



Using automated paraphrasing tools: Examining the grammatical structure of generated paraphrases of scientific abstracts

Alvin Ping Leong

Nanyang Technological University, Singapore

ABSTRACT

The proliferation of automated paraphrasing tools (APTs) raises questions about how generated paraphrases differ from the original texts. In this study, 50 abstracts published in Nature were compared with paraphrases generated by QuillBot, Jasper, and Copilot. Tactic and logico-semantic relations were analyzed using a modified version of the Hallidayan clause-complexing framework. The findings revealed that the APT that most closely matched Nature was QuillBot; no significant differences were found in any of the categories. Collectively, Jasper and Copilot leaned toward clausal complexity, using fewer paratactic extensions and more hypotactic elaborations. This study highlights general features of generated paraphrases that are not commonly addressed in the literature. Explicit paraphrasing instructions are recommended to avert any misuse of APTs.

ARTICLE HISTORY

Received
8 December 2024
Accepted
7 March 2025

KEYWORDS

paraphrasing; science
writing; clause
complexing; abstracts;
QuillBot; Jasper;
Microsoft Copilot.

1. Introduction

Paraphrasing is a familiar concept to students and teachers alike; it is typically understood as “using different words to express the same meaning” (Sharpe, 2024: 179). This definition of paraphrasing is common in the literature, appearing in both guidebooks for writers (Hopkins & Reid, 2024; Sharpe, 2024) and scholarly articles (Bhagat & Hovy, 2013; Sun & Yang, 2015). However, characterizing paraphrasing as

involving mere word-based changes is simplistic and may lead writers to make superficial changes to the source text.

Instances of uncited, superficial changes in scholarly articles have been pointed out in PubPeer, a forum where users review and discuss scientific papers, among other publications. These minor changes allowed PubPeer users to easily locate the original articles (Nerita vitiensis, 2024a, 2024b), suggesting that plagiarism had occurred. Several of the identified cases also contained *tortured phrases*, where awkward expressions were used in place of the original wording. These included *computerized reasoning* for the more widely accepted *artificial intelligence*, and *huge information* for *big data* (Muraltia serpylloides, 2024). Such phrases were quite likely generated by automated paraphrasing tools (APTs) in an effort to avoid plagiarism detection. In some cases, the use of tortured phrases resulted in article retractions (Kincaid, 2023; Marcus, 2024).

Ideally, then, paraphrasing should result in a properly cited version that is different enough from the original. Changing words and expressions forms only one aspect of the paraphrasing process; wherever possible, effort should also be taken to change the grammatical structure of the text. The grammatical aspect of paraphrasing was highlighted as far back as 1877, in Davidson and Alcock's landmark publication on the analysis of English and what the authors term *a treatise on Paraphrasing* (Davidson & Alcock, 1877: v). They note:

It should also be remembered that a good paraphrase does not consist in the mere substitution of one word for another, even though the meaning conveyed be precisely the same. Having mastered the full sense of the original passage, the student is to express the same ideas by using different language, that is, besides a mere change of words, there should also be a change of phrases and idioms and AN ALTERATION OF THE STRUCTURE OF THE SENTENCES. (Davidson & Alcock, 1877: 220, emphasis mine)

This short extract highlights two important points. First, paraphrasing requires the writer to comprehend the meaning (or *full sense*) of the original text. In this present age of artificial intelligence (AI) and APTs, however, no human writer is involved, and text comprehension is instead facilitated by algorithms that recognize patterns. But algorithms at present do not yet have the ability to comprehend messages in the way humans do (Priyadarshini & Cotton, 2022), and errors have been known to occur. An over-reliance on AI and algorithms may thus not only compromise a needed skill for writers but also result in algorithms learning incidental errors and propagating them in generated texts (Wu et al., 2019).

Second, paraphrasing involves not merely the replacement of words and phrases, but changes to the grammatical structure of the source text as well. While some studies have acknowledged this grammatical aspect (Keck, 2010; Yahia & Egbert, 2023), they are few in number. The issues covered have also tended to include clause-internal constituents such as the grammatical subject, main verb, and direct object (Keck, 2010), or grammatical errors (Liu & Lin, 2022). These considerations are naturally insightful when analyzing short paraphrases comprising a sentence or two. In academic writing, however, paraphrases tend to be longer, and other grammatical features, such as the use of different types of clauses and inter-clausal relations, should thus also be investigated.

Such grammatical issues in paraphrases are less obvious, as compared to word substitutions, which may partly account for the lack of studies in this area. The pervasiveness of online tools also warrants a closer look at how generated paraphrases compare with the original versions. In this paper, abstracts published in a top-tier journal were compared with paraphrases generated from three APTs. The original and paraphrased abstracts were analyzed for their use of clauses and inter-clausal relations based on an adapted version of Halliday's clause-complexing framework (Halliday & Matthiessen, 2014: 428–556).

The rest of this paper is organized as follows. Section 2 presents a literature review on scholarly work related to paraphrasing involving academic writing, including a brief description of APTs. Section 3 outlines the Hallidayan clause-complexing framework and related studies involving academic writing. Section 4 details the corpus and method of analysis. Section 5 discusses the findings, followed by a concluding section summarizing the key results and implications of the study.

2. Literature review

As a skill that requires writers to not merely understand the meaning of the source text but to reframe it to fit the larger discourse in which it is located, paraphrasing can be especially challenging to non-native users of English when writing a research paper (Shi et al., 2018; Yahia & Egbert, 2023). This often results in these writers opting for a path of least resistance, making minimal (and hopefully adequate) changes to the original passage to avoid plagiarism detection. Indeed, novice and L2 writers have been observed to rely on the source text, making more near copies of the original passage than L1 writers (Keck, 2006, 2014; Shi, 2004). In a more extreme case, Hirvela and Du (2013: 93) report how one of their interviewees, an L2 writer, contemplated abandoning paraphrasing entirely in favor of using direct quotations, which led to the following response from the researchers: “[t]his, we maintain, is where paraphrasing instruction is especially crucial.”

Providing paraphrasing instructions and guided tasks is perhaps the most direct way to help students better understand the mechanics of paraphrasing. Small-scale studies have shown that incorporating such instructions in a classroom setting is beneficial to some extent (Choy & Lee, 2012; McDonough et al., 2014; Wette, 2010), particularly to students identified as being highly motivated (Ahn, 2022). However, while the instructions in these studies are reported to cover, as might be expected, both word-level and syntactic changes (Ahn, 2022; McDonough et al., 2014), these are worded in non-specific terms.

By contrast, Bhagat and Hovy (2013) provide a lengthy list of categories based on their analysis of paraphrases from two corpora—the multiple-translations corpus and the Microsoft research paraphrase corpus. The categories, numbering 25 in all, range from the usual synonym substitution to changes based on world knowledge [e.g., *We must work hard to win this election* ~ *The Democrats must work hard to win this election* (Bhagat & Hovy, 2013: 469)]. This list is extensive and helpful, but as the authors

themselves point out, the categories address only lexical substitutions; possible changes at the level of the clause and beyond are not covered.

Most recently, the advent of AI and the growth of APTs have considerably eased the actual task of paraphrasing. APTs originated as text-spinning, an abusive means to achieve search engine optimization to boost the ranking of targeted web pages (Zhang et al., 2014). Spinning replaces expressions on the web page, producing a version that is different enough to avoid plagiarism detection. Automated tools propagate spinning by creating multiple versions of the same original web page, which are then used by spammers to point back to the targeted web page, thus improving its page rank. Today, APTs have developed by leaps and bounds, and their applications go beyond mere page ranks. Trained using large datasets and natural language processing (NLP) algorithms, APTs are generally able to produce human-like content with a click of the mouse.

