

UC Riverside

UC Riverside Previously Published Works

Title

Bridging or broadening Gaps? AI-Assisted professional writing among native and non-native English Writers

Permalink

<https://escholarship.org/uc/item/0q91t81t>

Journal

Computers in Human Behavior, 177

ISSN

0747-5632

Authors

Shin, Inyoung

Choung, Hyesun

Choi, Mina

Publication Date

2026-04-01

DOI

10.1016/j.chb.2025.108897

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Highlights

Bridging or Broadening Gaps? AI-Assisted Professional Writing among Native and Non-Native Writers

- Study examines AI use in professional writing by native and non-native English writers.
- Participants revised a cover letter using simple or advanced AI-generated text.
- With high self-efficacy, native and non-native writers show similar AI adoption.
- With low self-efficacy, non-native writers prefer simple text, natives prefer advanced one.
- Sociolinguistic gaps persist, especially among lower-skilled writers.

Bridging or Broadening Gaps? AI-Assisted Professional Writing among Native and Non-Native Writers

Abstract

This study examines how AI influences existing differences in professional writing between native and non-native English writers (NEWs and NNEWs) in the United States, reflecting individuals' task-related proficiency and language-based social positioning. We compare how these two groups integrate AI-generated content into their writing and include writing self-efficacy as a moderator to examine whether perceived task proficiency shapes the differences in AI use between the two groups. We also test an underlying social-psychological mechanism by examining the mediating role of perceived superiority: the extent to which participants judged the AI-generated text as better than their own writing. In an online experiment, 327 NEWs and NNEWs were recruited to write a job application cover letter for a hypothetical scenario. Participants were randomly assigned to receive AI-generated content written at either a simple or advanced lexical level and were asked to revise their letter as they chose, potentially incorporating the AI-generated content. Using natural language processing techniques, we measured the extent to which participants integrated AI-generated content. Findings show that NNEWs favored simpler content, while NEWs tended to incorporate advanced content. This difference was pronounced among those with lower writing self-efficacy. However, the perceived superiority of AI-generated text did not explain this pattern. These findings show that while AI can support lower-skilled groups in specific tasks, sociolinguistic gaps still remain even with AI-assistance. We conclude by calling for AI training frameworks that scaffold learning to address both proficiency gaps and social constraints.

Keywords: AI-Assisted Writing, AI Divide, Sociolinguistic gaps, Writing Self-Efficacy, Language Proficiency, Skill-Biased Technological Change, Digital Divide

1. Introduction

The advancement of Artificial Intelligence (AI), especially based on large language models (LLMs), such as ChatGPT, has shown remarkable capabilities in generating complex verbal content comparable to human output. Alongside these developments, there is growing anticipation that LLMs can significantly reduce the time and effort required for content production. AI tools are increasingly viewed as innovative assistive technologies that may help individuals with lower proficiency overcome skill barriers and gain access to knowledge-intensive tasks (Brynjolfsson et al., 2025), including not only general writing (Noy and Zhang, 2023) but also programming (Peng et al., 2023), data analysis and interpretation (Jansen et al., 2025), and even management consulting (Dell’Acqua et al., 2023).

Empirical evidence on whether and how AI improves the productivity of users with limited proficiency is still emerging and remains mixed. Historically, digital technologies such as computers and the internet have been discussed as forms of “skill-biased technological change,” (Bresnahan et al., 2002) which disproportionately benefit individuals with advanced skills and experience. More recently, however, scholars have suggested that AI may reverse this pattern, with individuals at lower proficiency levels potentially experiencing greater performance gains than those already highly skilled (Brynjolfsson et al., 2025). In a similar vein, generative AI is expected to improve the productivity of second-language learners or individuals who do not regularly use their first language in professional or social contexts (Ito et al., 2023; Wang et al., 2024; Tian et al., 2025) by helping them produce and refine professional writing more efficiently. In the United States (U.S.), millions of people identify a language other than English as their primary language (Dietrich and Hernandez, 2022). Non-native English speakers often experience anxiety and inhibition in communication, which can ultimately lower productivity. AI has the potential to alleviate some of these challenges and enhance the productivity of non-native English speakers.

Nonetheless, it may be premature to conclude that AI advancements reduce existing difficulties that non-native individuals face in professional practice. Scholarship framing AI as a reversal of skill-biased technological change has mostly focused on comparisons between novices and experts within specific organizational contexts (e.g., Brynjolfsson et al. 2025 on customer service workers; Peng et al. 2023 on novice vs. advanced programmers), providing

limited insights into how AI adoption and use interact with social structural and psychological factors. As the “digital divide” research has long suggested (Norris, 2001; Gonzales, 2016; Reisdorf et al., 2020), technology adoption and its use are embedded in individuals’ social constraints (Norris, 2001; Gonzales, 2016; Reisdorf et al., 2020). Supporting this view, a recent large-scale study in Denmark (Humlum and Vestergaard, 2025) suggested that structural inequalities hinder the positive impacts of AI by showing that women and lower-earning workers had slower adoption rates.

We presume that AI adoption and use are influenced by not only individuals’ proficiency but also their social identities and the constraints they entail. To examine this, we specifically compare non-native English writers (NNEWs) and native English writers (NEWs) because English nativity functions as a dual proxy for both writing ability and language-based social positioning, particularly in the U.S. context (Schwartz and Stiefel, 2006). Compared to NEWs, NNEWs often report lower writing proficiency and less confidence in English (Silva, 1993). Yet the differences between NNEWs and NEWs do not stem only from differences in proficiency. Some NNEWs are fluent at levels comparable to NEWs but, because of their cultural and social positioning, still adopt different communication styles from NEWs (Maier, 1992) and face disadvantages in professional settings where native-like fluency is privileged (Piller, 2016).

In this study, we specifically focus on comparing the ways that NNEWs and NEWs use AI-generated content in job application cover letters, a high-stakes professional task central to employment decisions in the U.S. (Tian et al., 2025). Cover letters not only summarize qualifications but also serve as signals of social and communication skills (Schwartzberg, 2025). Because professional writing requires complex lexical choices, structural precision, and vocabulary aligned with workplace norms, NNEWs, who often have higher anxiety in English writing (Cheng, 2004) may turn to AI-generated text to meet these expectations. At the same time, this context allows us to examine how individuals respond to AI-generated content that may itself fall short of professional standards. Using an experimental design, we test how NNEWs and NEWs evaluate and adopt AI-generated text at varying levels of lexical complexity.

Although linguistic nativity provides a useful dual lens for writing ability and social

constraints, it also makes it difficult to determine whether differences between NNEWs and NEWs stem primarily from proficiency gaps or social identity constraints. To disentangle these effects, we include self-assessed writing efficacy as a moderating factor. While not an objective measure of ability, writing self-efficacy strongly influences how individuals approach writing tasks in real-world contexts (Bruning et al., 2013). By examining writing efficacy as a moderator, we test whether sociolinguistic gaps between NNEWs and NEWs persist once perceived competence is taken into account.

Beyond perceived proficiency, we also explore social-psychological mechanisms underlying differences between NNEWs and NEWs in AI adoption, focusing on how they compare their own writing with AI-generated text. NNEWs are more likely than NEWs to view their own writing as less competent (Mitchell et al., 2023) and to perceive AI-generated content as superior to their own work (Logg et al., 2019; Mendez et al., 2021). These comparative perceptions may explain why NNEWs are more likely than NEWs to adopt AI-generated content and ultimately rely on AI as a compensatory tool. To test this possibility, we examine the mediating role of perceived superiority, asking whether differences in AI adoption emerge through comparison processes in which NNEWs undervalue their own writing while overvaluing AI-generated text.

Overall, our study makes three key contributions to the literature on social impacts of AI. First, by comparing English nativity in the U.S., we examine not only whether AI is disproportionately accepted by groups perceived as less skilled but also how adoption interacts with structural social factors, integrating insights from research on skill-biased technological change and digital divide scholarship. Second, by incorporating writing self-efficacy as a moderator, we clarify the role of perceived competence in a task and how it interacts with structural social factors tied to language nativity in AI-assisted writing. Third, we explore a social-psychological mechanism underlying these group differences by testing the mediating effects of evaluations of AI-generated text compared to one’s own writing. Our findings offer practical implications for supporting NNEWs, aligning with sociocultural theory in cognitive development (Vygotsky, 1978), which emphasizes that education tools like AI should be approached as scaffolding within contexts shaped by structural inequalities. Taken together, our study provides a comprehensive approach that accounts for proficiency differences, so-

ciocultural disparities, and social-psychological processes in AI-assisted writing.

2. Backgrounds

2.1. Technology, Proficiency and Social Constraints

Historically, major technological advancements have coincided with large-scale social and economic changes (Hampton, 2016). Understanding how new technological advancements affect economic and societal structures is a critical concern for academic researchers, practitioners, and policymakers. In the case of contemporary digital technologies such as computers and the internet, two overlapping but distinct frameworks have been established: 1) skill-biased technological change, which emphasizes benefits for individuals with higher skills, and 2) the digital divide, which highlights unequal access to and use of technology across social groups.

The first framework, rooted in labor economics, shows how digital technologies disproportionately benefit higher-skilled workers by raising their productivity and wages more than those of lower-skilled workers. This phenomenon is described as skill-biased technological change (Autor et al., 1998; Bresnahan et al., 2002; Katz and Murphy, 1992). Most studies in this tradition compare workers with and without STEM or other technical specialties. Scholars have further emphasized that access to advanced skills is closely tied to socioeconomic status, including income, race, and education, such that digital technologies not only reward skill differences but also reinforce pre-existing social constraints individual workers face (Goldin and Katz, 2008).

