

# How Do EFL Learners Regulate Corpus Consultation? A Qualitative Study of Self-Regulated Learning in Data-Driven Writing

Ruiyi Wu<sup>1</sup>, Wei Lun Wong<sup>1\*</sup>, Peng Liu<sup>1</sup>

<sup>1</sup> Faculty of education, Universiti Kebangsaan Malaysia, 43600 Bangi Selangor, Malaysia

Ruiyi Wu: P161113@siswa.ukm.edu.my

Wei Lun Wong: colinw@ukm.edu.my

Peng Liu: P115525@siswa.ukm.edu.my

\*Corresponding e-mail: colinw@ukm.edu.my

## HOW TO CITE:

Ruiyi Wu, Wei Lun Wong, Peng Liu (2026). How Do EFL Learners Regulate Corpus Consultation? A Qualitative Study of Self-Regulated Learning in Data-Driven Writing  
International Journal of Special Education, 41(19s), 1373 - 1384

## ABSTRACT

Corpus-based language pedagogy has shown some positive results in EFL writing, but how learners organise the consultation process itself has rarely been studied. This paper investigates how learners regulate corpus consultation in data-driven learning under an examination-oriented writing task. Twelve Chinese undergraduates preparing for the CET-4, sampled from the higher and lower ends of the proficiency range, participated in eight Sketch Engine-supported argumentative writing sessions. Guided learner logs of 412 search episodes and retrospective interviews were analysed using hybrid thematic analysis, combining the stages of self-regulated learning as an a priori framework with inductive coding in the phases. The analysis found seven subthemes that describe consultation as a cycle in which learners frame searchable problems around multiword units, monitor and revise queries based on concordance evidence, and convert search results into modified routines. Higher-proficiency learners planned queries at the phrase level and weighed evidence across lines; lower-proficiency learners searched in fragments and readily accepted evidence, and individual trajectories showed that these patterns were modifiable habits rather than fixed attributes. Based on the above results, phraseological competence and regulatory skills develop together; therefore, teaching should focus on practising search decisions rather than tool operations and establish consultation strategy as enduring content for corpus-informed writing pedagogy.

## COPYRIGHT STATEMENT:

Copyright: © 2026 Authors.  
Open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

**Keywords:** Corpus-based language pedagogy; data-driven learning; phraseological competence; self-regulated learning; EFL writing; CET-4; Sketch Engine

## Introduction

### 1.1 Data-driven Learning and Corpus Consultation in EFL Writing

Corpus tools have gradually moved from the linguist's desk to the language classroom, and their pedagogical applications have matured into a well-known area of research with their own goals [1]. In this area, data-driven learning positions learners as researchers who directly query corpora and induce patterns of use from concordance evidence; such an arrangement has spread widely in computer-assisted language learning environments [2]. Reviews of classroom implementations report advantages for vocabulary, grammar and writing, but also recurring problems in query construction and evidence interpretation [3]. Writing has also been shown to respond well to these methods; corpus-supported revision can increase the fluency and lexical richness of EFL compositions [4], and corpus consultation can provide learners with a reference resource for academic writing tasks [5]. Much of what learners are looking for in a corpus is phraseological rather than lexical in the narrow sense; that is to say, the combinations that make writing sound natural, rather than individual word choices, are what distinguish stronger argumentative texts, and formulaic competence has been shown to directly affect the quality of EFL argumentative writing [6]. Consultation, in other words, is a search for acceptable word company, and multiword patterns are at the core of what learners do with corpus interfaces.

## **1.2 Self-regulated Learning in Technology-Mediated Writing and the Process Gap**

Since consultation puts control in the hands of the learners, its success depends on how they manage their own activity. A typical model of this kind of management is self-regulated learning; it is a cycle in which learners set goals, monitor the implementation and evaluate the results [7]. Research has been conducted using the above framework. Strategy use predicts EFL writing performance among Chinese university students [8], validated instruments now measure self-regulatory writing strategies at various school levels [9], and structural models confirm the consistency of these strategies for advanced learners [10]. Instructional research has added that teaching metacognitive strategies can improve students' writing ability and motivation

[11], and individual differences affect the strategies that students choose [12]. However, the above studies have treated writing as an organised activity and used questionnaires to measure strategies; thus, the moment-to-moment management of a corpus search has been neglected. A questionnaire can only confirm that a learner generally reports planning or monitoring; it cannot show what planning looks like when the object of regulation is a query rather than an essay, or what monitoring means when feedback arrives as a screen of concordance lines. Research on data-driven learning shows the same blind spot; that is, most of the existing studies are still outcome-oriented and have paid little attention to learners' interactions with corpora and their decision-making during consultation [13]. What occurs between opening a corpus tool and revising a sentence, what goals learners bring to a query, how they respond when evidence is lacking, and what they carry forward to the next search remain a black box at the intersection of the two well-developed literatures.

