומן כמוצר-AI. רמות הסיכון שלהלן מסבירות את משמעות כל טווח.

רמות סיכון

סיכון נמוך (0–30%): הטקסט מציג מאפיינים אופייניים לכתיבה אנושית. רוב כלי הזיהוי לא צפויים לסמן אותו.

סיכון בינוני (30–60%): לטקסט יש מספר דפוסים דמויי-AI. חלק מהמזהים עלולים לסמן קטעים. כדאי לעבור על הקטעים המסומנים ולנסח אותם באופן טבעי יותר.

סיכון גבוה (60–100%): הטקסט מציג סימני AI חזקים. רוב המזהים יסמנו אותו בסבירות גבוהה. ייתכן שיידרש שכתוב משמעותי.

Burstiness

מה זה מודד: Burstiness מודד את השונות באורכי משפטים לאורך המסמך. כותבים אנושיים יוצרים באופן טבעי טקסט "בעל פיצוצי תנופה" — משפט קצר וחד אחריו ארוך ואנליטי, ואז מעבר באורך בינוני. מודלי AI נוטים להפיק משפטים באורכים דומים, מה שמייצר מקצב מונוטוני ואחיד.

איך לקרוא את זה: טכנית, Burstiness הוא מקדם השונות (סטיית תקן / ממוצע) של מספרי המילים בכל משפט, מנורמל לסקאלת 0-1. Burstiness של 0% משמעו שכל המשפטים זהים באורכם (אופייני ל-AI). Burstiness של 60% ומעלה מציין שונות אנושית טבעית. חשוב: Burstiness נמוך בחלק מהסקציות הוא נורמלי, לא פגם. סקציות הבנויות על תבנית מבנית קבועה — Abstract, Research Objectives, השערות, רשימת מקורות, רשימות bullets — הן נמוכות-Burstiness מטבען גם כשנכתבו במלואן על ידי כותב אנושי. המדד מודד שונות אורכי משפטים, ורשימה של 4 מטרות מחקר בתבנית "לחקור... / לזהות... / להעריך..." תקבל בערך 25%-40% מעצם מבנה. יש לקרוא את הציון בהקשר של סוג הסקציה; סקציות אנליטיות (דיון, רקע תיאורטי) הן המקומות המשמעותיים לחפש בהם Burstiness נמוך כסימן ל-AI.

Perplexity Proxy

מה זה מודד: Perplexity אמיתי דורש הרצת הטקסט במודל שפה כדי למדוד עד כמה הוא "מופתע" מכל מילה. מאחר שאנו לא טוענים מודלי ML כבדים, מדד זה מקרב את ה-perplexity באמצעות פרוקסים מדידים: עושר אוצר מילים (כמה מילים ייחודיות יחסית לסך המילים), שימוש במילים ארוכות, ושונות במבני משפט.

איך לקרוא את זה: ערכים גבוהים מציינים טקסט מגוון ופחות צפוי — מאפיין של כתיבה אנושית. ערכים נמוכים מציינים דפוסים נוסחתיים וצפויים אופייניים לפלט AI. זהו קירוב סטטיסטי, לא מדידת perplexity אמיתית.

עושר אוצר מילים (Root TTR)

מה זה מודד: ה-root type-token ratio סופר מילים ייחודיות חלקי השורש הריבועי של סך המילים. זה מתקן לאורך מסמך – טקסטים ארוכים יותר חוזרים באופן טבעי על מילים יותר. טקסט שנוצר ב-AI נוטה לחזור על אותו אוצר מילים שוב ושוב, בעוד שכותבים אנושיים שואבים מלקסיקון רחב יותר עם בחירות מילים מגוונות. איך לקרוא את זה: הערכים מדווחים כציון גולמי. ערכים גבוהים יותר מציינים אוצר מילים עשיר ומגוון יותר. מדד זה ניזון לציון ה-Perplexity Proxy המורכב.

מגבלות

- ציון זיהוי ה-AI תלוי במודל chatgpt-detector-roberta, שייתכן שאינו מכסה כל שפה באותה מידה. התוצאות הכי אמינות בטקסט באנגלית.
- Burstiness ו-Perplexity Proxy הם קירובים סטטיסטיים. מסמכים קצרים (פחות מ-500 מילים) מפיקים ציונים פחות אמינים.
- טכנולוגיית הזיהוי מתפתחת במהירות. הציונים משקפים את מצב שיטות הזיהוי בזמן ההפקה.
- כתיבה אקדמית חולקת באופן טבעי דפוסים מסוימים עם פלט AI (מרשם פורמלי, טיעונים מובנים). הדבר עשוי להעלות ציונים מעל לנדרש.