The second line of research, developed largely within communication, media studies, and sociology, focuses more directly on the role of social economic status (SES) in technology adoption and use. Early studies showed that individuals with greater economic resources and higher educational attainment were typically the first to adopt and integrate digital technologies into daily life and professional practice (Norris, 2001; Rogers, 2001). This phenomenon is referred to as digital divide (Gonzales, 2016). Although access has broadened, disparities in digital skills and usage patterns remain, often referred to as second-level divides (DiMaggio and Hargittai, 2023; Dutton et al., 2013; Livingstone and Helsper, 2007). More recently, digital divide research connects to the narrative of skill-biased technological change

by introducing the concept of third-level divide that highlights how differences in skills and usage translate into unequal social and economic outcomes (Reisdorf et al., 2020; van Deursen and Helsper, 2015). Both the skill-biased technological change and digital divide frameworks emphasize that the impacts of new technologies on productivity and outcomes are not simply technical but embedded in social structures, such that new technologies often reward pre-existing advantages and reproduce broader inequalities.

The recent emergence of generative AI has raised optimism in the studies of skill-biased technological change, which have focused on the productivity and wage gaps between skilled and unskilled workers. Unlike traditional digital technologies, generative AI powered by LLMs can be used with minimal technical knowledge (Brynjolfsson et al., 2025). These tools simplify complex information and generate content tailored to user input through user-friendly interfaces. Early evidence shows higher adoption rates of generative AI among unskilled workers or those lacking prior experience (Lund et al., 2023). Some scholars have described this trend as an example of the reverse of skill-biased technological change, where AI tools benefit individuals with fewer traditional skills more than those with higher levels of conventional expertise (Brynjolfsson et al., 2025; Noy and Zhang, 2023; Peng et al., 2023). For example, Brynjolfsson et al. (2025) found that customer service workers with fewer or lower-level work-related skills tend to gain greater performance improvements from using AI assistants than more skilled workers. Similarly, Peng et al. (2023) showed that older and less experienced programmers gained more from using LLM-based programming assistants, such as GitHub Copilot, than their more experienced counterparts.

Despite growing optimism, whether and how AI actually mitigates existing gaps between workers remains an open question. Much of the empirical evidence relies on data from single organizations and simplified comparisons between “less-skilled” and “more-skilled” workers (Brynjolfsson et al., 2025; Peng et al., 2023). The positive impacts of AI suggested by this work may be limited by pre-existing socio-economic factors. For example, Brynjolfsson et al. (2025) noticed that there were two groups of customer service workers in their observation, one in the Philippines and the other in the U.S. and U.K., but compared them in terms of skill gaps rather than recognizing the structural or contextual differences underlying these groups. Humlum and Vestergaard (2025) further argue that systemic barriers and pre-

existing inequalities may blunt the positive effects of AI adoption. Using a national survey in Denmark, they found that higher-earning employees adopted generative AI earlier and benefited more, while women were less likely than men to use AI tools. Moreover, recent research shows the gaps in AI adoption and use across SES status (Bassignana et al., 2025; Daepf and Counts, 2024; Sidoti et al., 2025), extending the classic literature on the digital divide. However, connections to the second- and third-level digital divides, concerning how different groups use AI and how such usage shapes social and economic outcomes, have not yet been systematically examined in the context of AI.

In this study, we compare NNEWs and NEWs residing in the U.S., which relate to differences in writing skills, while at the same time reflecting social and cultural positioning such as immigrant status, educational attainment, and racial background in the U.S. (Batalova and Zong, 2016; Schwartz and Stiefel, 2006). By focusing on linguistic nativity, our study bridges insights from research on skill-biased technological change and the digital divide, showing how both task-related skills and social structures influence AI use. We further provide empirical evidence on how NNEWs and NEWs integrate AI-generated content into the process of composing a job application cover letter, a realistic writing task where AI assistance is increasingly common (Tian et al., 2025). Our approach demonstrates how AI use influences writing performance and productivity beyond descriptive accounts of AI access or adoption.

Research has shown that NNEWs and NEWs generally differ in English writing proficiency (Silva, 1993). Even when NNEWs report fluency comparable to NEWs, their writing tends to differ in style, tone, and rhetorical strategy (Maier, 1992). For example, NEWs are more likely to use formal, indirect politeness strategies in professional writing, whereas NNEWs often adopt more direct or informal approaches (Hou and Li, 2011). Since cover letters require standardized formats and culturally embedded norms, the differences between NNEWs and NEWs can significantly affect how each group composes such texts and how they might turn to generative AI for support (Tian et al., 2025). In general, NNEWs are more likely to adopt generative AI tools for writing (Ito et al., 2023; Wang et al., 2024). Based on this, we hypothesize:

H1: NNEWs are more likely than NEWs to integrate AI-generated content into their

own writing when composing job application cover letters.

Professional writing typically requires formal language aligned with college-level standards (Tian et al., 2025). However, NNEWs’ inclination toward AI-generated content could persist even when AI generates a cover letter using a relatively simple lexicon because of their lack of writing proficiency and/or sociolinguistic positioning. Prior research suggests that NNEWs tend to seek more external help with their writing because they often experience greater social pressure, including expectations to conform to native-like writing norms and concerns about being negatively evaluated, compared to NEWs (Piller, 2016). It is also possible that individuals prefer AI-generated assistance that aligns with their personal writing style when seeking help. NNEWs often favor shorter, simpler constructions, while NEWs are generally more comfortable using complex and nuanced language (Maier, 1992). In our online experiment, we manipulate the lexical complexity of AI-generated content to test how NNEWs and NEWs differ in their adoption and incorporation of simple versus advanced styles into their professional writing. Specifically, we propose:

H2: NNEWs are more likely than NEWs to adopt and integrate AI-generated content when it is written in a simpler, rather than a more advanced, lexical style.

2.2. Writing Self-Efficacy, Language Nativity, and AI Use in Professional Writing

Novice and lower-skilled individuals often adopt and rely more on external assistance when completing or improving their tasks, provided such resources are available. In learning research, this tendency is well recognized in the sociocultural theory of cognitive development (Vygotsky, 1978), which posits that individuals who doubt their ability are more likely to seek support from peers, teachers, or specialized assistants. Extending this logic to technology adoption, we anticipate that when tools compensate for gaps in task performance, individuals who perceive themselves as less capable are more eager to adopt and use such tools actively. Brynjolfsson et al. (2025)’s findings on the benefits of AI for novice and lower-skilled workers may thus be partly explained by their lower perceived competence in performing tasks, coupled with their willingness to improve through technological support.

In the domain of writing, the perception of one’s ability is referred to as writing self-efficacy, a construct derived from Bandura’s theory of self-efficacy (Bandura, 1997), which

emphasizes that individuals’ beliefs in their capacity to learn and perform tasks effectively are critical to performance success. Applied to writing, writing self-efficacy is defined as one’s belief in their capabilities as a writer (Bruning et al., 2013; Pajares, 2007). Higher writing self-efficacy has been associated with better writing quality (Shell et al., 1989), greater adherence to standards (Zimmerman and Bandura, 1994), and lower levels of writing apprehension (Pajares and Johnson, 1994).

Writing self-efficacy is not typically discussed as a factor influencing new technology adoption. Instead, research has focused on technology-related self-efficacy (e.g., confidence in one’s ability to use a new system), which has been known as a key antecedent for adopting tools such as computers and other digital technologies (Davis, 1985, 1989). However, because the core function of generative AI tools based on LLMs lies in producing text from user prompts, a growing body of research has begun to consider writing self-efficacy as an important factor in AI adoption. For example, Jelson and Lee (2024) found that college students with lower writing self-efficacy relied more heavily on AI writing tools (e.g., ChatGPT) for essay composition, whereas students with higher efficacy were less inclined to depend on them. They further showed that writing self-efficacy can be a stronger predictor of AI usage than traditional factors such as perceived usefulness or ease of use, which are typically emphasized in technology acceptance research (Davis, 1989).

Beyond acceptance, individuals’ perceived competence in a given task strongly influences their attitudes toward and trust in AI tools. In writing, individuals with lower writing confidence were found to trust LLM-generated suggestions or corrections more readily, often viewing these tools as compensatory for their own perceived limitations (Jelson and Lee, 2024). In a controlled trial with non-native English speaking graduate students, Nazari et al. (2021) showed that an AI-powered writing assistant significantly boosted students’ writing self-efficacy and reduced writing-related anxiety, which in turn fostered more positive attitudes toward the tool. By contrast, more confident writers often have a skeptical attitude toward AI-generated content and identify stylistic flaws or shallow reasoning in AI-generated content, viewing AI as a supplemental aid rather than an authoritative source (Nazari et al., 2021).

Writing self-efficacy, particularly in English writing, is strongly correlated with linguistic

nativity. In general, NNEWs report lower writing self-efficacy than NEWs (Mitchell et al., 2023; Prat-Sala and Redford, 2012). Prior research showing that NNEWs are more likely to use and rely on AI for writing support (Ito et al., 2023; Wang et al., 2024) may partly be explained by variation in writing self-efficacy. NNEWs not only objectively have lower competence in writing but also often underestimate their abilities, leading to higher writing anxiety and a stronger perceived need for assistance (Cheng, 2004; Woodrow, 2011). Yet, writing self-efficacy alone cannot fully explain the sociolinguistic gap between NNEWs and NEWs. Writing self-efficacy is distinct from linguistic nativity, as it only reflects individuals’ perceptions of their writing skills and abilities rather than broader social-structural factors tied to linguistic nativity (Bruning et al., 2013). Even when NNEWs and NEWs report comparable levels of writing self-efficacy, their writing styles and behaviors may differ due to social conditions associated with linguistic nativity (Mitchell et al., 2023).

We specifically propose that writing self-efficacy interacts with linguistic nativity in affecting the adoption of AI-generated content. A meta-analysis by Sun et al. (2021) found that the impact of writing self-efficacy on English writing achievement was significantly stronger for NNEWs than for NEWs. This suggests that, to achieve good performance, NNEWs more than NEWs must find ways to overcome low confidence and believe in their abilities. Accordingly, NNEWs are more likely than NEWs to adopt AI-generated content, especially when they have low self-efficacy. At the same time, even when NNEWs report high self-efficacy comparable to NEWs, they often display distinct writing habits, such as seeking external feedback or double-checking their work (Williams and Takaku, 2011). It is also possible that differences between NNEWs and NEWs appear in distinct ways among both low- and high-efficacy writers.