## **1.3 The present study**

This paper will explore how EFL learners manage corpus consultation in the course of data-driven argumentative writing through qualitative research. Twelve Chinese undergraduates preparing for the CET-4 participated in eight Sketch Engine-assisted writing sessions, and their guided logs and retrospective interviews were analysed according to the regulatory cycle. Four questions direct the inquiry, namely how learners set goals and plan queries before searching, how they monitor and adjust their searches during consultation of the corpus, how they evaluate search results and change subsequent strategies, and how these regulatory patterns differ for higher-proficiency and lower-proficiency learners. By tracing the regulation at the level of individual search episodes, this study hopes to open the consultation process to inspection and derive implications for corpus-based writing instruction.

## **2. Data and Methods**

### **2.1 Research Context and Participants**

The study took place in a CET-4 writing preparation course at a comprehensive university in southern China, where the teaching of argumentative writing has been using corpus tools more and more often [14]. The eight weekly classes of the course each had a short training session on Sketch Engine and an argumentative writing task that required students to use the corpus for drafting and revising. All students in the course had received at least six years of formal English education before university and none had prior experience with corpus consultation.

Twelve undergraduate students were selected by purposive extreme-case sampling to show the differences that are not obvious in a more

uniform group [15]. Six participants who scored 550 or higher in the CET-4 were designated as the high-proficiency group, and six who scored 480 or lower were placed in the low-proficiency group. The sample included students in the second or third year of engineering, business and humanities. Participants are referred to as S1 to S12 throughout, with S1 to S6 being in the upper group and S7 to S12 in the lower group. Table 1 shows the demographics and ability levels of the participants. The two ends of the score distribution were deliberately set apart to determine whether the regulation of corpus consultation differed by proficiency, and at the same time, a small enough sample was obtained for in-depth qualitative analysis.

**Table 1. Participant Profiles**

| ID  | Gender | Year | Major                    | CET-4 score | Proficiency group |
|-----|--------|------|--------------------------|-------------|-------------------|
| S1  | F      | 2    | English-related business | 592         | Higher            |
| S2  | M      | 3    | Computer engineering     | 575         | Higher            |
| S3  | F      | 2    | Accounting               | 568         | Higher            |
| S4  | M      | 2    | Civil engineering        | 561         | Higher            |
| S5  | F      | 3    | Chinese literature       | 556         | Higher            |
| S6  | M      | 2    | International trade      | 550         | Higher            |
| S7  | F      | 2    | Mechanical engineering   | 478         | Lower             |
| S8  | M      | 3    | Logistics management     | 471         | Lower             |
| S9  | M      | 2    | Materials science        | 463         | Lower             |
| S10 | F      | 2    | Public administration    | 452         | Lower             |
| S11 | M      | 3    | Electrical engineering   | 441         | Lower             |
| S12 | F      | 2    | Tourism management       | 436         | Lower             |

No

te. Scores refer to the most recent CET-4 administration before the course. None of the participants had previous experience with corpus consultation. All names are replaced by codes.

## 2.2 Data-driven Writing Tasks

All eight sessions were based on argumentative

prompts taken from past CET-4 exams, and well-known topics such as online classes, part-time work and environmental awareness were covered. Learners drafted a response of 120-180 words in class, meeting the requirements of the examination, and then revised it after consulting Sketch Engine, which was chosen for its simple

interface and the combination of concordancing, Word Sketch summaries, and collocation candidate lists that have shown good results in previous work on learner collocation use [16]. The preparatory training introduced three core operations. Learners practised running concordance searches to observe multiword patterns in context, generated Word Sketches to compare the collocational behaviour of near-synonyms, and used collocation candidate lists to check whether a combination produced in their draft was attested in the reference corpus.

Task sheets focused on phraseological units in argumentative prose, such as stance markers, cause-and-effect expressions, and verb-noun collocations, rather than individual vocabulary items. Learners independently determined when to use the corpus, what to search for, and whether to accept, adapt or discard the evidence they found. This open design maintained the self-

directed nature of the consultation that the study aimed to observe, and heavily scripted search procedures would have preempted the very regulatory decisions under investigation.

### 2.3 Data Collection: Learner Logs and Semi-structured Interviews

The two complementary data sources recorded the consultation process both during and after the event. In each session, the participants completed a guided learner log with four fields: the goal of that search, the exact query entered, how the retrieved examples were weighted and selected, and the resulting modification made to the draft. The four-field format kept the entries connected to concrete search cases and was light enough to be filled in while writing. Over the eight sessions, the twelve participants produced 96 logs that contained 412 documented search episodes. Table 2 shows an overview of the two data sources and their volumes.