If used appropriately as a tool to aid learning, APTs can bring much benefit to students, particularly those learning English as a second or foreign language. For instance, Chen et al. (2015) created PREFER (PREFabricated Expression Recognizer) to help Chinese students in their paraphrasing tasks. PREFER was a suggestion system that listed several options (e.g., *depend on*, *count on*) in response to a query (*rely on*). In other words, it required users to actively evaluate and choose competing options in context, thus facilitating learning.

Present-day APTs, however, generate entire paraphrases for the user, removing any need for evaluation or choice. It also dispenses with the need for paraphrasing instructions and exercises, since all the hard work is handled by the technology. The convenience offered by APTs and, of greater concern, their pervasiveness raise ethical concerns. It is now easier than ever for users to exploit APTs to create paraphrases in a matter of seconds and to pass them off as their own work. But this constitutes academic dishonesty, regardless of whether citations are inserted just to satisfy anti-plagiarism regulations. As Roe and Perkins (2022: 6) put it, “this does not change the core fact that the student’s submitted work was not their own.”

As the NLP architecture in APTs analyzes and remembers relevant information carried by words in the source text, it is tempting to assume that changes in the generated content will also be largely word-based. This, however, need not be so since such changes may also affect the use of different clause types and inter-clausal relations. These clausal phenomena, though, remain a largely unexplored area in paraphrasing. To more fully understand them, we turn next to Halliday’s clause-complexing framework.

3. Hallidayan framework on clause complexing

3.1. *Ranking clauses, embedded clauses, and clause complexes*

The Hallidayan framework separates grammatical units according to a RANK scale. In this hierarchy, the constituent occupying the highest rank is the clause, followed by the word phrase, and the word. This rank scale is illustrated in Figure 1.

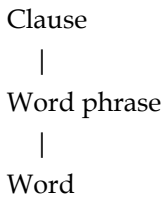


Figure 1: Halliday's rank scale (adapted)

Clauses, however, can also function as word phrases, i.e., at a rank below that of a clause. The Hallidayan framework thus makes a distinction between ranking clauses and EMBEDDED clauses. Ranking clauses are separate clauses; they can be either main or subordinate clauses, but they each function as a complete unit. On the other hand, embedded clauses, such as noun clauses or restrictive (defining) relative clauses, are always part of a larger clause. As they operate at a lower rank, they are sometimes also called downranked or rankshifted clauses. Examples of ranking and embedded clauses are provided in (1–2) below. Following the Hallidayan convention, clauses are separated using double vertical lines $\parallel$, and embedded clauses are enclosed using double square brackets $[[\dots]]$. A list of all symbols, including those for tactic and logico-semantic relations (see Sections 5.3 and 5.4), is given in the appendix. All examples in this paper are taken from the corpus.

- (1) *We have designed the iconic 6502 microprocessor in both technologies as a use case $\parallel$ to demonstrate $\parallel$ and expand the multi-project wafer approach.*
- (2) *So far it has proven challenging $[[$ to achieve detectable supercurrents through quantum Hall conductors. $]]$ (Nature 16, embedded clause functioning as delayed subject)*

Inter-clausal relations, if any, occur within the CLAUSE COMPLEX, a term used to refer to the orthographical sentence. The most basic clause complex, also known as a SIMPLEX, contains only a single main clause, but it is not uncommon for clause complexes to contain a mixture of ranking and embedded clauses, as in (3). Clause complexes in the Hallidayan framework are enclosed using triple vertical lines $\parallel\parallel\parallel\dots\parallel\parallel\parallel$.

- (3) $\parallel\parallel\parallel$ *Unlike pterosaurs, birds and bats, the wings of insects did not evolve from legs, $\parallel$ but are novel structures $[[$ that are attached to the body via a biomechanically complex hinge $[[$ that transforms tiny, high-frequency oscillations of specialized power muscles into the sweeping back-and-forth motion of the wings. $]]$ $]]$ $\parallel\parallel\parallel$ (Nature 45)*

Inter-clausal relations within clause complexes are categorized in both tactic and logico-semantic terms. The details of these relations are outlined in the next two subsections.

3.2. Tactic relations

We begin with tactic relations; these refer to the type of interdependency between clauses. The relation between clauses of equal status is known as PARATAXIS, and that between clauses of unequal status, HYPOTAXIS. Since parataxis captures merely the sequential ordering of equal-status clauses, Arabic numerals (i.e., ‘1’, ‘2’, ‘3’) are used. These clauses need not be main clauses only, since two or more subordinate clauses can also be paratactically related.

By contrast, in hypotaxis, Greek letters are used to signal the unequal relationship between clauses; the dominant clause is represented by ‘α’, and the other clauses dependent on it as ‘β’, ‘γ’, and so on. The example in (4) below illustrates both parataxis and hypotaxis in the same clause complex; here, both the subordinate clauses are paratactically related to each other.

- (4)
- α

β

|||

||

||

The γδ TCR associates with CD3 subunits,

initiating T cell activation

and holding great potential in immunotherapy.

|||

(Nature 01)

These tactic relations add a further layer of insight into how clauses are used at the text level. Unlike traditional grammar, which recognizes only the main-subordinate distinction between clauses, (4) shows that while the segment describing the result and significance of the association between the TCR and CD3 subunits has a subordinate status, it contains clauses that are of equal weighting to each other. That is to say, the outcome of the TCR-CD3 association is as important as its potential in immunotherapy – indeed, no potential is possible without a desired outcome. As Leong (2023a: 107) notes, “[l]abeling clauses as solely ‘main’ (‘independent’) or ‘subordinate’ (‘dependent’), as in traditional grammar, runs the risk of obscuring such relations.”

3.3. Logico-semantic relations

Logico-semantic relations capture the second way in which clauses are related. These relations are both logical and semantic in the sense that they take into account both the interdependency (i.e., the tactic relation) and meaning association between clauses (Halliday, 2006). In the mainstream Hallidayan framework, these semantic relations fall into two broad types, PROJECTION and EXPANSION.

Projected clauses are those that carry the locution or idea expressed by verbs of saying or thinking, as underlined in (5) below.

- (5)
- Overall, we propose

that ODA is a key reaction mechanism for complexity acceleration in the processing of DOM molecules, [...]
- (Nature 22)

The Hallidayan framework regards such projected clauses as ranking clauses. Such an analysis, however, is problematic for two reasons. The first is that if the projected clause is a full ranking clause, it then cannot be preceded by a noun phrase (e.g., *the fact*),

since this would make the projected clause a postmodifier of the head noun, rendering it an embedded clause. This simple test appears to work well for some verbs (e.g., **We say the fact that*), but not all. Verbs such as *state* or *assert* can be easily followed by *the fact*, but the projected clauses are nevertheless regarded as ranking clauses in the framework.

The second, which is harder for the framework to defend, is that the status of the projecting (i.e., initiating) clause as being a proper clause becomes unworkable. Since ranking clauses are, by definition, complete clauses, this would make *we propose* in (5) a non-clause. As Fawcett (2000: 29) aptly points out, the projecting clause is actually “an uncompleted clause that is ‘expecting’ [...] another element (which we may call a Complement).” For these reasons, projection is excluded in the present analysis. What the Hallidayan framework considers a projected clause is regarded in this paper as an embedded complement clause instead.

This paper, then, considers only expansion in its logico-semantic analysis. The Hallidayan framework divides expansion into three sub-types – EXTENSION, ENHANCEMENT, and ELABORATION.

In extension, the clause provides new information, an exception, or an alternative. Extending clauses, represented by the addition sign (+), typically contain conjunctions such as *and* or *but* in parataxis (6), or *whereas* in hypotaxis (7).