Combined with the lexical complexity of AI-generated content that we manipulate in the experiment, we expect a three-way interaction among lexical complexity, linguistic nativity, and writing self-efficacy.

H3: Writing self-efficacy moderates the effects of nativity and its interaction with the lexical complexity of AI-generated content on participants’ acceptance and integration of AI-generated text into their writing.

2.3. *Perceived Superiority and AI Acceptance*

As discussed, we hypothesize interactions between lexical complexity, linguistic nativity, and writing self-efficacy in the acceptance of AI-generated content. We additionally examine the mechanism underlying these interaction dynamics by focusing on a social-psychological process through which individuals compare AI-generated outputs to their own writing. We term this process *perceived superiority*, referring to the degree to which people judge AI-generated content as better than their personal efforts. Perceived superiority can serve as an amplifying factor, heightening the baseline tendency to prefer algorithmic over human-generated outcomes (Logg et al., 2019; Lee and See, 2004; Gunaratne et al., 2018). For example, Logg et al. (2019) found that people tend to trust algorithmic outputs because they perceive them as more objective, consistent, or neutral, even when algorithmic performance is comparable to or worse than human alternatives. This phenomenon, known as algorithmic appreciation, highlights how subjective perceptions or prior expectations often drive AI reliance more than actual performance differences.

Similarly, the Technology Acceptance Model (TAM) posits that users’ perceptions of a technology’s usefulness and competence directly influence their intention to adopt it (Davis, 1989). Recent studies applying TAM to AI systems have consistently found that perceived competence and perceived usefulness predict individuals’ willingness to adopt AI-based tools and other generative AI applications (Choung et al., 2023; Jeong et al., 2024). When individuals believe that AI-generated writing surpasses their own capabilities, they may perceive AI as both useful and competent, increasing the likelihood of accepting and relying on AI assistance.

Apart from the literature on technological adoption, the sociocultural theory of cognitive development (Vygotsky, 1978) suggests that perceived superiority may help explain some individuals’ tendency to accept and rely more on AI-generated content. According to Vygotsky (1978), individuals are more inclined to accept assistance from sources they perceive as more capable than themselves, a concept described as the “more knowledgeable other.” Within this framework, scaffolding refers to the strategic support provided by the more knowledgeable other to help learners accomplish tasks within their own zone of proximal development (ZPD), referring to the psychological range between what learners can do independently and

what they believe they can achieve with appropriate external assistance. Recent scholarship has begun to explore AI tools as a form of the more knowledgeable other, enabling users to perform tasks beyond their independent capabilities (Cai et al., 2024; Sætra, 2025). In writing contexts, if users perceive AI-generated content as providing more capable assistance than their own writing, they may be more inclined to adopt and integrate AI-generated content into their work. This mechanism of perceived superiority may be especially salient in high-stakes tasks, such as writing job application cover letters, where individuals seek to present their best possible work.

The ways that people perceive superiority of AI-generated content may vary across NNEWs and NEWs. Specifically, NNEWs, who often face greater challenges in formal English writing, may be more likely than NEWs to perceive AI-generated content as superior to their own writing. In other words, perceived superiority may mediate the relationship between writing self-efficacy and the adoption of AI-generated text. Even when AI-generated content includes relatively simple vocabulary, NNEWs possibly choose to adopt it because they perceive it as better than what they could produce on their own. Therefore, we hypothesize:

H4: The perceived superiority of AI-generated content mediates the effects of nativity, and its interaction with lexical complexity and writing efficacy, on the acceptance of AI-generated content.

3. Methods

To examine our hypotheses, we conducted a 2×2 mixed experiment (between- and within-subjects) embedded within a survey. Participants first completed a pre-treatment task by writing a job application cover letter for a fictional persona. For the treatment, they were shown a sample letter accompanied by a short video depicting ChatGPT generating the letter. The letter had been pre-generated at the assigned complexity level but was presented through a customized interface to appear as real-time output. In the post-treatment task, participants revised their original letter with access to both their draft and the AI-generated letter. We then compared the revised letters with the AI text, using textual similarity as a proxy for the reliance on and acceptance of AI-generated content.

3.1. Participants

With IRB approval, we recruited participants from the online crowdsourcing platform, Prolific. To align with the study’s purpose, participation was restricted to adults aged 18 and older residing in the U.S. Using Prolific’s pre-screeners, we recruited participants from two distinct pools: 1) those who identified English as their first language (i.e., NEWs) and 2) those who identified a language other than English as their first language (i.e., NNEWs). To ensure participant eligibility, we included additional screener questions at the beginning of the survey. Individuals who reported multiple first languages, including English, were excluded from the study. Several attention check questions were included throughout the experiment. Participants who failed any of these checks were automatically exited from the study and did not continue further. Upon completion, they received \$3 compensation.

A total of 327 individuals participated in the study. Three participants were excluded for not providing quality responses in the cover letter writing tasks. The final sample included 324 participants, residing across 48 states in the U.S. Of the final sample, 43.8% identified as NEWs, and 57.2% identified as non-native English writers. Among NNEWs, more than ten different first languages were reported. The most common first language was Spanish (15.2%), followed by Chinese (7.4%). Among all participants, 61.4% identified as female. Nearly half (49.3%) identified as White. While 46.6% reported currently working full-time, 62% of those had been in their current position for less than three years. A majority of participants indicated that the highest degree they had earned or expected to earn was above a bachelor’s degree. In terms of income, 46.6% reported an annual household income below \$75,000. Additionally, most participants (98.5%) reported that they had heard of or were currently using generative AI tools such as ChatGPT in their daily lives for English writing.

3.2. Experimental Materials

Generative AI systems such as ChatGPT do not reliably reproduce the same output even with identical prompts. To ensure that all participants were exposed to identical AI-generated text, we pre-produced cover letters through multiple rounds of prompting with ChatGPT-4. However, simply presenting static text could have reduced ecological validity,

as participants would not experience how AI typically produces content. To address this, we configured ChatGPT to display the finalized texts as if they were generated in real time and created a screen-capture video of the chat interface. The processes of creating the stimuli and the final outputs are shown in Figure 1.

[Insert Figure 1 here]

AI-generated letters: Our goal was to obtain pre-generated AI texts that included all required information to maintain experimental control, allowing for a fair comparison of whether NNEWs and NEWs, when believing the text was AI-generated, would revise their writing to be similar to the provided version. We first created background information for a fictional applicant, “Taylor Williams,” including education, skills, and work experience, and then established criteria to evaluate the AI-generated texts to be used as experimental stimuli (see Figure 1). We then provided this information to ChatGPT-4 along with the initial prompt described in Figure 1. The first response did not meet all criteria, so we issued brief re-prompts (e.g., asking it to include experience analyzing sales data or shorten the text to under 100 words) to ensure the output fully satisfied the requirements. The finalized version was saved as the “advanced” AI-generated cover letter. To create the “simpler” version, we provided ChatGPT with the finalized complex text and instructed it to simplify the language to approximately a 7th-grade Flesch–Kincaid reading level while preserving the content. This process was repeated until both readability and content requirements were satisfied. The resulting text was saved as the “simpler” AI-generated cover letter.

Video clip demonstrating AI-generated letter: OpenAI provides configuration functions that allow ChatGPT to be customized for specific uses. Using these functions behind the scenes, we created two ChatGPT environments, each set to display either the complex or the simple letter that we had pre-generated through prompting. Regardless of the prompt entered, each version was configured to return its assigned text, while the interface displayed the same instructions participants received for their pre-treatment task. We recorded this process with screen capture to present to participants, creating the impression of live text generation.

3.3. Experimental Procedures

[Insert Figure 2 here]

Upon joining the experiment, participants completed five steps. The detailed experimental procedures are shown in Figure 2 and described below. The average completion time was 15.34 minutes.

1. **Pre-experiment survey:** After accepting the study on Prolific, participants were directed to a Qualtrics survey. They were informed that they would complete three writing tasks followed by survey questions. Participants indicated whether English was their first language. If not, they specified their primary language. They then completed measures of writing self-efficacy and their use of and trust in commercial LLM tools for writing tasks.
2. **Pre-treatment task:** Participants assumed the role of a fictional persona and wrote a 100-word job application cover letter based on background information provided by the researchers. They were instructed to complete the letter independently and without external tools (See Figure 2a).
3. **Treatment:** One of the pre-generated AI letters (simple vs. advanced complexity), along with a short video clip showing ChatGPT producing a sample text from the same prompt, was randomly presented to participants (See Figures 2b and 2c). After viewing the stimuli, participants rated the perceived superiority of the AI-generated text compared to their own draft.
4. **Post-treatment task:** Participants rewrote the cover letter for the same persona. At the top of the screen, they saw two reference texts: their original draft and the AI-generated letter. They could copy, revise, or resubmit either letter. A minimum time requirement was enforced to ensure attention (See Figure 2d).
5. **Post-experiment survey:** Finally, participants completed demographic questions and additional measures.

3.4. Measures

Textual Similarity: To capture how much participants accepted and integrated AI-generated content into their writing, we measured textual similarity, or the extent to which

participants composed their final letters similarly to the AI-generated text when they were allowed to reference it. Using the Python machine learning package, scikit-learn’s Term Frequency–Inverse Document Frequency (TF-IDF) vectorization (Pedregosa et al., 2011), we quantified the textual similarity between the cover letter each participant submitted for the final task and the AI-generated letter provided as a reference. Cosine similarity based on TF-IDF vectorization is one of the most widely used natural language processing techniques for capturing surface-level textual overlap (Manning et al., 2008). It measures similarity by comparing word usage patterns, taking into account both how frequently a word appears in a specific document (term frequency) and how rare it is across all documents (inverse document frequency) (see Appendix for details). Since cover letters typically follow a standardized structure and participants were given the same background information about the fictional character, cosine similarity is well-suited for detecting meaningful surface-level textual overlap. In this context, we are primarily interested in the degree of lexical reuse rather than deeper semantic or stylistic variation. Similarity scores originally ranged from 0 to 1 and were rescaled to a 0–100 scale for interpretation ($M = 69.3$, $SD = 18.7$). To check the validity of the similarity scores, we randomly selected 10 samples and qualitatively compared the participant-written letters with the AI-generated letters. Higher scores indicate greater reuse of words or phrases from the AI-generated letter. In the same way, we calculated similarity scores between participants’ original drafts, submitted before exposure to the AI-generated text, and the AI-generated text ($M = 69.3$, $SD = 18.7$).