**Table 2. Overview of Data Sources**

| Data source                | Timing                                              | Structure                                                                                                                     | Volume                                                                                                                       |
|----------------------------|-----------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------|
| Guided learner logs        | Completed during each of the eight writing sessions | Four fields per search episode: search goal; query entered; selection among retrieved examples; resulting change to the draft | 96 logs; 412 documented search episodes                                                                                      |
| Semi-structured interviews | Within two weeks of the final session               | Reconstruction of logged searches; responses to unhelpful results; evaluation of consultation over the course                 | 12 interviews of 30–40 minutes; approximately 58,000 Chinese characters transcribed, translated with back-translation checks |

About two weeks after the last class, each student took part in a one-on-one interview of about 30–40 minutes. Interviews were conducted to review some representative log entries and, along with the participants, reconstructed the reasons for their searches, described how they responded when a query returned no results, and evaluated the effect of consultation on their writing over the past eight weeks. Interviews were conducted in Chinese so that the participants could speak freely without the pressure of language; they were

audio-recorded, transcribed verbatim, and then translated into English by the researcher, and a sample of the transcripts underwent back-translation to ensure semantic consistency. The entire interview corpus was about 58,000 Chinese characters before translation.

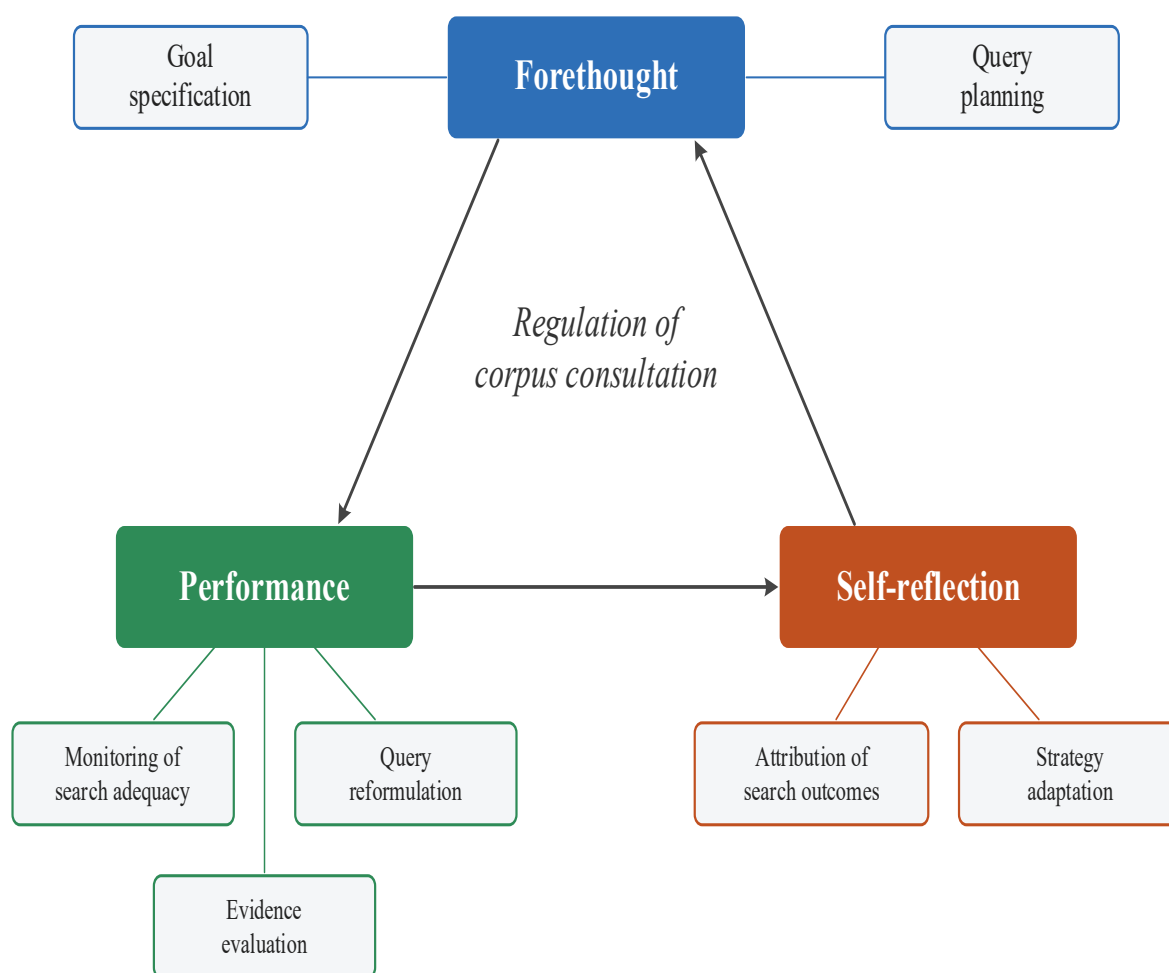
### 2.4 Data Analysis

Both deduction and induction were used in the analysis, and the stages of analysis followed those put forward by Braun and Clarke [17].