- (6) 1 ||| *Here we describe the structure of RAD52*
+2 || *and define the mechanism of annealing.* ||| (Nature 21)
- (7) α ||| *Annealing is driven by the RAD52 N-terminal domains,*
+β || *whereas the C-terminal regions modulate the open-ring conformation and RPA interaction.* ||| (Nature 21)

Enhancement offers circumstantial information, and so ‘enhances’ the overall description in the clause complex. Such information, being adverbial in nature, relates to aspects such as time, manner, reason, and the like. Enhancing clauses are represented by the multiplication sign (×).

- (8) α ||| *Here we describe the growth of diamond crystals and polycrystalline diamond films with no seed particles*
×β || *using liquid metal but at 1 atm pressure and at 1,025 °C,*
×γ || *breaking this pattern.* ||| (Nature 06)
- (9) α ||| *The spike timing of these PAC neurons was coordinated with frontal theta activity*
×β || *when cognitive control demand was high.* ||| (Nature 39)

The final subtype, elaboration, restates the information by offering “a further characterization of one that is already there” (Halliday & Matthiessen, 2014: 461). The information in the elaborating clause can be equated, as it were, with that in the primary clause. Elaborating clauses are sometimes introduced using punctuation marks such as a colon or dash (10), or specific clause types such as the non-restrictive (non-defining) relative clause (11). Elaborating clauses are marked off using the equality sign (=).

- (10) 1 ||| *Recent research sheds light on the process:*
 =2 || *activation of the pore-forming protein GSDMD by the cytosolic lipopolysaccharide (LPS) sensor caspase-11 – but not by TLR4-induced cytokines – mediates BBB breakdown in response to circulating LPS or during LPS-induced sepsis.* ||| (Copilot 42)
- (11) α ||| *This is despite pervasive and ongoing gene flow with one parent, *Heliconius pardalinus*,*
 =β || *which homogenizes 99% of their genomes.* ||| (Nature 31)

At this stage, a distinction is needed to separate restrictive from non-restrictive relative clauses. The former serves an identifying function within the noun phrase, and so plays a key role in distinguishing the head noun from other nouns. Restrictive relative clauses, that is to say, are a crucial part of the noun phrase and are therefore embedded.

Non-restrictive relative clauses, on the other hand, do not identify, but describe the noun that is already taken to be fully specific. The parenthetical nature of such clauses is reflected in how punctuation marks such as commas, dashes, or brackets are used to separate them from the rest of the clause. As Lakoff (1968: 45) notes, non-restrictive relative clauses “serve no limiting function so there is no reason to believe that they are associated with noun phrases in deep structure.”

As tactic and logico-semantic relations are inter-clausal phenomena, they apply to all clauses at the same rank, regardless of whether the clauses themselves are ranking or embedded. The example in (12) illustrates an analysis involving a mix of ranking and embedded clauses; embedded inter-clausal relations are marked using the superscript ‘E’.

- (12) 1 ||| *We simultaneously tracked single-cell mtDNA heteroplasmy and ancestry,*
 +2 || *and found*
 ×β^E [[*that, although the population heteroplasmy shifts,*
 α^E || *the heteroplasmy of individual cell lineages remains stable, [...]]] ||| (Nature 09)*

3.4. Related studies on clause complexing

Early work on clause complexing focused on a variety of genres involving both speech and writing. These included spoken narrative texts in Australian English, where extension was found to be prevalently used (Nesbitt & Plum, 1988). In a later comparative study involving a range of spoken and written categories such as conversations, handwritten letters, and academic writing, among others, Greenbaum and Nelson (1995) found that only spontaneous conversations differed from the other categories, having the highest proportion of simplexes. All the other categories, whether

spoken or written, did not display any marked differences. Other genres investigated included aphasic discourse (Armstrong, 1992) and student essays (Leong & Wee, 2005).

Where scholarly writing is concerned, research has been far more modest. Nevertheless, three studies deserve mention. The work of Sellami-Baklouti (2011), involving 120 abstracts in the sciences and social sciences, found that science writing favored the use of simplexes, whereas social science writing preferred hypotactic relations. Her work represents a welcome foray into the grammar of scholarly writing, showing explicitly how the two broad disciplines differ in terms of clause usage. Sellami-Baklouti's characterization of science writing was confirmed a decade later in the work of Leong (2021), who analyzed full scholarly articles instead of abstracts. Leong's study also added a further insight into humanities writing, which was found to contain more embedded clauses.

Most recently, in view of the overwhelming popularity of chatbots, Leong (2023a) compared original scientific abstracts with those generated using Google's Bard, OpenAI's ChatGPT, and Quora's Poe Assistant. He found that none of the chatbots matched the original abstracts in all clause-complexing categories; only ChatGPT came closest, although there were still distinct differences in the use of expanding and elaborating clauses.

Generating a new abstract, however, is not quite the same as paraphrasing the original. As pointed out in Section 1, paraphrases should retain the meaning of the original text as much as possible. Paraphrasing also involves more than mere lexical substitutions; grammatical changes are expected as well. The resulting text, to reiterate, should be different enough from the original version.

Given the paucity of studies on clause complexing, however, it remains an open question whether AI-generated paraphrases differ sufficiently from the source text. This study sought to address this gap by comparing original abstracts published in a prestigious journal with paraphrases generated by three APTs. It examined whether the use of tactic and logico-semantic relations differed between them.

4. Methodology

4.1. Corpus

The corpus comprised 50 original abstracts, with each abstract paraphrased by three APTs—QuillBot, Jasper, and Microsoft's Copilot. The total number of texts in the corpus was therefore 200, i.e., 50 original abstracts and 150 generated paraphrases.

The original abstracts were published in *Nature*, a highly cited journal in the sciences. According to Scimago (2023), it was ranked first in the field of multidisciplinary science. The abstracts were taken from the 50 most recent publications at the time of analysis. The APTs originally selected were QuillBot and Jasper; they received favorable user reviews and appeared as the top two paraphrasing tools in a few lists (Aayush, 2024; Gulati, 2024). A third tool, Copilot, was subsequently added due to its ease of access and seeming ubiquity. Starting from January 2024, Copilot can be easily activated via a

dedicated key on Windows-compatible keyboards (Warren, 2024), making it even more convenient for users to generate content seamlessly.

Using the APTs to generate the paraphrases was fairly straightforward, given the effort put in by the developers to make the process as easy and intuitive as possible for the user. In QuillBot, paraphrasing required only the click of a button. The setup in Jasper and Copilot, however, was different, and so slightly dissimilar prompts were used:

The above is an abstract of a research paper. Paraphrase it as a single paragraph using formal English. (Jasper; prompt placed beneath original abstract)

The following is an abstract of a research paper. Paraphrase it as a single paragraph using formal English. The abstract is as follows: "<original abstract>" (Copilot)

4.2. Method of analysis

Each text was broken up into ranking and embedded clauses, and the frequency counts of tactic and logico-semantic relations (as described in Sections 3.2-3.3) were recorded. Microsoft Excel was used to keep track of the frequency counts; each ranking or embedded clause was placed on a separate row.

This is illustrated in Figure 2, taking the first nine clauses of Nature 01 as an example. In the analysis, double-angle brackets <<...>> are used to mark off the inserted clause (Clause 2).