Self-Efficacy for English Writing: Participants’ self-efficacy for writing was measured using the Self-Efficacy for Writing Scale (SEWS) developed by Bruning et al. (2013). The original scale consists of 16 items across three sub-constructs: ideation, conventions, and self-regulation. Because the self-regulation items emphasized general behavioral regulation rather than English writing, we focused on the two sub-constructs more directly related to English writing skills. For brevity, we selected three items from each of these sub-constructs with the highest factor loadings (Bruning et al., 2013). Example items include: “I can think of many ideas for my writing” and “I can write grammatically correct sentences.” Participants answered them on 7-point Likert scales. The selected items demonstrated acceptable internal consistency (Cronbach’s $\alpha = .914$). For analysis, we computed the average score across the

6 items ($M = 52.9$, $SD = 11.1$).

Perceived Superiority of AI-Generated Letter: Drawing on prior research that measures individuals’ perceived performance evaluation of algorithmic outputs compared to human performance (Logg et al., 2019; Lee, 2018), we asked participants to evaluate the AI-generated cover letter in comparison to the one they had written in the first task. They responded to the question: “Compared to the cover letter you wrote, how would you evaluate the quality of the letter written by Editor Bot?” Responses were recorded on a 7-point scale ranging from “Much Worse than Mine (-3)” to “Much Better than Mine (3).” ($M = .98$, $SD = 1.55$)

Performance Trust in Commercial LLMs: We assessed participants’ trust in commercial LLMs for English writing performance by adapting sub-scales from the Multi-Dimensional Measure of Trust (Malle and Ullman, 2021), specifically targeting capability trust in commercial LLMs. Participants were asked to evaluate a commercial generative AI chatbot (e.g., ChatGPT, Gemini, etc) they had used most within the past month in the context of English writing. Using six-point scales (0-5), they rated how well the following six adjectives described the AI tool: reliable, consistent, predictable, capable, skilled, and competent ($M = 3.46$, $SD = 1.12$, Cronbach’s $\alpha = .961$).

3.5. Analysis

After preprocessing data, we conducted Ordinary Least Squares (OLS) regression and a moderation analysis using SPSS PROCESS (Model 1) to test H1 and H2 (Hayes, 2013). To address H3, we employed a three-way moderation model (Model 3) in SPSS PROCESS. For H4, we performed a moderated mediation analysis using SPSS PROCESS (Model 12).

Demographic variables such as gender, education, income, employment status, and performance trust of LLMs were initially included as control variables. Among them, only performance trust in LLMs was statistically significant and was retained as a control variable in all models. The detailed characteristics of NNEWs and NEWs are reported separately in Table 1.

[Insert Table 1 here]

4. Results

[Insert Table 2 here]

4.1. Experimental Effects

To check whether the experimental stimulus we prepared affected participants' writing in the post-treatment task, we compared the textual similarity of the original draft from the pre-treatment task with the revised letter from the post-treatment task. A paired-samples t-test revealed a significant difference between the two ($t = -15.84$, $p < .001$; $M = 52.90$, $SD = 11.05$ for original draft; $M = 69.30$, $SD = 18.66$ for revised letter), indicating that our experimental manipulation was successful.

To further ensure the effects of the experimental treatments, we also adjusted the text similarity scores by subtracting the score for original drafts from the score for the revised letters and then replicated the analyses. There was no significant difference between the models based on either the raw text similarity scores or the adjusted scores. For ease of interpretation, we report the results based on the raw text similarity scores.

Table 2 presents the results of regression analyses examining the effects of nativity, lexical complexity of the AI-generated text, and their interaction on textual similarity. After accounting for performance trust in generative AI as a control variable, as shown in Block 2 in Table 2, nativity did not have a significant main effect on textual similarity, indicating no overall difference between non-native and native English writers in how much they integrated the AI-generated letter into their own writing. Therefore, H1, which proposed that NNEWs would produce letters more similar to AI-generated content than NEWs, was not supported.

In response to H2, which examined the two-way interaction between language nativity and the lexical complexity of AI-generated content, we found a significant interaction effect ($b = 9.645$, $p < .05$), as shown in Block 3 in Table 2. Thus, H2 was supported. We probed this interaction using the SPSS PROCESS macro (see Figure 2 for details). Specifically, when participants were presented with the simplified version of the AI-generated letter, NNEWs produced cover letters that were, on average, about 7% more similar to the AI-generated content than those written by NEWs ($p < .05$). However, when the advanced version was presented, no significant difference was observed between the two groups.

[Insert Figure 3 here]

4.2. Moderation Effect of Writing Efficacy

In H3, we anticipated that writing efficacy moderated the effect of language nativity and lexical complexity on textual similarity. Because we found a significant interaction between nativity and lexical complexity but not their main effects, we examined a three-way interaction among writing efficacy, nativity, and lexical complexity. The results are reported in Block 4 in Table 2. We found a significant three-way interaction among nativity, lexical complexity, and writing efficacy ($b = -12.856$, $p < .01$). H3 was supported.

[Insert Figure 4 here]

To probe this three-way interaction, we divided participants into two groups based on their writing efficacy: those with scores below the mean and those with scores above the mean (see Figure 4). A significant difference emerged only among participants with relatively low writing efficacy. Specifically, among those whose writing efficacy was below the mean, NNEWs produced letters that were, on average, more than 10% more similar to the AI-generated content than those written by NEWs, when the AI-generated letter used simple language ($p < .05$). In contrast, when the AI-generated text was written in advanced language, NEWs were more likely than NNEWs to produce letters that were more similar to the AI-generated content by over 10% on average. No significant differences were observed between the two groups when writing efficacy was high ($p < .05$).

4.3. Mediation Effect of Perceived Superiority

Through H4, we proposed that the effects of nativity on integrating AI-generated content into participants' own writing are mediated by perceived superiority, even when its effects interact with lexical complexity and writing self-efficacy. Since we found significant three-way interaction effects among nativity, lexical complexity, and writing self-efficacy, we conducted a mediated moderation analysis using SPSS PROCESS. Results are presented in Figure 5.

[Insert Figure 5 here]

Although perceived superiority was a strong predictor of textual similarity ($b = 3.617$, $p < .001$), we found no significant association between the three-way interaction term (nativity $\times$ lexical complexity $\times$ writing efficacy) and perceived superiority. The three-way

interaction remained a significant predictor of textual similarity even after including perceived superiority as an additional covariate in the model. Moreover, the coefficients for the interaction terms were similar with and without perceived superiority in the model. These findings indicate that perceived superiority did not mediate the experimental or moderated effects identified in the earlier analyses. H4 was not supported.

5. Discussion

Our study examined differences in AI content adoption between NNEWs and NEWs to understand how gaps in English writing skills and social constraints associated with linguistic nativity influence their use of AI and AI-assisted content production. Results show that NNEWs and NEWs differed in how they evaluated, accepted, and integrated AI-generated content when writing job application cover letters. By testing writing self-efficacy as a moderator and perceived superiority of AI-generated content as a mediator, we found that these patterns reflect both task-related proficiency gaps and sociolinguistic positioning tied to English nativity, which influence individuals' reliance on external support. Taken together, our findings suggest that pre-existing differences in writing between NNEWs and NEWs are likely to persist even with AI use, emphasizing the need for scaffolding efforts that not only build technical writing skills but also address the sociocultural constraints individuals face.

5.1. Interpretation of the Results

In our experiment, we did not find support for H1, which predicted a general difference in how NNEWs and NEWs accepted and integrated AI-generated content into their writing. Instead, we found that the difference between NNEWs and NEWs depended on the lexical complexity of the AI-generated content, supporting H2 regarding a two-way interaction between nativity and lexical complexity. When the AI-generated content was written at a simpler level (approximately 7th-grade), NNEWs were more likely than NEWs to adopt and integrate it into their writing. However, when the content was written at a more advanced level (approximately college graduate level), both groups adopted and integrated the AI-generated content to a similar extent. These findings support prior research suggesting

that NNEWs tend to rely on AI-generated content in writing tasks (Ito et al., 2023; Wang et al., 2024), and further reveal that NEWs also rely on AI assistance when AI demonstrates competent performance.

Given that writing self-efficacy reflects one’s perceived competence in writing and is conceptually separated from linguistic nativity (Mitchell et al., 2023; Sun et al., 2021), we further examined whether the differences in AI integration between NNEWs and NEWs resulting from the experimental manipulation (i.e., lexical complexity of AI-generated content) were moderated by individuals’ writing self-efficacy. Specifically, to address H3, we tested a three-way interaction: whether writing self-efficacy moderated the interaction between nativity and lexical complexity in predicting the extent of AI integration. Our findings indicate that among individuals with high writing self-efficacy, both NNEWs and NEWs showed similar patterns of writing characterized by relatively low textual similarity with AI-generated content. However, among individuals with lower writing self-efficacy, both groups tended to adopt AI-generated content for their writing tasks, but in different ways. When presented with AI-generated content written in simpler language, NNEWs were more likely than NEWs to integrate this AI content into their writing. In contrast, when the AI-generated content featured more advanced vocabulary, NEWs were more likely than NNEWs to produce writing more similar to the AI content. These findings confirm earlier research showing that individuals with low writing self-efficacy are generally more receptive to AI-generated content (Dumbauld et al., 2014; Rahmawati, 2019; Zainab et al., 2017) and further support Brynjolfsson et al. (2025)’s argument that AI benefits lower-skilled workers more. However, we also found that within the low-writing-efficacy group, NNEWs and NEWs differed in their preferences for lexical complexity. NNEWs tended to integrate simpler AI-generated text, while NEWs were more selective and showed a stronger preference for content with higher linguistic sophistication. This suggests that sociolinguistic gaps, stemming from social and cultural backgrounds and experiences among NNEWs and NEWs, influence AI use in ways that extend beyond gaps in writing proficiency.