Zimmerman's cyclical model provided three a priori categories, and thus every meaningful part of the logs and transcripts was classified as forethought, performance or self-reflection according to the stage of consultation it described. Subthemes in each stage were derived inductively from multiple readings, rather than being predefined. Segments coded as 'performance' gradually divided into monitoring of search adequacy, reformulation of failed queries, and evaluation of concordance evidence; forethought segments split into goal specification and query planning; and self-reflection segments split into attribution of search outcomes and strategy adaptation.

Coding proceeded in three passes through the data set. An initial pass generated phase

assignments, a later pass created candidate subthemes within the phases, and a final pass merged overlapping subthemes and selected representative excerpts. Log entries and interview segments from the same participant were linked during analysis to read retrospective accounts in light of the contemporaneous record of the search they described. Figure 1 shows the analytical framework, with the three-phase cycle and the subthemes developed in each phase; Table 3 shows the full coding scheme with definitions and sample coded extracts. Divergent cases that did not fit the main pattern were kept and coded instead of being excluded, as they provided a reference for comparison among the proficiency groups.



A priori phases (colored) follow the cyclical model of self-regulated learning; subthemes (light boxes) were generated inductively.

Figure 1. Analytical Framework Combining A Priori Regulatory Stages with Inductively Generated Subthemes

Table 3. Coding Scheme

| Phase (a priori) | Subtheme (inductive)           | Definition                                                                                 | Sample coded extract                                                                                                |
|------------------|--------------------------------|--------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------|
| Forethought      | Goal specification             | Identifying the linguistic problem a search is meant to solve before entering a query      | I wanted to know which verb goes with “impact”, so that was what I searched for. (S2, interview)                    |
| Forethought      | Query planning                 | Deciding the form of the query, including the choice between single words and word strings | Search goal: check if “put forward a solution” is used; query: put forward. (S5, log, session 4)                    |
| Performance      | Monitoring of search adequacy  | Judging during a search whether the returned lines match the intended problem              | The lines were all about news, not student writing, so I felt the examples did not fit my sentence. (S8, interview) |
| Performance      | Query reformulation            | Revising a failed or noisy query by shortening, substituting, or changing search mode      | Too many results for “make”; changed to Word Sketch and looked only at objects. (S1, log, session 3)                |
| Performance      | Evidence evaluation            | Weighing retrieved examples before accepting, adapting, or rejecting a pattern             | Many lines showed “play a role in”, so I did not use “play roles in”. (S11, interview)                              |
| Self-reflection  | Attribution of search outcomes | Explaining why a search succeeded or failed after the writing task                         | That search failed because I typed the whole sentence from Chinese. (S10, interview)                                |
| Self-reflection  | Strategy adaptation            | Carrying a revised search routine into subsequent tasks                                    | After session five I always searched the noun first and read the verbs around it. (S6, interview)                   |

Note. Phases were deduced from the cyclical model of self-regulated learning, and subthemes were inductively generated within these phases.

## 2.5 Trustworthiness

Several measures supported the reliability of the analysis and aligned with the current discussions on quality practice in thematic analysis [18]. A second coder independently coded 20% of the data and reached an 87% agreement with the first coder at the subtheme level; all discrepancies were resolved through discussion before the scheme was finalised. Member checking was performed by returning the translated interview summaries to the participants for confirmation,

and an audit trail of coding decisions and scheme revisions was kept throughout. Interview transcripts are not published to protect the privacy of the participants.

## 3. Results

Seven subthemes were identified in the three stages of the regulatory cycle through analysis, and Figure 2 shows how they are organised and what differences exist among proficiency groups. The following accounts move through the cycle in sequence, drawing from interview excerpts and log data, and participant codes indicate the origin of each excerpt