Nature 01	Main Clause	Subordinate clause	Embedded clause	Parataxis				Hypotaxis			
				+	x	=		+	x	=	
1 Gamma delta (γδ) T cells, a unique T cell subgroup, are crucial in various immune responses and immunopathology.	1										
2 The γδ T cell receptor (TCR), << generated by γδ T cells, >>		1						1			1
3 recognizes a diverse range of antigens independently of the major histocompatibility complex.	1										
4 The γδ TCR associates with CD3 subunits,	1										
5 initiating T cell activation		1						1		1	
6 and holding great potential in immunotherapy.		1		1	1						
7 Here, we report the structures of two prototypical human Vγ9Vδ2 and Vγ5Vδ1 TCR-CD3 complexes,	1										
8 unveiling two distinct assembly mechanisms		1						1		1	
9 [[that depend on Vγ usage.]]			1								

Figure 2: Partial sample analysis table of Nature 01

4.3. Statistical analysis

As each text contained a different number of clauses, normed rates of occurrence per 100 words were used for comparison (Biber & Jones, 2009). Real Statistics Resource Pack (Zaiontz, 2022), a Microsoft Excel add-in, was used to conduct statistical tests. These included the Welch one-way analysis of variance (ANOVA) test and the Games-Howell post-hoc test for all significant ANOVA results. The significance level for all statistical tests was $\alpha=.05$. In this paper, a single asterisk is used for $p<.05$, and double asterisks for $p<.01$.

5. Results and discussion

5.1. AI detection, plagiarism, and authorial markers

In a preliminary scan of the generated abstracts, Copyleaks, an originality detection tool, was used to ascertain if the abstracts were able to bypass basic checks on academic integrity. In the main, the generated abstracts did not escape AI or plagiarism detection (see Table 1). The best-performing APT in evading AI detection was QuillBot, where 20 abstracts returned a nil score for the use of AI. By contrast, all of Jasper’s abstracts were assessed to be AI-generated.

Table 1: Mean and median Copyleaks scores for APT abstracts

APT	Use of AI			Plagiarism		
	Mean	Median	No AI [‡]	Mean	Median	No plag. [‡]
QuillBot (n=50)	60	100	20	78.92	87.5	3
Jasper (n=50)	100	100	–	70.76	79	5
Copilot (n=50)	80	100	10	75	88.50	4

[‡] Frequency count of articles with a 0% score

Most of the generated articles were also picked out as containing plagiarized content. Sixteen abstracts from QuillBot, and 13 abstracts each from Jasper and Copilot had a score of 100%. In fairness to the APTs, much of the plagiarized content was attributed to uncited paraphrasing, which can be easily rectified by including appropriate in-text citations. Nevertheless, instances of direct copying from other sources were not difficult to locate. The following is a screenshot of the plagiarism report for Copilot 04, where the darker shaded segments are duplicates from the original abstract itself.

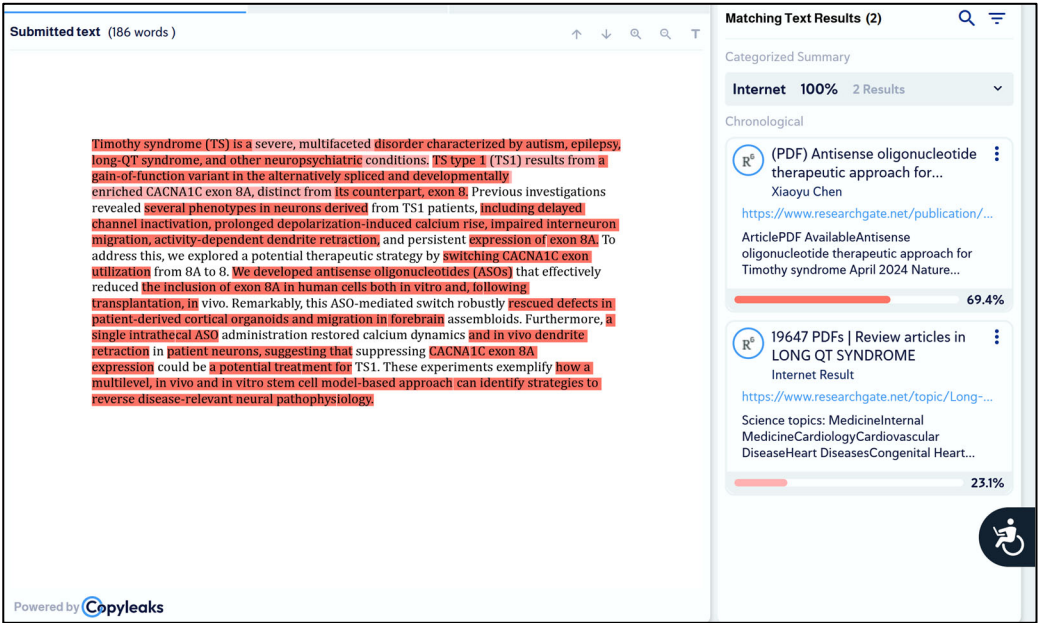


Figure 3: Plagiarism report of Copilot 04

Two other remarks about the corpus are necessary. First, unlike the situation encountered by Rogerson and McCarthy (2017), who found that paraphrasing tools often produced confusing expressions and grammatical errors (often referred to as *word salads*), the paraphrases in this study were generally coherent and free of language errors. The only expression that came across as being somewhat odd was *cancer beginning studies* (QuillBot 18) for *cancer initiation studies*. Apart from this anomaly, the overall good language quality in the generated texts illustrates the advancements made over the years to ensure that the output is at least grammatically and semantically sound.

The second comment relates to authorial markers. If an author paraphrases her/his own original text to avoid self-plagiarism, retaining the first-person pronouns is perfectly legitimate. In most cases, however, paraphrasing involves texts written by *someone* else. If so, the authorial markers in the original text must be changed; for instance, the use of *we* should be replaced by *the researchers* or the like. Yet, only four paraphrases – Copilot 02, Copilot 40, Copilot 42, and Jasper 48 – made this change. On this count, then, the APTs did not appear to be sensitive enough to make adjustments to the writing with regard to this very fundamental issue. More worryingly, continuing the use of *we* in the paraphrase represents a different type of dishonesty for the writer, i.e., claiming to have done or found something when this is patently false.

5.2. Use of ranking and embedded clauses

Table 2 presents the summary statistics concerning the distribution of clauses in the corpus. The proportions of main (72.12%) and subordinate clauses (27.88%) to ranking clauses, and embedded clauses (33.67%) to total clauses are comparable to those reported in Leong (2023a), who also included Nature abstracts in his study. This consistency in the proportions across two different corpora suggests that they could serve as an identifying feature of science writing, at least where abstracts are concerned.

Table 2: Word length and frequency counts of clauses

	Nature	QuillBot	Jasper	Copilot	Total
Words	10,010	11,636	9,061	8,993	30,716
(A) Main clauses	476	560	332	461	1,829
(B) Subordinate clauses	184	206	250	215	855
Total A+B	660	766	582	676	2,684
(C) Embedded clauses	335	468	333	246	1,382
Total A+B+C	995	1,234	915	922	4,066

Table 3: Mean occurrence rates of main, subordinate, and embedded clauses

	Mean	SD	Games-Howell comparisons (<i>p</i>)		
			Nature	QuillBot	Jasper
Main clauses , $F(3, 106.72)=53.70, p<.001^{**}$					
Nature	4.76	1.04			
QuillBot	4.82	0.79			
Jasper	3.66	0.56	<.001**	<.001**	
Copilot	5.12	0.68			<.001**
Subordinate clauses , $F(3, 108.65)=8.25, p<.001^{**}$					
Nature	1.83	1.10			
QuillBot	1.80	1.07			
Jasper	2.77	1.20	<.001**	<.001**	
Copilot	2.43	1.29		.044*	
Embedded clauses , $F(3, 107.44)=7.58, p<.001^{**}$					
Nature	3.38	1.71			
QuillBot	4.04	1.72			
Jasper	3.70	1.35			
Copilot	2.80	1.15		<.001**	<.001**

The mean occurrence rates (per 100 words) of main, subordinate, and embedded clauses are presented in Table 3. Only the *p* values of statistically significant results are reported. Of the three APTs, QuillBot and Copilot did not differ statistically from Nature in the use of both ranking and embedded clauses. As compared to Nature, Jasper used fewer main clauses and more subordinate clauses. Differences in the use of embedded clauses involved only Copilot.