Our final analysis focused on whether perceived superiority of AI-generated content mediated the interactive relationship among nativity, lexical complexity, and writing self-efficacy (H4). We hypothesized that group differences would be explained by participants’ percep-

tions of AI-generated content as superior to their own writing, such that NNEWs would view AI-generated content as better even when it used simpler vocabulary, while NEWs would perceive it as superior primarily when it featured more complex lexicon. However, we did not find support for the mediation effect of perceived superiority. Although participants who rated AI-generated content as better than their own were generally more likely to adopt it, perceived superiority did not fully account for the differences observed between NNEWs and NEWs. In fact, in our analysis, NEWs with low writing self-efficacy did not perceive AI-generated content with advanced lexicon as superior. Similarly, NNEWs, while more likely than NEWs to see AI-generated content as better overall, did not consistently judge the simpler version as superior to their own writing, even though they were more inclined to integrate it into their revisions. These findings suggest that NNEWs and NEWs did not adopt AI-generated content simply because they believed it was better than their own.

5.2. Theoretical and Practical Implications

Our study advances the literature on the effects of AI use on task performance by bridging two dominant frameworks: skill-biased technological change and digital divide research. Prior studies of AI adoption have largely emphasized either 1) skill gaps between novice and expert workers (Brynjolfsson et al., 2025) or 2) disparities in access and general usage across socioeconomic groups (Bassignana et al., 2025; Humlum and Vestergaard, 2025). By focusing on linguistic nativity, we connect these two areas of scholarship and show that differences in AI-assisted writing between NNEWs and NEWs cannot be explained by skill differences alone. Instead, they are conditioned by both perceived competence and sociolinguistic identity. On one hand, our findings extend skill-biased technological change research by demonstrating that disproportionate effects of AI on task performance arise not only from task proficiency but also from social and structural constraints. On the other hand, we advance digital divide research beyond access and general usage to focus on professional writing outcomes. Our contributions provide a more nuanced account of how AI may reproduce existing gaps related to writing skills and sociolinguistic positioning in professional contexts.

Additionally, our study raises an important unresolved question: Why do native and non-

native English writers prefer different types of AI-generated content for the same writing task, even when they both perceive the AI-generated text as superior to their own? We propose that the key to understanding this difference lies in how individuals evaluate AI-generated content as a form of scaffolding, which is temporary support that helps bridge the gap between current ability and task demands. According to sociocultural learning theory (Vygotsky, 1978), individuals are more likely to accept external assistance and complete new tasks when the support they receive falls within their ZPD. The ZPD refers to the cognitive space where individuals can perform tasks with guidance: tasks that are slightly beyond their independent capabilities, but not so difficult that they become unattainable. Vygotsky (1978) emphasized the importance of scaffolding in learning, arguing that support from more knowledgeable others can enhance an individual’s ability to complete tasks. However, for scaffolding to be effective, the assistance must be appropriately matched to the learner’s current ability (Palincsar et al., 2018). If the help from technology is perceived as either too simplistic or too advanced, the learner may not engage with or accept the assistance. When it comes to AI assistance, the key is that the support should fall just beyond the learner’s current capabilities so that it fosters confidence and learning without causing undue frustration (Kim and Hannafin, 2011).

We suggest that ZPD may help interpret our findings regarding the non-significant mediating effect of perceived superiority on the acceptance and integration of AI-generated content. When AI-generated content appears overwhelmingly superior, individuals may not perceive it as a useful or relatable reference for their own writing. In such cases, the AI-assisted output may feel too far removed from their personal capabilities, making it less likely to be adopted. It is also highly plausible that NNEWs and NEWs operate within different ZPDs in English writing. Prior research suggests that NEWs and NNEWs often employ different writing styles, regardless of fluency levels (Maier, 1992; Tian et al., 2025). In this context, NNEWs may have felt more comfortable integrating AI-generated content with simpler lexicons because it aligned more closely with their typical writing patterns, allowing them to feel they were enhancing their writing with the assistance of AI, rather than merely copying it. Conversely, NEWs may have found content with more advanced vocabulary more compatible with their writing style and capabilities, making it more acceptable

and useful for integration. Because our study does not directly operationalize the concept of ZPD, it offers only preliminary evidence that the adoption and perceived usefulness of AI assistance may relate to individuals' ZPD. This zone may be shaped through interactions between cultural and linguistic identity, writing self-efficacy, and actual writing capabilities. Nonetheless, our findings provide important implications for future research regarding the need to conceptualize AI assistance in the production task as a form of scaffolding that users are more likely to accept when it falls within their "sweet spot".

Finally, our study highlights the need for AI tools and educational programs designed especially for NNEWs with lower proficiency and skills. Our findings showed that NNEWs with lower writing efficacy may benefit less than NEWs with similarly low efficacy. This is because NNEWs tend to rely more on simpler lexical choices, even though advanced vocabulary is often expected in professional writing contexts. To be effective, AI assistants for professional tasks should function as scaffolding, offering support that aligns with users' current abilities and pushes them just enough to promote growth without overwhelming them. In the case of professional writing, NNEWs may need higher-order writing strategies as well, as they are more inclined to adopt AI-generated content that focuses more on foundational skills such as grammar and potentially miss out on the importance of higher-order writing strategies that are cognitively manageable for them. AI literacy programs should also include guidance on how to critically evaluate, adapt, and integrate AI-generated content in ways that support personal skill development. In our study, NNEWs appeared less likely to demonstrate the nuanced rhetorical and strategic elements typically expected in professional writing. This suggests that educational programs should raise awareness of the continued importance of areas where humans currently hold a clear advantage over AI, such as human judgment, contextual sensitivity, and the ability to communicate intent with nuance and audience awareness. Designing AI literacy programs with these goals in mind can help ensure that generative AI empowers users rather than displacing uniquely human competencies.

5.3. Limitations

We acknowledge several limitations of this study. First, the experiment necessarily simplified user–AI interactions to control for confounding factors. For example, we used pre-generated AI texts created through detailed prompts specifying applicant information, word length, and lexical level to ensure that all participants were exposed to the same content. In real-world use, however, generative AI does not reliably reproduce identical outputs, and individuals typically interact with the system iteratively, refining prompts, tailoring outputs, and responding to feedback. While standardization strengthened experimental control in our design, it also limits reproducibility in naturalistic contexts. Future research should examine dynamic, ongoing user–AI interactions and how these evolve in relation to users’ social and cultural backgrounds. Second, this study focused on a specific professional writing task: composing a job application cover letter. This task was selected because writing is a foundational function of LLMs, and job applications represent high-stakes, high-engagement use cases where people are especially motivated to seek AI assistance. However, the task was fictionalized; participants wrote on behalf of an imagined person, not themselves. Future research should aim to replicate these findings in real-world contexts where participants generate writing for their own job applications. Third, while we focused on writing, LLMs are used in a wide range of productive tasks beyond writing, such as coding, design, and video generation, which are domains that often require specialized expertise. Investigating whether AI use in these other domains influences social skill gaps would offer a broader understanding of AI’s role in shaping digital inequality. Finally, although we propose ZPD as a useful framework for interpreting differential reliance on AI, we did not directly measure this construct. Future work should aim to operationalize and measure ZPD in the context of human–AI interaction to better test its role in mediating AI acceptance and use.

5.4. Conclusion

Despite its limitations, our study contributes to a growing body of research examining how AI use supports the productivity and efficiency of individuals who live within existing social constraints. Our study is among the few that explore how AI adoption and use intersect with sociolinguistic factors and perceived task proficiency, focusing on English nativity in

the U.S. context. Overall, our findings do not fully support the claim that AI benefits low-skilled workers more (Brynjolfsson et al., 2025). Rather, we demonstrate that AI adoption is embedded in users’ social experiences and identities, and in how individuals perceive their competence and access to support. We also propose a scaffolding approach to AI adoption, emphasizing that people adopt and integrate AI-generated content in ways that align with their current skill levels. Future research should continue to investigate these complex social and cultural dynamics, taking into account how factors such as individuals’ skills, social identity, access to resources, and perceived legitimacy of AI assistance shape engagement with AI tools.

References

- Autor, D.H., Katz, L.F., Krueger, A.B., 1998. Computing inequality: Have computers changed the labor market? *The Quarterly Journal of Economics* 113, 1169–1213. doi:10.1162/003355398555874.
- Bandura, A., 1997. *Self-efficacy: The Exercise of Control*. W. H. Freeman, New York, NY.
- Bassignana, E., Curry, A.C., Hovy, D., 2025. The ai gap: How socioeconomic status affects language technology interactions. URL: <https://arxiv.org/abs/2505.12158>, arXiv:2505.12158.
- Batalova, J., Zong, J., 2016. *Language Diversity and English Proficiency in the United States*. Report. Migration Policy Institute. Washington, DC.
- Bresnahan, T.F., Brynjolfsson, E., Hitt, L.M., 2002. Information technology, workplace organization, and the demand for skilled labor: Firm-level evidence. *The quarterly journal of economics* 117, 339–376. doi:10.1162/003355302753399526.
- Bruning, R., Dempsey, M., Kauffman, D.F., McKim, C., Zumbrunn, S., 2013. Examining dimensions of self-efficacy for writing. *Journal of educational psychology* 105, 25. doi:10.1037/a0029692.
- Brynjolfsson, E., Li, D., Raymond, L., 2025. Generative AI at work. *The Quarterly Journal of Economics* 140, 889–942. doi:10.1093/qje/qjae044.
- Cai, L., Msafiri, M.M., Kangwa, D., 2024. Exploring the impact of integrating AI tools in higher education using the zone of proximal development. *Education and Information Technologies* , 1–74.
- Cheng, Y.S., 2004. A measure of second language writing anxiety: Scale development and preliminary validation. *Journal of second language writing* 13, 313–335. doi:10.1016/j.jslw.2004.07.001.
- Choung, H., David, P., Ross, A., 2023. Trust in AI and its role in the acceptance of AI technologies. *International Journal of Human–Computer Interaction* 39, 1727–1739.