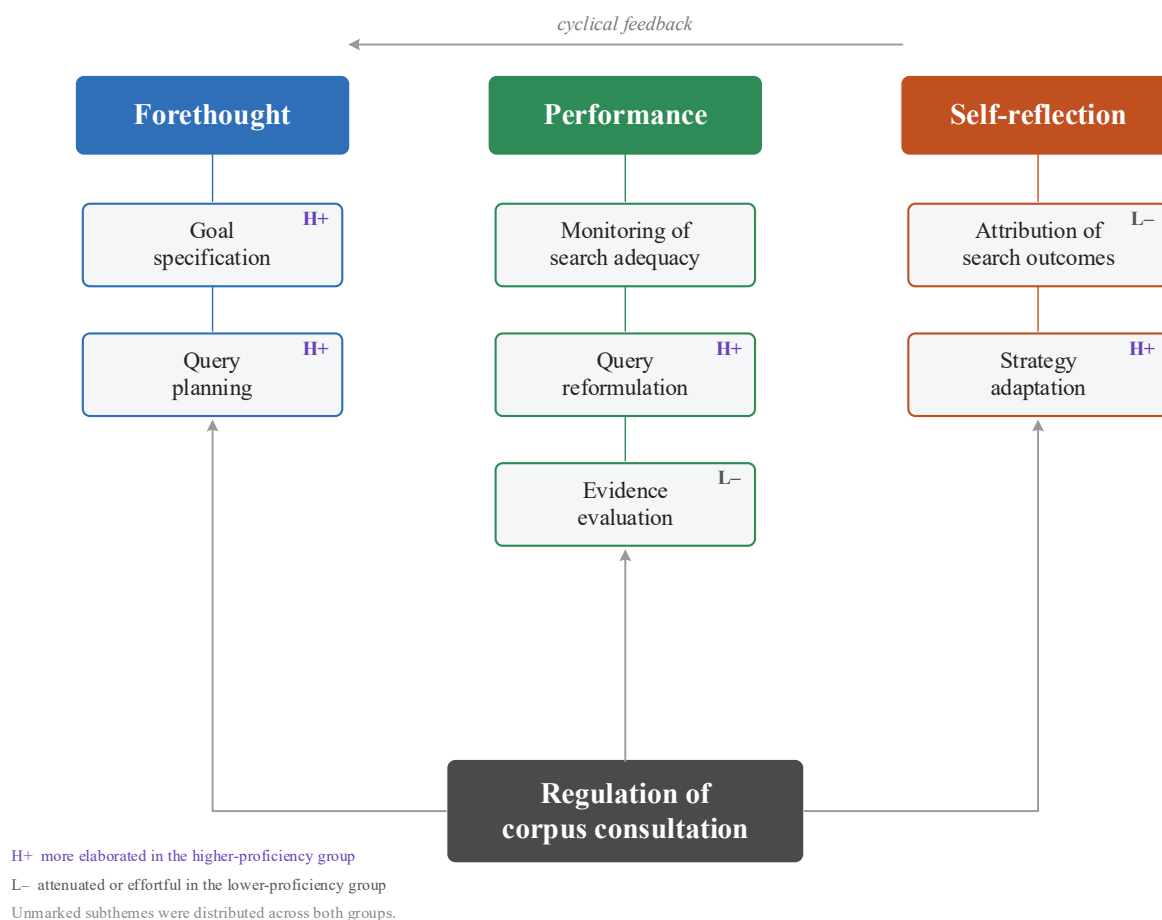


Figure 2. Thematic Map of Regulatory Subthemes with Proficiency-Group Differences

### 3.1 Forethought: Goal Setting and Query Planning

Regulation started before any query was typed in. High-proficiency learners generally began their work with the corpus by formulating a problem at the level of a word combination rather than a single word. S5 explained that she had rehearsed the target expression mentally before opening the tool. "Before searching, I already had a guess, maybe 'raise people's awareness', and I went to the corpus mainly to confirm whether native writers actually put these words together." S2 described a similar hypothesis-testing attitude. "I did not search to find new words. I looked up the words I wanted to use and checked the verb before 'impact', because I was not sure if it was 'exert'."

Goal setting for the lower-proficiency group was often too general and difficult to convert into an actionable query. S12 recalled that her early searches were based on a general desire rather

than a specific object. "At the beginning, I only wanted to improve the expression of my essay, so I typed the topic word 'environment' and hoped that good expressions would appear by themselves." S7 reported that she framed goals word by word from a Chinese draft sentence. "I first wrote the sentence in Chinese in my head, then looked up each English word separately, so my goal was indeed several small goals."

Planning the form of the query also showed the same contrast. High-proficiency participants often reduced their intended expression to a shorter string to keep the results manageable, as shown in one log entry.

Log excerpt (S5, session 4). Goal recorded as checking whether "put forward a solution" is natural; query entered as "put forward"; selection noted that most lines continued with "proposal" or "suggestion" rather than "solution"; change recorded as replacing the phrase with "propose a solution".

S6 summarised the reason for the above trimming. "If I typed the whole expression and it was wrong, nothing would be returned; so I searched for the shorter part and let the corpus complete it."

### **3.2 Performance: Monitoring, Query Reformulation and Evidence Evaluation**

Once a search was running, attention turned to judging the returned lines against the writing problem at hand. Monitoring of search adequacy occurred in both groups, but the reasons for triggering were different. S1 checked whether the register of the concordance lines matched that of argumentative prose. "When the lines came from fiction, I ignored them, because a phrase natural in a novel might look strange in an exam essay." S8 used a rougher signal, the sheer readability of the display. "When the page showed forty lines, I could not read them all; I only felt that the search was good when there were few and short lines."

Unsatisfactory results led to reformulation, and here the higher-proficiency group showed a wider repertoire. Participants shortened the strings, substituted lemmas, and switched search modes instead of giving up. One log entry records a sequence of three moves in a single episode.

Log excerpt (S3, session 6). Goal recorded as finding the verb for "conclusion"; query moved from "draw a conclusion that" to "draw a conclusion" and then to a Word Sketch for "conclusion"; selection noted "draw" and "reach" as the strongest verbs; change recorded as keeping "draw" and deleting "that".