Given that paraphrasing should ideally involve grammatical changes, the similarity in the use of ranking clauses in Nature, QuillBot, and Copilot presents a complication that many accounts of paraphrasing do not often address adequately. In (13a–c), although lexical substitutions are evident, the paraphrases (13b–c) clearly mirror the structure of the original text (13a). In each case, two simplexes are used, and Copyleaks was able to accurately trace both paraphrases to the Nature abstract. This highlights the crucial need for paraphrases to be different enough from the original, as alluded to in Section 3.4.

(13a) ||| *Phenotypic variation among species is a product of evolutionary changes to developmental programs.*

||| *However, how these changes generate novel morphological traits remains largely unclear.* ||| (Nature 23)

(13b) ||| *The phenotypic heterogeneity observed among species is a result of evolutionary modifications to developmental programs.*

||| *Nevertheless, the mechanisms by which these alterations give rise to new morphological characteristics remain mainly obscure.* ||| (QuillBot 23)

(13c) ||| *Phenotypic variation across species arises from evolutionary changes in developmental processes.*

||| *However, the mechanisms underlying the emergence of novel morphological traits remain poorly understood.* ||| (Copilot 23)

The only APT that differed from Nature in the distribution of clauses was Jasper. The low occurrence rate of main clauses in Jasper is particularly noteworthy. As seen in Table 2, the proportion of main clauses to ranking clauses in Jasper abstracts was about 57%, which is 15 percentage points lower than that for Nature. In consequence, Jasper's preference for subordinate clauses adds a layer of complexity (Biber & Gray, 2016) to its paraphrases. This complexity is generally manifested in two possible ways – through the use of ranking subordinate clauses or embedded clauses. In relative terms, Table 3 suggests that Jasper relied more on subordinate clauses, since the occurrence rate of embedded clauses did not differ statistically from that of either Nature or QuillBot.

There is one other consequence of Jasper's reliance on subordinate clauses. We see this in its paraphrase (14b) of the original passage (14a) below.

(14a)	1	<i>The compound was well tolerated in mice</i>
	×2	<i>and led to robust tumour regression in multiple MSI-H colorectal cancer cell lines and patient-derived xenograft models.</i>
	1	<i>Our work shows an allosteric approach for inhibition of WRN function</i> [[<i>that circumvents competition from an endogenous ATP cofactor in cancer cells, </i>]]
	+2	<i>and designates VVD-133214 as a promising drug candidate for patients with MSI-H cancers.</i> (Nature 10)
(14b)	×β	<i>Demonstrating considerable tolerance in mouse models and substantial tumor reduction in various MSI-H colorectal cancer lines and xenografts</i> [[<i>derived from patients, </i>]]
	α	<i>the study proposes VVD-133214 as an innovative allosteric strategy</i>
	×γ	<i>to inhibit WRN function.</i> (Jasper 10)

In (14a), the result of the study is expressed as two clause complexes, and the clauses in both clause complexes are paratactically related. This makes the writing more straightforward and thus easier to process. By contrast, the paraphrase in (14b) is far more compressed. This denseness is seen in how the entire first sentence in (14a) is reframed as the opening clause in (14b). Further, as the opening clause in (14b) is subordinate to the main clause, the reader is forced to hold the information in memory until the appearance of the latter. This adds to the overall complexity of the paraphrase.

We have already seen earlier the similarity between Copilot and Nature where ranking clauses are concerned. However, differences in the occurrence rates of subordinate clauses for Copilot and Jasper were also found to be statistically insignificant. To examine if the seeming preference for clausal complexity by Jasper and Copilot had a bearing on their respective uses of simplexes, the occurrence rates of simplexes were compared. The results are presented in Table 4.

Table 4: Mean occurrence rates of simplexes

	Mean	SD	Games-Howell comparisons (<i>p</i>)		
			Nature	QuillBot	Jasper
Simplexes, F(3, 108.27)=24.73, <i>p</i> <.001**					
Nature	2.19	0.87			
QuillBot	2.51	0.74			
Jasper	1.33	0.76	<.001**	<.001**	
Copilot	2.49	1.00			<.001**

The results show that Jasper had the lowest occurrence rate of simplexes among the abstract groups. The use of simplexes in Copilot, on the other hand, was comparable to that in Nature and QuillBot. We therefore see a ‘balanced’ approach in the way Copilot handled clausal complexity vis-à-vis Jasper –while it resembled Jasper in its use of subordinate clauses, it eased away from intensifying this complexity by using more simplexes. Jasper, on the other hand, favored a condensed, complex style of writing, using more subordinate clauses and fewer main clauses and simplexes.

5.3. Use of tactic relations

A comparison of the use of tactic relations in the corpus reinforces the findings in the preceding section, particularly the close match between Nature and QuillBot, and the preference for clausal complexity in Jasper and Copilot. The results for both parataxis and hypotaxis are listed in Table 5. It should be noted that the rates of hypotactic relations in Table 5 are lower than the rates of subordinate clauses in Table 3; this is because subordinate clauses can also be paratactically related to each other, as shown earlier in example (4).

The statistical insignificance between Nature and QuillBot in the use of clauses, whether ranking or embedding, parallels the findings reported in Table 3. The results here also reiterate Jasper’s and Copilot’s propensity to generate complex texts, rather than simpler paratactic versions. The occurrence rates of paratactic ranking clauses in Jasper (0.32) and Copilot (0.44) were only about half that in Nature (0.80). As Halliday and Matthiessen (2014: 452) note, the relationship between paratactic clauses is one of sequence, not dependence, thus allowing for information to be understood quickly. This is not unlike the grammar of spoken discourse, where “short clause-like chunks [are] chained together in a simple incremental way for ease of processing” (Leech, 2000: 699).

Table 5: Mean occurrence rates of tactic relations

	Mean	SD	Games-Howell comparisons (<i>p</i>)		
			Nature	QuillBot	Jasper
Paratactic ranking clauses , F(3, 107.19)=7.15, <i>p</i> <.001**					
Nature	0.80	0.68			
QuillBot	0.60	0.59			
Jasper	0.32	0.41	<.001**	.029*	
Copilot	0.44	0.50	.017*		
Paratactic ranking clauses					
Nature	0.11	0.23			
QuillBot	0.24	0.36			
Jasper	0.01	0.24			
Copilot	0.13	0.37			
Hypotactic ranking clauses , F(3, 108.30)=10.64, <i>p</i> <.001**					
Nature	1.72	0.93			
QuillBot	1.68	0.95			
Jasper	2.69	1.14	<.001**	<.001**	
Copilot	2.33	1.21	.026*	.016*	
Hypotactic embedded clauses , F(3, 108.12)=3.95, <i>p</i> =.010*					
Nature	0.39	0.40			
QuillBot	0.52	0.52			
Jasper	0.50	0.45			
Copilot	0.26	0.36		.029*	.026*

The clausal complexity in Jasper’s paraphrases is a layered one, involving not just hypotactic ranking clauses, but embedded clauses as well. (This is unlike Copilot, which used significantly fewer hypotactic embedded clauses than Jasper.) Even though Nature and Jasper did not differ in their use of hypotactic embedded clauses, these clauses contributed to the overall complexity in several of Jasper’s paraphrases. A case in point is (15).