- Daepp, M.I.G., Counts, S., 2024. The emerging ai divide in the united states. CoRR abs/2404.11988. URL: <https://doi.org/10.48550/arXiv.2404.11988>.
- Davis, F.D., 1985. A technology acceptance model for empirically testing new end-user information systems: Theory and results. Ph.D. thesis. Massachusetts Institute of Technology.
- Davis, F.D., 1989. Perceived usefulness, perceived ease of use, and user acceptance of information technology. MIS Quarterly 13, 319–340. doi:10.2307/249008.
- Dell’Acqua, F., McFowland III, E., Mollick, E.R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraymer, L., Candelon, F., Lakhani, K.R., 2023. Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality. Working Paper 24-013. Harvard Business School Technology & Operations Mgt. Unit & The Wharton School. doi:10.2139/ssrn.4573321. available at SSRN: <https://ssrn.com/abstract=4573321>.
- van Deursen, A.J.A.M., Helsper, E.J., 2015. The third-level digital divide: Who benefits most from being online?, in: Communication and Information Technologies Annual. Emerald Group Publishing Limited. volume 10, pp. 29–52. doi:10.1108/S2050-206020150000010002.
- Dietrich, S., Hernandez, E., 2022. Nearly 68 million people spoke a language other than english at home in 2019. U.S. Census Bureau Stories Online. Published December 6, 2022. Available at <https://www.census.gov/library/stories/2022/12/languages-we-speak-in-united-states.html>.
- DiMaggio, P., Hargittai, E., 2023. From the ‘Digital Divide’ to ‘Digital Inequality’: Studying internet use as penetration increases. URL: <https://doi.org/10.31235/osf.io/rhqmu>, doi:10.31235/osf.io/rhqmu.
- Dumbauld, J., Black, M., Depp, C.A., Daly, R., Curran, M.A., Winegarden, B., Jeste, D.V., 2014. Association of learning styles with research self-efficacy: Study of short-term research training program for medical students. Clinical and translational science 7, 489–492.

- Dutton, W.H., Blank, G., Groselj, D., 2013. Cultures of the Internet: The Internet in Britain. Oxford Internet Survey 2013 Report. Technical Report. Oxford Internet Institute, University of Oxford. Oxford, United Kingdom. URL: <http://oxis.oii.ox.ac.uk/>.
- Goldin, C.D., Katz, L.F., 2008. The race between education and technology. harvard university press.
- Gonzales, A., 2016. The contemporary us digital divide: from initial access to technology maintenance. *Information, Communication & Society* 19, 234–248. doi:10.1080/1369118X.2015.1050438.
- Gunaratne, J., Zalmanson, L., Nov, O., 2018. The persuasive power of algorithmic and crowdsourced advice. *Journal of Management Information Systems* 35, 1092–1120. doi:10.1080/07421222.2018.1523534.
- Hampton, K.N., 2016. Persistent and pervasive community: New communication technologies and the future of community. *American Behavioral Scientist* 60, 101–124. doi:10.1177/0002764215601714.
- Hayes, A.F., 2013. Introduction to mediation, moderation, and conditional process analysis: A regression-based approach. Guilford Press, New York, NY.
- Hou, H.I., Li, M.Y., 2011. A contrastive rhetoric analysis of internship cover letters written by taiwanese and canadian hospitality majors. *International Journal of Linguistics* 3, 43–57. doi:10.5296/ijl.v3i1.1135.
- Humlum, A., Vestergaard, E., 2025. The unequal adoption of chatgpt exacerbates existing inequalities among workers. *Proceedings of the National Academy of Sciences* 122, e2414972121. doi:10.1073/pnas.2414972121.
- Ito, T., Yamashita, N., Kuribayashi, T., Hidaka, M., Suzuki, J., Gao, G., Jamieson, J., Inui, K., 2023. Use of an ai-powered rewriting support software in context with other tools: A study of non-native english speakers, in: *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology (UIST '23)*, Association for Computing Machinery. pp. Article 45, 1–13. doi:10.1145/3586183.3606810.

- Jansen, J.A., Manukyan, A., Al Khoury, N., Akalin, A., 2025. Leveraging large language models for data analysis automation. *PloS one* 20, e0317084. doi:10.1371/journal.pone.0317084.
- Jelson, A., Lee, S.W., 2024. An empirical study to understand how students use chatgpt for writing essays and how it affects their ownership, in: *Proceedings of the Third Workshop on Intelligent and Interactive Writing Assistants*, Association for Computing Machinery, New York, NY, USA. p. 26–30. doi:10.1145/3690712.3690720.
- Jeong, S., Kim, S., Lee, S., 2024. Effects of perceived ease of use and perceived usefulness of technology acceptance model on intention to continue using generative AI: Focusing on the mediating effect of satisfaction and moderating effect of innovation resistance, in: *International Conference on Conceptual Modeling*, pp. 99–106. doi:10.1007/978-3-031-75599-6_7.
- Katz, L.F., Murphy, K.M., 1992. Changes in relative wages, 1963–1987: supply and demand factors. *The quarterly journal of economics* 107, 35–78.
- Kim, M.C., Hannafin, M.J., 2011. Scaffolding problem solving in technology-enhanced learning environments (teles): Bridging research and theory with practice. *Computers & Education* 56, 403–417. doi:10.1016/j.compedu.2010.08.024.
- Lee, J.D., See, K.A., 2004. Trust in automation: Designing for appropriate reliance. *Human factors* 46, 50–80. doi:10.1518/hfes.46.1.50.30392.
- Lee, M.K., 2018. Understanding perception of algorithmic decisions: Fairness, trust, and emotion in response to algorithmic management. *Big Data & Society* 5, 1–16. doi:10.1177/2053951718756684.
- Livingstone, S., Helsper, E., 2007. Gradations in digital inclusion: children, young people and the digital divide. *New Media & Society* 9, 671–696. doi:10.1177/1461444807080335.
- Logg, J.M., Minson, J.A., Moore, D.A., 2019. Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes* 151, 90–103. doi:10.1016/j.obhdp.2018.11.010.

- Lund, S., Manyika, J., Chui, M., Bughin, J., Woetzel, J., Krishnan, A.M., Seong, J., Muir, M., George, K., 2023. The Economic Potential of Generative AI: The Next Productivity Frontier. Technical Report. McKinsey Global Institute. URL: <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier>.
- Maier, P., 1992. Politeness strategies in business letters by native and nonnative english speakers. *English for Specific Purposes* 11, 189–205. doi:10.1016/S0889-4906(05)80009-2.
- Malle, B.F., Ullman, D., 2021. Chapter 1 - a multidimensional conception and measure of human-robot trust, in: Nam, C.S., Lyons, J.B. (Eds.), *Trust in Human-Robot Interaction*. Academic Press, pp. 3–25. doi:<https://doi.org/10.1016/B978-0-12-819472-0.00001-0>.
- Manning, C.D., Raghavan, P., Schütze, H., 2008. *Introduction to Information Retrieval*. Cambridge University Press.
- Mendez, G., Galárraga, L., Chiluiza, K., 2021. Showing academic performance predictions during term planning: effects on students’ decisions, behaviors, and preferences, in: *Proceedings of the 2021 CHI conference on human factors in computing systems*, pp. 1–17. doi:<https://doi.org/10.1145/3411764.344571>.
- Mitchell, K.M., Zumbrunn, S., Berry, D.N., Demczuk, L., 2023. Writing self-efficacy in postsecondary students: A scoping review. *Educational Psychology Review* 35, 82.
- Nazari, N., Ebersole, M., Schunn, C.D., Koedinger, K.R., 2021. Application of artificial intelligence powered digital writing assistant in higher education: Randomized controlled trial. *Heliyon* 7, e07014. doi:10.1016/j.heliyon.2021.e07014.
- Norris, P., 2001. *Digital Divide: Civic Engagement, Information Poverty, and the Internet Worldwide*. Cambridge University Press.
- Noy, S., Zhang, W., 2023. Experimental evidence on the productivity effects of generative artificial intelligence. *Science* 381, 187–192.

- Pajares, F., 2007. Empirical properties of a scale to assess writing self-efficacy in school contexts. *Measurement and evaluation in Counseling and development* 39, 239–249.
- Pajares, F., Johnson, M.J., 1994. Confidence and competence in writing: The role of self-efficacy, outcome expectancy, and apprehension. *Research in the Teaching of English* 28, 313–331.
- Palincsar, A.S., Fitzgerald, M.S., Marcum, M.B., Sherwood, C.A., 2018. Examining the work of “scaffolding” in theory and practice: A case study of 6th graders and their teacher interacting with one another, an ambitious science curriculum, and mobile devices. *International Journal of Educational Research* 90, 191–208. URL: <https://doi.org/10.1016/j.ijer.2017.11.006>, doi:10.1016/j.ijer.2017.11.006.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., Édouard Duchesnay, 2011. Scikit-learn: Machine learning in python. *Journal of Machine Learning Research* 12, 2825–2830. URL: <http://jmlr.org/papers/v12/pedregosa11a.html>.
- Peng, S., Kalliamvakou, E., Cihon, P., Demirer, M., 2023. The impact of AI on developer productivity: Evidence from github copilot. *arXiv preprint arXiv:2302.06590*. URL: <https://arxiv.org/abs/2302.06590>, doi:10.48550/arXiv.2302.06590. accessed: 6 Jan 2025.
- Piller, I., 2016. *Linguistic Diversity and Social Justice: An Introduction to Applied Sociolinguistics*. Oxford University Press. URL: <https://books.google.com/books?id=K7V1CwAAQBAJ>.
- Prat-Sala, M., Redford, P., 2012. Writing essays: Does self-efficacy matter? the relationship between self-efficacy in reading and in writing and undergraduate students’ performance in essay writing. *Educational psychology* 32, 9–20. doi:0.1080/01443410.2011.621411.
- Rahmawati, R.N., 2019. Self-efficacy and use of e-learning: A theoretical review technology