S4 considered mode switching to be a deliberate economy. "Concordance gave me too many sentences with 'take', so I opened the Word Sketch instead and looked only at the object column; that answered my question in one screen." Although the collocation candidate list was introduced in training, it rarely appeared in the logs, and consultation settled on concordance and Word Sketch as the working pair for both groups. Reformulation in the lower-proficiency group more often stopped after a single adjustment. S11 thought that persistence was

about the first page of results. "If the first page didn't show what I wanted, I would usually close it and use a word from my textbook; trying another search felt like wasting writing time."

Evaluation of retrieved evidence showed the largest gap in this stage. Higher-proficiency participants considered frequency and consistency across lines as reasons to choose. S3 mentioned that she had weighed two patterns. "Many lines showed 'result in' followed by bad things, such as pollution, but 'lead to' appeared with both good and bad results, so I chose 'lead to' because my sentence was about a positive change." Several lower-proficiency participants treated any attested example as sufficient license instead. S9 said he would accept the top line without reading further. "The first sentence already contained 'more and more serious', so I copied its shape and did not read the other lines."

### **3.3 Self-reflection: Attribution and Strategy Adaptation**

After each task, the participants differed in how they explained why a search had worked or not. High-proficiency students tended to look for the causes in the form of the query itself. S2 attributed his improved hit rate to a change in the search unit. "My searches worked better in the later weeks because I stopped searching single words and started searching two or three words together; this is what the corpus is good at." Attribution among lower-proficiency participants was more frequently directed at the tool or at English in general, although the logs sometimes preserved the evidence for a more specific diagnosis. The entry S10 made in an early session recorded a query that consisted of a full translated sentence about the advantages of online learning, which returned no results and did not lead to a revision. Her interview account returned to that episode with a clearer explanation. "That search failed because I typed the whole sentence in Chinese, and now I understand that the corpus cannot read my sentence; it can only match the words I give it."

Reflection fed forward to modified routines. S6 institutionalised a search order that persisted until the end of the course. "My rule after the middle

weeks was noun first, then read which verbs appeared around it, and only then decide my own verb." Adaptation for the lower-proficiency group sometimes took the form of retreat rather than refinement. S12 reduced consultation to a last resort. "Later, I only used the corpus when the teacher marked a phrase as incorrect; otherwise, searching during writing would make me lose my thoughts."

### 3.4 Regulation patterns across proficiency levels

Across the cycle, the difference between the two groups lay less in whether they regulated and more in where the regulation was concentrated. High-proficiency participants distributed their efforts across all three stages, making phrase-level hypotheses in the search phase and closing tasks with transferable routines; lower-proficiency participants focused their efforts on the performance stage, and evaluation remained a relatively weak link, as shown by the shaded markers in Figure 2. The patterns inside each group were not uniform. S9 showed first-line copying in the early sessions and, by the end of the last two weeks, had developed a systematic comparison habit. "In the last two essays I read at least ten lines and chose the pattern that appeared most frequently, although reading them was slow for me." A higher-proficiency participant, in contrast, described planning as something that could be dropped once the topics became familiar. S1 reported that routine topics no longer required preparation. "For the easy topics, I stopped planning searches altogether; I just wrote and only checked when a phrase felt uncertain." Regulation thus appeared as a repertoire responsive to task familiarity and confidence rather than a fixed attribute of proficiency.

## 4. Discussion

Corpus consultation is, in fact, a whole cycle of regulation rather than a simple look-up function for writing. The search episodes started with problem definition, were continuously observed and modified, and concluded with attribution adjustments to future searches; thus, they followed a recursive pattern of self-regulated

learning, and reflection on one search informed the goals for the next [19]. What the data add to this general picture is the observation that the cycle operated at the grain of individual queries. Participants did not merely regulate a writing task that happened to include corpus use; rather, they regulated each search as a mini-task with its own goal, execution and evaluation, and thus a single essay could contain several complete regulatory loops. This granularity helps explain why decades of research on data-driven learning have produced consistently positive but strikingly varied results among learners and in different environments [20]. If the benefit of consultation is determined by dozens of small regulatory decisions, then two learners given the same tools and the same training may derive different values from the same concordance output, and outcome measures alone will not explain this difference. The current account places the source of the variation in the search stage, such as shortening a query before inputting it or reading past the first line and then accepting a pattern.