- (15)
- α
- ||| The Triton dataset’s detailed account allows for an unprecedented examination of Cenozoic pelagic macroevolution

=β || *where the global biogeographic responses of functional communities and richness exhibit a nuanced disconnection during climatic shifts of the era,*

+γ || *suggesting*

α^E [[*that the universal reaction of functional groups to comparable abiotic stressors may be influenced by the prevailing climatic conditions, whether greenhouse or icehouse,*

=β^E || *to which the groups were adapted.*]] || (Jasper 38)

This paraphrase, comprising three ranking and two embedded clauses, could have been broken up into two sentences. Instead, two elaborating clauses and an extending clause are used, complicating the structure somewhat. The first elaborating clause also interrupts the transition between what the details in the Triton dataset allow for and what they imply.

5.4. Use of logico-semantic relations

Logico-semantic relations offer another way of characterizing the interaction of clauses in texts. We begin with parataxis. Overwhelmingly, extension was used as the paratactic logico-semantic relation in all abstracts. The number of enhancing and elaborating clauses—one and seven, respectively—was far too small for fair comparisons to be made; they were therefore excluded from the analysis. The rates of paratactic extension are presented in Table 6.

Table 6: Mean occurrence rates of paratactic extension

	Mean	SD	Games-Howell comparisons (<i>p</i>)		
			Nature	QuillBot	Jasper
Ranking clauses, $F(3, 107.50)=7.41, p<.001^{**}$					
Nature	0.79	0.64			
QuillBot	0.58	0.57			
Jasper	0.31	0.41	<.001**	.034*	
Copilot	0.42	0.49	.009**		
Embedded clauses					
Nature	0.24	0.35			
QuillBot	0.24	0.36			
Jasper	0.10	0.24			
Copilot	0.13	0.37			

The observed differences among embedded clauses were not found to be statistically significant. In ranking clauses, as might be expected from the discussion in Section 5.3, Jasper and Copilot had lower occurrence rates of paratactic extending clauses. This again highlights their shift from the relatively straightforward coordination of clauses toward a dependency (or hypotactic) relationship.

The hypotactic logico-semantic relations found are listed in Table 7. These involve only hypotactic enhancements and elaborations; there were too few instances of hypotactic extensions for comparisons to be made.

Table 7: Mean occurrence rates of hypotactic enhancement and elaboration

	Mean	SD	Games-Howell comparisons (<i>p</i>)		
			Nature	QuillBot	Jasper
Enhancement, ranking clauses					
Nature	1.19	0.91			
QuillBot	1.19	0.89			
Jasper	1.27	0.95			
Copilot	1.32	1.06			
Enhancement, embedded clauses					
Nature	0.26	0.35			
QuillBot	0.33	0.42			
Jasper	0.29	0.37			
Copilot	0.17	0.29			
Elaboration, ranking clauses, $F(3, 105.69)=14.20, p<.001^{**}$					
Nature	0.45	0.54			
QuillBot	0.49	0.40			
Jasper	1.29	0.90	<.001**	<.001**	
Copilot	0.83	0.60	.010**	.008**	.016*
Elaboration, embedded clauses					
Nature	0.12	0.22			
QuillBot	0.18	0.30			
Jasper	0.21	0.33			
Copilot	0.09	0.21			

The statistical insignificance involving hypotactic enhancements in both ranking and embedded clauses is perhaps unsurprising. Hypotactic enhancements are the equivalent of what traditional grammar terms ‘adverbial clauses’ (Halliday & Matthiessen, 2014:

481); they are essential to provide the needed circumstantial information surrounding scientific facts, methods, and findings. While the same circumstantial information can be expressed paratactically, as seen in (16a–b), paratactic enhancements are exceedingly rare in the corpus. In the case of (16b), paratactic enhancement would also have required the two clauses to be re-sequenced.

- (16a) α ||| *This is achieved*
 ×β || *by using a novel comprehension of carrier recombination and transport in single-crystal Cu₂O thin films.* ||| (QuillBot 14)
- (16b) 1 ||| *A novel comprehension of carrier recombination and transport in single-crystal Cu₂O was used*
 ×2 || *and so achieved a superior performance of Cu₂O photocathodes.* |||

The only significant differences in Table 7 involved hypotactic elaborations in ranking clauses. In the Hallidayan framework, such elaborations are manifested as non-restrictive relative clauses, whether finite or non-finite (Halliday & Matthiessen, 2014: 464). The occurrence rate was highest for Jasper, and together, Jasper and Copilot used more of such clauses than Nature and QuillBot. In several instances, more than one elaborating clause was used in the same clause complex (17–18).

- (17) α ||| *In this detailed investigation, the focus is on Heliconius elevatus,*
 =β || *which emerged as a distinct species through hybridization,*
 =γ || *coexisting with its parent species for over 180,000 years.*
 ||| (Jasper31)
- (18) =β ||| *In the field of palaeontology, functional groups – << defined by consistent ecological and morphological traits across a clade’s evolutionary history >> –*
 α || *offer distinct insights into biodiversity dynamics*
 ×γ || *compared to species and genera,*
 =δ || *which are relatively short-lived in evolutionary terms.* ||| (Copilot 38)

This use of hypotactic elaborations is an interesting finding, revealing one of the means by which Jasper and Copilot compressed and repackaged information in their paraphrases. This is exemplified in (19a–b).

- (19a) ||| *Our biochemical and biophysical assays further corroborate the dimeric structure.*
- ||| *Importantly, the dimeric form of the V γ 5V δ 1 TCR is essential for T cell activation.* ||| (Nature 01)
- (19b) = β ||| *This dimerization, << validated by biochemical and biophysical tests, >>*
- α *is crucial for T cell activation.* ||| (Jasper 01)

Here, the original extract (19a), comprising two simplexes, is condensed into a single clause complex in (19b). The elaborating clause in (19b) re-frames the information about biochemical and biophysical assays as a gloss on dimerization; in other words, it is cast as an incidental piece of information, a mere clarification about dimerization’s validation. This information in the original text, though, is new, not incidental, but compressing the original version using elaboration does succeed in laying the correct focus on the important role of the V γ 5V δ 1 TCR in activating T cells. In this light, the paraphrase captures not only the meaning of the original version, but more crucially its underlying significance.

6. Conclusion

This study compared original abstracts with paraphrases generated by QuillBot, Jasper, and Copilot. Tactic and logico-semantic relations were analyzed based on a modified version of the Hallidayan clause-complexing framework. The major findings are summarized as follows.

- (a) Prior to the analysis, Copyleaks was used to determine whether the APT paraphrases could evade AI and plagiarism detection. The majority could not. The best-performing AI-evasion APT was QuillBot; by contrast, all of Jasper’s paraphrases were picked out as being AI-generated. A total of 42 abstracts, constituting 28% of the generated paraphrases, had a plagiarism score of 100%.
- (b) As it is more typical for a person to paraphrase someone else’s writing, authorial markers in the original text should be amended as necessary. Only four APT paraphrases made such changes; the first-person pronoun continued to be used or implied in the rest of the paraphrases.
- (c) The APT that most closely matched Nature was QuillBot. No significant differences between them were found in all categories.
- (d) The distribution of main and subordinate clauses was similar among Nature, QuillBot, and Copilot. As compared to Nature, Jasper used fewer main clauses and more subordinate clauses.
- (e) Jasper and Copilot used fewer paratactic extensions and more hypotactic elaborations in ranking clauses than Nature, suggesting their greater propensity

for clausal complexity. Copilot, however, departed from Jasper in using more simplexes, reducing the extent of this complexity.

The close match between Nature and QuillBot could lead to plagiarism concerns, as illustrated in example (13). An interesting question, though, presents itself. Since paraphrasing should ideally result in grammatical changes, and Jasper and Copilot appear to have done so, do their paraphrases count as the standard? There is no easy answer to this question since deciding whether a paraphrase is good is often highly subjective. Uemlianin (2000: 348) points out that “assessments of student’s paraphrases and understanding can seem quite arbitrary and mysterious, almost aesthetic.”