- acceptance model (TAM). *American Journal of Humanities and Social Sciences Research* 3, 41–55. URL: <https://www.ajhssr.com/wp-content/uploads/2019/10/Z1941055.pdf>.
- Reisdorf, B.C., Fernandez, L., Hampton, K.N., Shin, I., Dutton, W.H., 2020. Mobile phones will not eliminate digital and social divides: How variation in internet activities mediates the relationship between type of internet access and local social capital in detroit. *Social Science Computer Review* 40, 288–308. doi:10.1177/0894439320909446.
- Rogers, E.M., 2001. The digital divide. *Convergence* 7, 96–111. doi:10.1177/135485650100700406.
- Sætra, H.S., 2025. Scaffolding human champions: AI as a more competent other. *Human Arenas* 8, 56–78. doi:doi.org/10.1007/s42087-022-00304-8.
- Schwartz, A.E., Stiefel, L., 2006. Is there a nativity gap? new evidence on the academic performance of immigrant students. *Education Finance and Policy* 1, 17–49. doi:10.1162/edfp.2006.1.1.17.
- Schwartzberg, J., 2025. Cover letters still matter—even if they’re not required. URL: <https://hbr.org/2025/03/cover-letters-still-matter-even-if-theyre-not-required>. Harvard Business Review (online).
- Shell, D.F., Murphy, C.C., Bruning, R.H., 1989. Self-efficacy and outcome expectancy mechanisms in reading and writing achievement. *Journal of educational psychology* 81, 91.
- Sidoti, O., Park, E., Gottfried, J., 2025. About a quarter of U.S. teens have used chatgpt for schoolwork, double the share in 2023. URL: <https://www.pewresearch.org/short-reads/2025/01/15/about-a-quarter-of-us-teens-have-used-chatgpt-for-schoolwork-double-the-share-in-2023/>. accessed: April 9 2025.
- Silva, T., 1993. Toward an understanding of the distinct nature of l2 writing: The esl research and its implications. *TESOL Quarterly* 27, 657–677. doi:10.2307/3587400.
- Sun, T., Wang, C., Lambert, R.G., Liu, L., 2021. Relationship between second language english writing self-efficacy and achievement: A meta-regression analysis. *Journal of Second Language Writing* 53, 100817. doi:<https://doi.org/10.1016/j.jslw.2021.100817>.

- Tian, X., Hahn, H.M., Leon Carangui, M.K., Si, Y., Chen, J., Li, X., Fussell, S.R., 2025. Assessing human evaluations of cover letters written or edited by ai and non-native english speakers, in: Companion Proceedings of the 2025 ACM International Conference on Supporting Group Work, Association for Computing Machinery, New York, NY, USA. p. 76–82. doi:10.1145/3688828.3699644.
- Vygotsky, L.S., 1978. *Mind in Society: The Development of Higher Psychological Processes*. volume 86. Harvard University Press.
- Wang, C., Zou, B., Du, Y., Wang, Z., 2024. The impact of different conversational generative ai chatbots on efl learners: An analysis of willingness to communicate, foreign language speaking anxiety, and self-perceived communicative competence. *System* 127, 103533. doi:10.1016/j.system.2024.103533.
- Williams, J.D., Takaku, S., 2011. Help seeking, self-efficacy, and writing performance among college students. *Journal of Writing Research* 3, 1–18. doi:10.17239/jowr-2011.03.01.1.
- Woodrow, L., 2011. College english writing affect: Self-efficacy and anxiety. *System* 39, 510–522. doi:10.1016/j.system.2011.10.017.
- Zainab, B., Awais Bhatti, M., Alshagawi, M., 2017. Factors affecting e-training adoption: An examination of perceived cost, computer self-efficacy and the technology acceptance model. *Behaviour & Information Technology* 36, 1261–1273. doi:10.1080/0144929X.2017.1380703.
- Zimmerman, B.J., Bandura, A., 1994. Impact of self-regulatory influences on writing course attainment. *American educational research journal* 31, 845–862.

Table 1: Participant Characteristics by Group (NNEWs= 184; NEWs= 143)

Characteristic	NNEWs	NEWs
Female	62.1%	60.6%
Non-white	81.3%	31.7%
Full-time employment	43.3%	50.7%
Bachelor’s degree or higher	64.7%	60.0%
Household income below \$75,000	44.5%	61.3%
Use of LLMs	98.8%	99.3%
Trust in LLMs (0–5 scale)	3.57 (1.12)	3.39 (1.10)
Writing self-efficacy (0–6 scale)	4.41 (1.12)	5.24 (0.86)
Perceived superiority of AI text (–3 to 3 scale)	1.20 (1.43)	0.70 (1.65)
Textual similarity: Revised letter vs. AI text (0–100 scale)	70.12 (18.78)	68.25 (18.51)

Values are means unless otherwise indicated; values in parentheses represent standard deviations.

LLMs include ChatGPT, Gemini, Claude, and others.

Table 2: OLS Regressions Predicting Textual Similarity

	Coefficient	SE	ΔR^2
Block 1: Controls			$\Delta.020^*$
Constant	61.161***	3.375	
AI Trust	2.356*	0.928	
Block 2: Main Effects			$\Delta.003$
Nativity (1 = Native)	-1.714	2.087	
Lexicon Complexity(1 = Advanced)	1.296	2.070	
Block 3: Two-way Interaction			$\Delta.016^*$
Nativity $\times$ Lexicon Complexity	9.645*	4.141	
Block 4: Three-way Interaction			$\Delta.042^*$
Writing Efficacy	-4.790**	1.709	
Nativity $\times$ Writing Efficacy	5.724	3.014	
Lexicon Complexity $\times$ Writing Efficacy	5.294	2.440	
Nativity $\times$ Lexicon Complexity $\times$ Writing Efficacy	-12.856**	4.284	
R²	.083		

* $p < .05$, ** $p < .01$, *** $p < .001$

	Advanced-lexicon Letter (College-level)	Simple-lexicon Letter (6-7th Grade)
Evaluation Criteria	- Include all requested information (e.g., a fictional persona with specified education and skills). - Exclude any extraneous information beyond what the researchers specified. - Be written in formal cover-letter format, including addresses. - Stay within 100 words. - Avoid contractions. - Meet college-level vocabulary standards based on Flesch-Kincaid readability	- Maintain the same contents as the advanced version. - Be written in formal cover-letter format, including addresses. - Stay within 100 words. - Use contractions. - Meet 6-7th grade vocabulary standards based on Flesch-Kincaid readability
Initial Prompt	[Applicant Information: Taylor Williams — Education: B.A. in Business from Best University, with a focus on marketing and sales management; Skills and Experience: Worked in electronic sales for 5 years, successfully introduced several new products to the market, and received the “Best Salesperson Award” in the department for 2 consecutive years. Strong communication and relationship-building skills with customers. Proficient in sales tracking programs such as Salesforce and HubSpot, and experienced in analyzing sales data to improve performance.] [Task: Imagine you are Taylor Williams. Write a 100-word cover letter for a Sales Representative position at ABC Electronic Inc.. Use college-level vocabulary, as defined by the Flesch-Kincaid Grade Level test.]	Please rewrite the following cover letter [Final output for advanced text] using 6th–7th grade vocabulary as defined by the Flesch-Kincaid Grade Level test. Make sure to maintain the content and keep the length at approximately 100 words.
Sample Refinement Prompts	- The experience in analyzing sales data to improve performance is missing. Please add it and rewrite the letter. - Please shorten the letter to within 100 words. - Please use college-level vocabulary, as defined by the Flesch-Kincaid Grade Level test. - Please do not use contractions.	- You also missed: “Taylor won the Best Salesperson Award for two consecutive years.” Include this. - Some words are above a 6–7th grade vocabulary level. Please edit the text to a 6–7th grade reading level (Flesch-Kincaid) while keeping the content the same. - Keep the letter at approximately 100 words. - Please use contractions.
Final Output	Hiring Manager ABC Electronics Inc. 123 Tech Drive Spring Field, CA 90210 Dear Hiring Manager, I am eager to apply for the Sales Representative position at ABC Electronics Inc. With a Bachelor of Arts in Business, specializing in Marketing and Sales Management from Best University, and five years of experience in electronics sales, I have developed a robust skill set in product introduction and client relationship management. My proficiency in sophisticated sales tracking systems such as Salesforce and HubSpot, coupled with my analytical skills in interpreting sales data, has enabled me to consistently exceed targets and earn the Best Salesperson Award for two consecutive years. I am enthusiastic about the opportunity to contribute my expertise and passion to your esteemed team. Sincerely, Taylor Williams	Hiring Manager ABC Electronics Inc. 123 Tech Drive Spring Field, CA 90210 Dear Hiring Manager, I'm excited to apply for the Sales Representative job at ABC Electronics Inc. I have a degree in Business from Best University, where I studied Marketing and Sales Management. I've spent five years selling electronics and am good at introducing new products. I enjoy talking to customers and building strong relationships with them. I know how to use sales programs like Salesforce and HubSpot to keep track of my sales and customer information. I can understand and use sales data to help me sell more. I've won the Best Salesperson Award for two years in a row. I look forward to bringing my skills and experience to your team. Sincerely, Taylor Williams