The contrast between the two proficiency groups makes this interpretation clearer. High-proficiency participants used phrase-level hypotheses to address the corpus and judged the evidence based on frequency and consistency; low-proficiency participants searched word by word and considered any attested example sufficient support. This asymmetry is in line with the evidence that the processing of formulaic sequences becomes more holistic with increased proficiency, and stronger learners can perceive and manipulate multiword units as wholes [21]. A learner who mentally represents "put forward a proposal" as a single unit can formulate a query at that level and recognise whether concordance lines support or contradict it; a learner who operates word by word must construct the same judgment from fragments. Developmental studies of phrase-frames also point in the same direction, showing that the variability and functional range of learners' recurrent frames expand gradually rather than appearing suddenly after instruction [22]. The regulation observed here is therefore not a free-standing skill layered on top of language knowledge. Phraseological

competence provides the basic units for planning, evaluation and attribution, and thus the quality of regulation and phraseological development are mutually constitutive; each provides material support for the other. The paths of the individual participants also support this reciprocal reading; the lower-proficiency student who learned systematic line comparison late in the course first changed a regulatory habit, and that habit itself required perceiving recurrence across lines.

These observations have some applications in corpus-based language teaching. Although the training in this study covered tool operation, the differences that mattered were regulatory; therefore, instruction should practice decision-making rather than clicks. Explicit modelling of query trimming, of reading beyond the first line, and of diagnosing failed searches would target exactly the moves that separated the two groups, and such strategy-focused instruction has shown good results in writing classrooms when integrated into regular tasks rather than delivered as a standalone module [23]. The finding that reflection sometimes led to retreat rather than refinement deserves special attention. A student who decides that consultation interrupts writing has in fact drawn a reasonable lesson from a real difficulty. Teachers who foresee this trend can validate the check after drafting as a legitimate form of consultation, rather than considering reduced mid-writing search a failure. Generally speaking, the design lesson is that pedagogical materials should scaffold the regulatory cycle itself; in this regard, recent demands to close the gap between corpus research and classroom practice have focused on what learners actually do with corpus tools [24]. The development of generative chatbots has not changed these reasons. Conversational systems can now answer many phraseological questions directly, but the arguments for the continued value of corpus consultation are based on verifiable and frequency-grounded characteristics of corpus evidence [25], and learners themselves describe chatbot feedback as an addition to, rather than a substitute for, their own checking routines [26]. The regulatory repertoire documented here sets

up a specific question, judges the returned evidence, and diagnoses failure; this is precisely the repertoire that transfers to the critical use of any language resource and positions consultation strategy as enduring curriculum content rather than tool-specific technique.

## 5. Conclusion

This paper investigates how EFL learners regulate corpus consultation in the process of data-driven argumentative writing. Based on learner logs and interviews with twelve Chinese undergraduates preparing for the CET-4, the analysis traced regulation across forethought, performance, and self-reflection, and identified seven subthemes that together describe consultation as a cycle of framing searchable problems, adjusting queries according to retrieved evidence, and converting the results of individual searches into revised routines.

The main contribution of this paper is to shift the research focus from whether corpus consultation can improve writing to how learners use it. Regulation worked at the level of a single search episode, and its quality changed systematically with proficiency. Higher-proficiency learners planned queries around multiword units and weighed evidence across concordance lines; lower-proficiency learners searched in smaller fragments and accepted evidence more readily, but individual trajectories showed that these patterns were habits open to change rather than fixed endowments. Phraseological competence and regulatory skills have emerged as interdependent, and each provides the units and routines for the development of the other.

Corpus training focused on tool operations does not change the decision-making layer, according to the study. Preparing learners to trim queries, read beyond the first line, and diagnose failed searches addresses the moves on which productive consultation actually depended, and this decision-oriented preparation retains its value as learners extend the same habits to new language resources.

Some deficiencies are also present. The participants were all from the same course at one

university, the data reflect eight weeks of practice with a single corpus platform, and self-reports may not cover all regulatory changes. In the future, one can extend the time of observation for the regulation, directly obtain the search

behaviour via screen recording, and determine whether instruction aimed at the regulatory cycle changes both the quality of consultation and written results.