The findings are also mixed. In some cases, the shift toward complexity allowed Jasper to better capture the core meaning in the original abstract [example (19b)], but in other cases, it led to denseness in the description, making the writing less straightforward [example (14b)]. Some, however, may point out that such a compressed style of writing is a characteristic feature of science writing (Biber et al., 2022; Leong, 2023b).

While this may be true, there have been recent calls for science writing to be made more readable. For instance, in a study involving abstracts published between 1881 and 2015, Plavén-Sigray et al. (2017) note that science writing has become more difficult to read, due primarily to the increased use of jargon and longer sentences. Making language overcomplicated can isolate science writing, making it even less accessible. As Chawla (2020) notes, “[n]ot only does such overcomplicated language alienate non-scientists and the media, it can also make life difficult for junior researchers and those transitioning to new fields.” As paraphrasing is intended to re-express the content of the original text, there is therefore little need for it to add to the complexity, as that runs the risk of obscuring the original meaning. Hence, while paraphrasing should include grammatical changes, current indications suggest that such changes should be done to clarify, not complicate the description.

Seen in this light, paraphrasing is a complex task. Writers are expected to retain the meaning of the original text by using different words and grammatical structures without complicating the description. There is every temptation to use APTs to ease this process, but as this study has shown, it is not difficult for a detection tool to pick out instances of AI use or plagiarism. Even though such detection tools are not foolproof – as, indeed, several paraphrases in this study fell through the cracks – they appear to work rather effectively, and with continuous testing and development, they can only get better over time. Writers thinking of using these tools in place of the actual effort that goes into the writing process are, therefore, taking a gamble. APTs, in other words, cannot (and should not) replace the work that they themselves need to put in. Rogerson and McCarthy (2017: 4) aptly remark:

The fact remains that taking another author’s work, processing it through an online paraphrasing tool then submitting that work as ‘original’ is not original work where it involves the use of source texts and materials without acknowledgement.

Given the challenges involved in paraphrasing, the obvious remedy is to provide explicit instructions to show how paraphrasing is done. Available studies have shown that paraphrasing instructions are helpful (e.g., Ahn, 2022; Choy & Lee, 2012). In scientific writing, specifically, the challenge lies not only in avoiding the (over)use of jargon, but utilizing tactic and logico-semantic relations appropriately to avoid unnecessary complexity. Writers cannot be assumed to know how to paraphrase by simply reading enough scientific papers and using a thesaurus.

Much further work in this area remains to be done. This study is limited by its focus on scientific writing and the use of only three APTs. There are a host of APTs, many of them free to use, that are widely accessible on the Internet. An in-depth look at the use of clauses and clause-complexing relations in the paraphrases generated by these other APTs will provide us with a better understanding of how they differ from source texts, and how paraphrasing instructions need to be modified and/or expanded to mitigate potential abuses of such tools. In future work, different genres of writing will need to be included as well, since plagiarism affects all disciplines. Papers in the humanities, social sciences, and multi-disciplinary fields (e.g., medical humanities) are bound to differ, given their varied focus areas and research methodologies. Extensive studies in these areas will help scholars, and particularly educators, keep pace with how AI and technology impact language education. The focus is to enhance the skills needed for writing, not to replace them.

References

- Aayush (2024). 8 best paraphrasing tools in 2024 (compared). *Elegant Themes*. <https://www.elegantthemes.com/blog/business/best-paraphrasing-tools> [accessed on 29 Nov 2024]
- Ahn, Soojin (2022). The effect of paraphrasing instruction on Korean EFL learners' attempts of paraphrasing. *Journal of Asia TEFL* 19(2): 637–646. <https://doi.org/10.18823/asiatefl.2022.19.2.16.637>
- Armstrong, Elizabeth M. (1992). Clause complex relations in aphasic discourse: A longitudinal case study. *Journal of Neurolinguistics* 7(4): 261–275. [https://doi.org/10.1016/0911-6044\(92\)90018-R](https://doi.org/10.1016/0911-6044(92)90018-R)
- Bhagat, Rahul, Eduard Hovy (2013). What is a paraphrase? *Computational Linguistics* 39(3): 463–472. https://doi.org/10.1162/COLI_a_00166
- Biber, Douglas, James K. Jones (2009). Quantitative methods in corpus linguistics. Lüdeling, Anke, Merja Kytö, eds., *Corpus Linguistics: An International Handbook* (Vol. 2). Berlin: Mouton de Gruyter, 1286–1304.
- Biber, Douglas, Bethany Gray (2016). *Grammatical Complexity in Academic English: Linguistic Change in Writing*. Cambridge: Cambridge University Press.
- Biber, Douglas, Bethany Gray, Shelley Staples, Jesse Egbert (2022). *The Register-Functional Approach to Grammatical Complexity: Theoretical Foundation, Descriptive Research Findings, Application*. New York: Routledge.
- Chawla, Dalmeet Singh (2020). Science is getting harder to read. *Nature*. <https://www.nature.com/nature-index/news/science-research-papers-getting-harder-to-read-acronyms-jargon>

- Chen, Mei-Hua, Shih-Ting Huang, Jason S. Chang, Hsien-Chin Liou (2015). Developing a corpus-based paraphrase tool to improve EFL learners' writing skill. *Computer Assisted Language Learning* 28(1): 2–40. <https://doi.org/10.1080/09588221.2013.783873>
- Choy, Siew Chee, Mun Yee Lee (2012). Effects of teaching paraphrasing skills to students learning summary writing in ESL. *Journal of Teaching and Learning* 8(2): 77–89. <https://doi.org/10.22329/jtl.v8i2.3145>
- Davidson, William, Joseph Crosby Alcock (1877). *Complete Manual of Analysis and Paraphrasing*. London: T. J. Allman.
- Fawcett, Robin P. (2000). *A Theory of Syntax for Systemic Functional Linguistics*. Amsterdam: John Benjamins.
- Greenbaum, Sidney, Gerald Nelson (1995). Clause relationships in spoken and written English. *Functions of Language* 2(1): 1–21. <https://doi.org/10.1075/fol.2.1.02gre>
- Gulati, Mayank (2024). Top 5 best online paraphrasing tools for 2024—rephrase your content. *ContentNinja*. <https://www.contentninja.in/dojo/content-marketing/top-5-online-paraphrasing-tools-to-make-your-content-unique/> [accessed on 29 Nov 2024]
- Halliday, Michael A. K. (2006). On the grammatical foundations of discourse. Webster, Jonathan J., ed. *Studies in Chinese Language*. London: Continuum, 346–363.
- Halliday, Michael A. K., Christian M. I. M. Matthiessen (2014). *Halliday's Introduction to Functional Grammar* (4th edn.). London: Routledge.
- Hirvela, Alan, Qian Du (2013). “Why am I paraphrasing?”: Undergraduate ESL writers' engagement with source-based academic writing and reading. *Journal of English for Academic Purposes* 12(2): 87–98. <https://doi.org/10.1016/j.jeap.2012.11.005>
- Hopkins, Diana, Tom Reid (2024). *The Academic Skills Handbook: Your Guide to Success in Writing, Thinking and Communicating at University*. London: Sage.
- Keck, Casey (2006). The use of paraphrase in summary writing: a comparison of L1 and L2 writers. *Journal of Second Language Writing* 15(4): 261–278. <https://doi.org/10.1016/j.jslw.2006.09.006>
- Keck, Casey (2010). How do university students attempt to avoid plagiarism? A grammatical analysis of undergraduate paraphrasing strategies. *Writing & Pedagogy* 2(2): 193–222. <https://doi.org/10.1558/wap.v2i2.193>
- Keck, Casey (2014). Copying, paraphrasing, and academic writing development: A re-examination of L1 and L2 summarization practices. *Journal of Second Language Writing* 25: 4–22. <https://doi.org/10.1016/j.jslw.2014.05.005>
- Kincaid, Ellie (2023). Turmoil at Sage journal as retractions mount. *Retraction Watch*. <https://retractionwatch.com/2023/09/14/turmoil-at-sage-journal-as-retractions-mount/>
- Lakoff, George (1968). *Deep and Surface Grammar*. Bloomington, IN: Indiana University Linguistics Club.
- Leech, Geoffrey (2000). Grammars of spoken English: New outcomes of corpus-oriented research. *Language Learning* 50(4): 675–724. <https://doi.org/10.1111/0023-8333.00143>
- Leong, Ping Alvin (2021). Writing in the sciences and humanities: A clause-complex perspective. *Word* 67(2): 137–158. <https://doi.org/10.1080/00437956.2021.1909866>
- Leong, Ping Alvin (2023a). Clause complexing in research-article abstracts: Comparing human- and AI-generated texts. *ExELL* 11(2): 99–132. <https://doi.org/10.2478/exell-2023-0008>
- Leong, Ping Alvin (2023b). Discipline-specific writing and embedded clauses: The case of cell biology and classics. *Word* 69(4): 362–379. <https://doi.org/10.1080/00437956.2023.2270875>
- Leong, Ping Alvin, Bee Geok Wee (2005). Investigating the clause complex: An analysis of exposition-type essays written by secondary school students in Singapore. *ITL – International Journal of Applied Linguistics* 150: 47–76. <https://doi.org/10.2143/ITL.150.0.2004372>