Figure 1: Process for generating a job application cover letter via ChatGPT

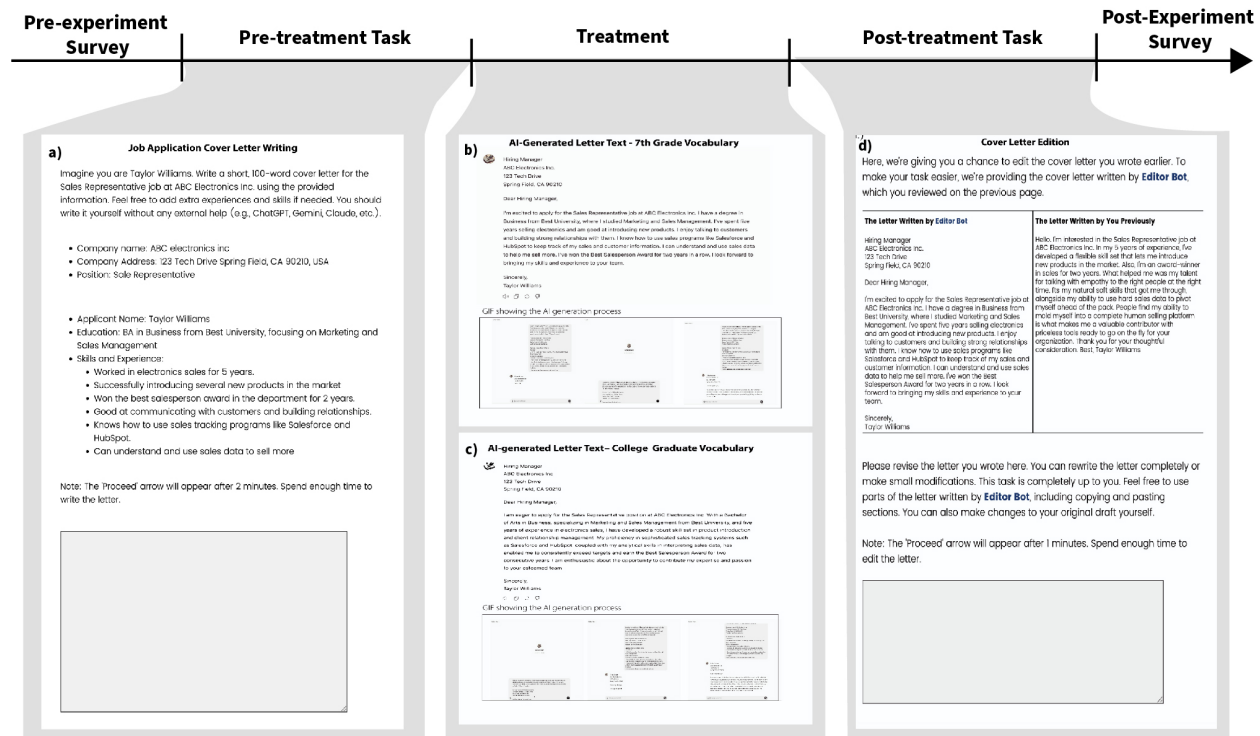


Figure 2: Timeline of Experiment. The study began with pre-experiment surveys on writing efficacy and generative AI use. Participants then completed a) a pre-treatment task: writing a cover letter for a fictional persona without external help. Next, they were randomly assigned to b) view the advanced text or c) the simple text, both pre-generated but presented as if produced by ChatGPT. Finally, d) they revised their original letter using both drafts as references before completing the study.

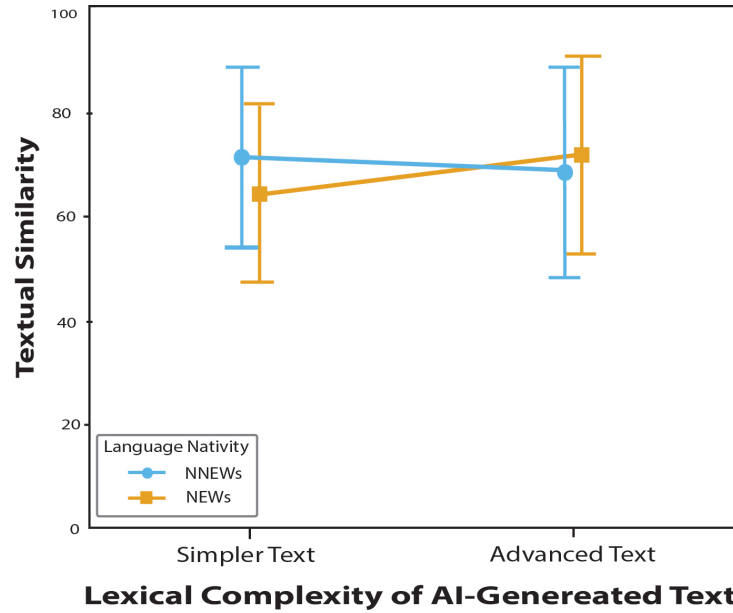


Figure 3: Two-way interaction between language nativity and the lexical complexity of AI-generated content. Each bar represents the mean adoption score for each group by content complexity (error bars = ± 1 SD).

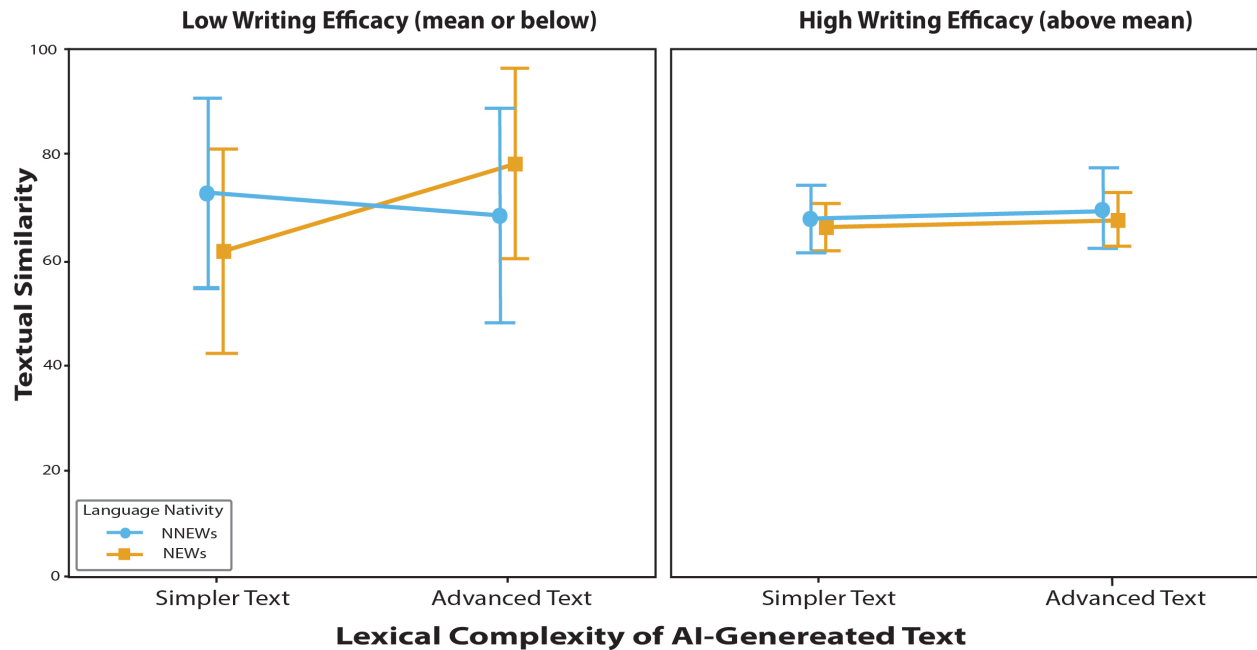


Figure 4: Three-way interaction among nativity, lexical complexity, and writing efficacy. Each bar represents the mean adoption score for each group by content complexity (error bars = ± 1 SD).

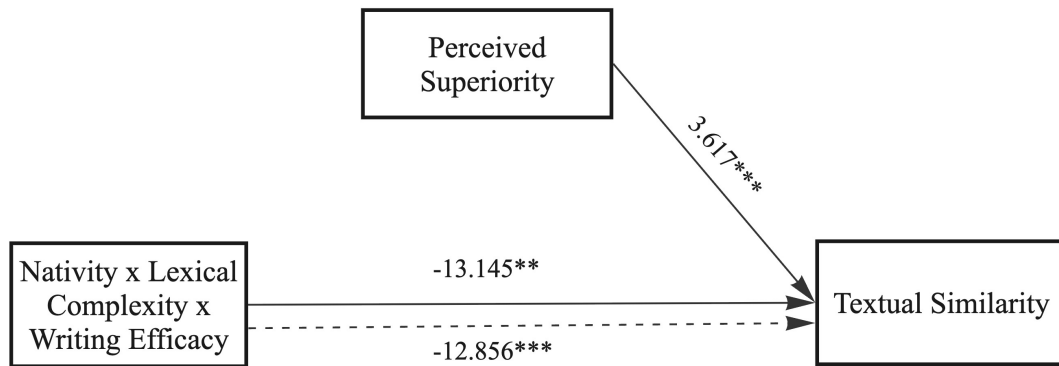


Figure 5: Mediated moderation model with perceived superiority mediating the interaction between Nativity and Text Level. Unstandardized coefficients shown; dashed line indicates interaction effect without the mediator. Trust in general LLMs was controlled.

Appendix A. Textual Similarity Calculation

To quantify the extent to which participants incorporated AI-generated content into their own writing, we calculated cosine similarity between the participant’s final submitted letter and the AI-generated reference letter. Both texts were first transformed into vector representations using the TF-IDF weighting scheme implemented in `scikit-learn`. TF-IDF is computed as follows:

$$\text{TF}_{t,d} = \frac{f_{t,d}}{\sum_k f_{k,d}} \quad (\text{A.1})$$

$$\text{IDF}_t = \log \left(\frac{N}{1 + n_t} \right) \quad (\text{A.2})$$

$$\text{TF-IDF}_{t,d} = \text{TF}_{t,d} \times \text{IDF}_t \quad (\text{A.3})$$

where $f_{t,d}$ is the frequency of term t in document d , $\sum_k f_{k,d}$ is the total number of terms in document d , N is the total number of documents, and n_t is the number of documents containing term t .

Once each document is vectorized, the cosine similarity between the two TF-IDF vectors $\vec{A}$ and $\vec{B}$ is computed as:

$$\text{CosineSimilarity}(\vec{A}, \vec{B}) = \frac{\vec{A} \cdot \vec{B}}{\|\vec{A}\| \|\vec{B}\|} \quad (\text{A.4})$$

This yields a similarity score ranging from 0 (completely dissimilar) to 1 (identical). For interpretability, we scaled the resulting score to a 0–100 range. This metric captures surface-level textual overlap based on shared vocabulary and word usage patterns, rather than semantic similarity.