## References

- O'Keeffe, A. (2021). Data-driven learning: A call for a broader research gaze. *Language Teaching*, 54(2), 259–272. <https://doi.org/10.1017/S0261444820000245>
- Pérez-Paredes, P. (2022). A systematic review of the uses and spread of corpora and data-driven learning in CALL research during 2011–2015. *Computer Assisted Language Learning*, 35(1–2), 36–61. <https://doi.org/10.1080/09588221.2019.1667832>
- Lusta, A., Demirel, Ö., & Mohammadzadeh, B. (2023). Language corpus and data-driven learning (DDL) in language classrooms: A systematic review. *Heliyon*, 9(12), e22731. <https://doi.org/10.1016/j.heliyon.2023.e22731>
- Şahin Kızıl, A. (2023). Data-driven learning: English as a foreign language writing and complexity, accuracy and fluency measures. *Journal of Computer Assisted Learning*, 39(4), 1382–1395. <https://doi.org/10.1111/jcal.12807>
- Emir, G., & Yangın Ekşi, G. (2023). Corpus used as a data-driven learning tool in L2 academic writing: Evidence from Turkish contexts. *TEFLIN Journal*, 34(2), 209–225. <https://doi.org/10.15639/teflinjournal.v34i2/209-225>
- Li, H., & Yao, Y. (2023). Formulaic competence in college-level Asian English learner's argumentative writing: Examining the effects of language background and topic. *The Asia-Pacific Education Researcher*, 32(6), 793–803. <https://doi.org/10.1007/s40299-022-00695-w>
- Zimmerman, B. J. (2002). Becoming a self-regulated learner: An overview. *Theory Into Practice*, 41(2), 64–70. [https://doi.org/10.1207/s15430421tip4102\\_2](https://doi.org/10.1207/s15430421tip4102_2)
- Shen, B., & Bai, B. (2024). Chinese university students' self-regulated writing strategy use and EFL writing performance: Influences of self-efficacy, gender, and major. *Applied Linguistics Review*, 15(1), 161–188. <https://doi.org/10.1515/applirev-2020-0103>
- Teng, M. F., Wang, C., & Zhang, L. J. (2022). Assessing self-regulatory writing strategies and their predictive effects on young EFL learners' writing performance. *Assessing Writing*, 51, Article 100573. <https://doi.org/10.1016/j.asw.2021.100573>
- Wang, X., Ma, J., Li, X., & Shen, X. (2023). Validation of self-regulated writing strategies for advanced EFL learners in China: A structural equation modeling analysis. *European Journal of Investigation in Health, Psychology and Education*, 13(4), 776–795. <https://doi.org/10.3390/ejihpe13040059>
- Han, L. (2024). Metacognitive writing strategy instruction in the EFL context: Focus on writing performance and motivation. *SAGE Open*, 14(2). <https://doi.org/10.1177/21582440241257081>
- Umamah, A., El Khoiri, N., Widiati, U., & Wulyani, A. N. (2022). EFL university students' self-regulated writing strategies: The role of individual differences. *Journal of Language and Education*, 8(4), 182–193. <https://doi.org/10.17323/jle.2022.13339>
- Yang, Y., & Ren, H. (2025). The efficacy of the corpus-based error correction method on revision in writing classrooms. *PLOS ONE*, 20(3), e0317574. <https://doi.org/10.1371/journal.pone.0317574>
- Zhao, Y., & Dong, J. (2026). Tracing the diachronic effects of data-driven learning on lexical complexity in EFL learners' argumentative writing. *ReCALL*, 38(3), 324–344. <https://doi.org/10.1017/S0958344026100494>
- Patton, M. Q. (2015). *Qualitative research & evaluation methods: Integrating theory and practice* (4th ed.). SAGE.
- Du, X., Afzaal, M., & Al Fadda, H. (2022). Collocation use in EFL learners' writing across multiple language proficiencies: A corpus-driven study. *Frontiers in Psychology*, 13, 752134. <https://doi.org/10.3389/fpsyg.2022.752134>
- Braun, V., & Clarke, V. (2021). *Thematic analysis: A practical guide*. SAGE.
- Braun, V., & Clarke, V. (2021). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative Research in Psychology*, 18(3), 328–352. <https://doi.org/10.1080/14780887.2020.1769238>

- Panadero, E. (2017). A review of self-regulated learning: Six models and four directions for research. *Frontiers in Psychology*, 8, 422. <https://doi.org/10.3389/fpsyg.2017.00422>
- Boulton, A., & Vyatkina, N. (2021). Thirty years of data-driven learning: Taking stock and charting new directions over time. *Language Learning & Technology*, 25(3), 66–89. <https://doi.org/10.64152/10125/73450>
- Fan, K., & Wang, H. (2024). The productive processing of formulaic sequences by second language learners in writing. *Frontiers in Psychology*, 15, 1281926. <https://doi.org/10.3389/fpsyg.2024.1281926>
- Yan, J., Li, Q., & Jiang, J. (2024). The development of phrase-frames in EFL learners' essays: Variability, structures, and functions. *International Journal of Applied Linguistics*, 34(4), 1689–1708. <https://doi.org/10.1111/ijal.12589>
- Sun, J., Motevalli, S., Chan, N. N., & Khodi, A. (2024). Implementation of self-regulated learning writing module: Amplifying motivation and mitigating anxiety among EFL learners. *Forum for Linguistic Studies*, 6(3), 89–109. <https://doi.org/10.30564/fls.v6i3.6600>
- Crosthwaite, P. (Ed.). (2024). *Corpora for language learning: Bridging the research-practice divide*. Routledge. <https://doi.org/10.4324/9781003413301>
- Crosthwaite, P., & Baisa, V. (2023). Generative AI and the end of corpus-assisted data-driven learning? Not so fast! *Applied Corpus Linguistics*, 3(3), 100066. <https://doi.org/10.1016/j.acorp.2023.100066>
- Teng, M. F. (2024). "ChatGPT is the companion, not enemies": EFL learners' perceptions and experiences in using ChatGPT for feedback in writing. *Computers and Education: Artificial Intelligence*, 7, Article 100270. <https://doi.org/10.1016/j.caeai.2024.100270>