- Liu, Sen, Dunlai Lin (2022). Developing and validating an analytic rating scale for a paraphrase task. *Assessing Writing* 53: 100646. <https://doi.org/10.1016/j.asw.2022.100646>
- Marcus, Adam (2024). Cureus retracts paper for plagiarism following Retraction Watch inquiries. *Retraction Watch*. <https://retractionwatch.com/2024/04/30/cureus-retracts-paper-for-plagiarism-following-retraction-watch-inquiries/>
- McDonough, Kim, William J. Crawford, Jindarat de Vleeschauwer (2014). Summary writing in a Thai EFL university context. *Journal of Second Language Writing* 24: 20–32. <https://doi.org/10.1016/j.jslw.2014.03.001>
- Muraltia serpylloides (2024). This IEEE article contains several tortured phrases. *PubPeer*. <https://pubpeer.com/publications/36448420457798DEC2236D7048CF3C#1>
- Nerita vitiensis (2024a). Several passages in this article appear to be copied and machine-paraphrased. *PubPeer*. <https://pubpeer.com/publications/630A7C3D4F1EB3394D4E9CF6C99609#2>
- Nerita vitiensis (2024b). Some of text in this article appears to be plagiarized and machine-paraphrased from the 2017 article. *PubPeer*. <https://pubpeer.com/publications/36448420457798DEC2236D7048CF3C#3>
- Nesbitt, Christopher, Guenter Plum (1988). Probabilities in a systemic-functional grammar: The clause complex in English. Fawcett, Robin, P., D. Young, eds., *New Developments in Systemic Linguistics, Volume 2: Theory and Application*. London: Pinter, 6–38.
- Plavén-Sigra, Pontus, Granville James Matheson, Björn Christian Schiffler, William Hedley Thompson (2017). The readability of scientific texts is decreasing over time. *eLIFE* 6: e27725. <https://doi.org/10.7554/eLife.27725>
- Priyadarshini, Ishaani, Chase Cotton (2022). AI cannot understand memes: Experiments with OCR and facial emotions. *Computers, Materials & Continua* 70(1): 781–800. <https://doi.org/10.32604/cmc.2022.019284>
- Roe, Jasper, Mike Perkins (2022). What are automated paraphrasing tools and how do we address them? A review of a growing threat to academic integrity. *International Journal for Educational Integrity* 18: 15. <https://doi.org/10.1007/s40979-022-00109-w>
- Rogerson, Ann M., Grace McCarthy (2017). Using Internet based paraphrasing tools: Original work, patchwriting or facilitated plagiarism? *International Journal for Educational Integrity* 13(2): 2. <https://doi.org/10.1007/s40979-016-0013-y>
- Scimago (2023). Scimago journal and country rank. <http://www.scimagojr.com/>
- Sellami-Baklouti, Akila (2011). The impact of genre and disciplinary differences on structural choice: Taxis in research article abstracts. *Text & Talk* 31(5): 503–523. <https://doi.org/10.1515/text.2011.025>
- Sharpe, Pamela, J. (2024). *TOEFL IBT Premium*. Lauderdale, FL: Kaplan North America.
- Shi, Ling (2004). Textual borrowing in second-language writing. *Written Communication* 21(2): 171–200. <https://doi.org/10.1177/0741088303262846>
- Shi, Ling, Ismaeil Fazal, Nasrin Kowkabi (2018). Paraphrasing to transform knowledge in advanced graduate student writing. *English for Specific Purposes* 51: 31–44. <https://doi.org/10.1016/j.esp.2018.03.001>
- Sun, Yu-Chih, Fang-Ying Yang (2015). Uncovering published authors' text-borrowing practices: paraphrasing strategies, sources, and self-plagiarism. *Journal of English for Academic Purposes* 20: 224–236. <https://doi.org/10.1016/j.jeap.2015.05.003>
- Uemlianin, Ivan A. (2000). Engaging text: Assessing paraphrase and understanding. *Studies in Higher Education* 25(3): 347–358. <https://doi.org/10.1080/713696160>
- Warren, Tom (2024). Microsoft's new Copilot key is the first big change to Windows keyboards in 30 years. *The Verge*. <https://www.theverge.com/2024/1/4/24023809/microsoft-copilot-key-keyboard-windows-laptops-pcs>

- Wette, Rosemary (2010). Evaluating student learning in a university-level EAP unit on writing using sources. *Journal of Second Language Writing* 19(3): 158–177. <https://doi.org/10.1016/j.jslw.2010.06.002>
- Wu, Lijun, Xu Tan, Tao Qin, Jianhuang Lai, Tie-Yan Liu (2019). Beyond error propagation: Language branching also affects the accuracy of sequence generation. *IEEE/ACM Transactions on Audio, Speech, and Language Processing* 27(12): 1868–1879. <https://doi.org/10.1109/TASLP.2019.2933727>
- Yahia, Intissar, Joy L. Egbert (2023). Supporting non-native-English speaking graduate students with academic writing skills: A case study of the explicit instructional use of paraphrasing guidelines. *Journal of Writing Research* 14(3): 305–341. <https://doi.org/10.17239/jowr-2023.14.03.01>
- Zaiontz, Charles (2022). Real statistics using Excel. <http://www.real-statistics.com/>
- Zhang, Qing, David Y. Wang, Geoffrey M. Voelker (2014). DSpin: Detecting automatically spun content on the web. *NDSS Symposium 2014*, San Diego, California. <https://www.ndss-symposium.org/ndss2014/ndss-2014-programme/dspin-detecting-automatically-spun-content-web/>

Author's address:

Alvin Ping Leong
Language and Communication Centre
School of Humanities, Nanyang Technological University
48 Nanyang Avenue, Singapore 639818
Singapore
e-mail: alvin.leong@ntu.edu.sg

Appendix

Symbols used in the analysis

Symbol	Description
	Clause complex / simple
	Clause separator
<<...>>	Inserted clause
[[...]]	Embedded clause
1, 2, 3, ...	Parataxis
$\alpha, \beta, \gamma, \dots$	Hypotaxis
=	Elaboration
+	Extension
×	Enhancement
'E' superscript (e.g., α^E, β^E)	Embedded clause relations