



The International Journal of Language Studies

ISSN: 3078-2244 (Online)

2026, Vol. 3, No. 2, pp. 1–20

DOI: <https://doi.org/10.60087/ijls.vol3.n2.001>

Journal homepage: <https://ijlangstudies.org>



A Comparative Corpus-Based Study of Vocabulary, Sentence Structure, and Tense-Aspect Usage in the Dialogues of “Global Success” and “Family and Friends” for Vietnamese Primary Learners

Dương Minh Hoàng^{1*}

¹ Faculty of Foreign Languages, Dalat University, Da Lat, Vietnam

*Corresponding author: Dương Minh Hoàng — e-mail: duongminhhoang119@gmail.com

Abstract

This study presents a corpus-based comparison of dialogue content in two primary English textbook series widely used in Vietnam: Family and Friends and Global Success (Grades 3–5). Examining three dimensions of linguistic input—vocabulary range, sentence complexity, and tense-aspect usage—the study analyzes a corpus of 1,143 utterances drawn exhaustively from nine Student’s Books. Utterances were analyzed using AntConc 4.3.1, a Moving-Average Type-Token Ratio (MATTR) procedure, and manual coding of tense-aspect and communicative functions. Results indicate that Family and Friends offers a wider lexical range, longer and more structurally diverse utterances, and a broader functional deployment of verb forms. Conversely, Global Success concentrates on the Present Simple, repeats a limited set of question-and-answer patterns, and confines the Past Simple almost exclusively to asking and answering. Interpreted through Systemic Functional Linguistics and Leech’s tense-aspect theory, the series represent complementary pedagogical designs: one oriented toward communicative breadth and the other toward patterned consolidation. These findings provide teachers, materials developers, and curriculum designers with an evidence-based foundation for textbook selection and adaptation in the Vietnamese primary EFL context.

Keywords: corpus-based textbook analysis; primary EFL; vocabulary range; tense-aspect; Systemic Functional Linguistics; Vietnam

Article information: Received: 21/03/2026;

Accepted: 15/05/2026;

Published: 25/06/2026

Copyright © 2026 The Author(s). Published by the International Journal of Language Studies. This is an open-access article distributed under the Creative Commons Attribution 4.0 International License (CC BY 4.0).

1. Introduction

English occupies a central place in Vietnam’s national education system, and at the primary level, where learners first encounter the language in a structured classroom setting, the textbook series adopted plays a particularly consequential role in shaping the input available to them (Richards, 2001). Two series dominate current practice. Global Success, produced by the Vietnam Education Publishing House in close alignment with the 2018 national curriculum, is used in most large urban centres; Family and Friends, published internationally by Oxford University Press, continues to be used in smaller cities and provinces, including Da Lat. Both target the A1–A² band of the Common European Framework of Reference for Languages (CEFR) and both aim to develop communicative competence, yet the two differ substantially in pedagogical orientation: Global Success organises its units around national curricular topics and everyday Vietnamese scenarios with an emphasis on systematic repetition, while Family and Friends follows an international spiral syllabus built around recurring fictional characters and a wide range of culturally varied situations. Because dialogue is the primary

vehicle through which both series model spoken interaction and how utterances are constructed, how turns are exchanged, and how tense and aspect express everyday meaning is understanding how the two series differ in their dialogue content is directly relevant to teachers, materials developers, and curriculum designers.

Despite the widespread and simultaneous use of both series, the empirical literature offers little direct evidence about how their dialogue content compares as a source of spoken input. Existing Vietnamese textbook research has concentrated on reading and listening texts (Mai et al., 2024; Nguyet & Long, 2020; To, 2018; Minh & Thuy, 2021), leaving primary-level dialogue, the most immediate model of communicative language available to young learners, comparatively unexamined. A further limitation is the tendency to treat vocabulary, sentence structure, and tense usage as isolated features rather than as interrelated dimensions of a single communicative system, and the near-total absence of direct, corpus-informed comparison between *Global Success* and *Family and Friends* specifically. Dung (2016) evaluated *Family and Friends 3* through teacher survey data, but no study has subjected the dialogues of the two series to systematic, multi-dimensional corpus analysis.

The present study addresses these gaps through a corpus-informed comparison of the dialogue content of *Global Success* and *Family and Friends* across Grades 3 to 5, guided by two research questions:

RQ1: What similarities and differences can be found in the vocabulary range, sentence structure complexity, and tense–aspect usage in the dialogues of the *Global Success* and *Family and Friends* textbook series for young learners?

RQ2: What pedagogical implications for young learners can be hypothesised from the observed linguistic similarities and differences between the two textbook series?

The study contributes to the field in two respects. Theoretically, it extends corpus-informed textbook analysis, most of which has been applied to written genres, to primary-level spoken dialogue, and demonstrates how Halliday’s Systemic Functional Linguistics (SFL) and Leech’s Tense and Aspect framework can be combined with corpus quantification to produce an analysis that is simultaneously formal and functional. Practically, the findings give teachers concrete, level-specific evidence of the vocabulary, sentence patterns, and tense-aspect categories each series emphasises, evidence that can inform decisions about supplementation and adaptation, and offer curriculum reviewers and materials researchers a replicable evaluative model that goes beyond impressionistic assessment.

Literature Review

2.1. Textbook Analysis in ELT

Textbooks are widely regarded as one of the most consequential variables in English Language Teaching, since they are frequently the dominant, and often the exclusive, source of structured target-language input in EFL contexts such as Vietnam (Cunningsworth, 1995; Richards, 2001). The field has developed complementary evaluative frameworks for assessing how well instructional materials serve learners: Cunningsworth (1995) argues that coursebooks should be treated as “good servants but poor masters”, correspondence with learners’ communicative needs being the primary criterion; Tomlinson (2011) extends this into a materials-development theory in which effective materials must achieve

impact and help learners feel at ease; and Vitta (2021) frames textbook analysis as a quality-control mechanism combining quantitative and qualitative assessment. Corpus linguistics supplies the quantitative arm of this enterprise, enabling precise documentation of lexical and grammatical patterns at a scale unavailable to manual inspection, while Tomlinson (2011) cautions that frequency data alone cannot capture the conditions under which learners process what they encounter which is an argument that underlies the present study's combination of corpus quantification with functional interpretation. Dialogue occupies a position of particular significance within this architecture, since young learners engage principally with spoken rather than written language and rely on meaningful discourse to construct understanding (Cameron, 2001); textbook dialogue that restricts learners to a narrow range of formulaic exchanges accordingly forecloses opportunities for meaning-making that a richer functional repertoire would open up.

2.2. Empirical Studies of Textbook Dialogue and Content

A foundational strand of the literature documents the gap between textbook dialogue and authentic spoken interaction. Gilmore (2004) compared coursebook service-encounter dialogues with matched authentic interactions and found the former markedly shorter, more lexically dense, and stripped of features characteristic of natural speech, such as false starts, hesitation, and back-channelling; Römer (2007) corroborated this from a corpus perspective, showing that spoken-type coursebook English structurally resembles written language more than authentic speech. Within the young-learner strand, Nordlund (2016) compared two Swedish EFL series across school years 4–6 and found considerable vocabulary variation both within and between series, with roughly a third of word types falling outside the 2,000 most frequent English words—the closest published international parallel to the present study, though restricted to vocabulary alone. In the Vietnamese context, Dang and Seals (2018) evaluated national primary textbooks from a sociolinguistic perspective but conducted no frequency analysis or grammatical profiling; Mai et al. (2024) and Nguyet and Long (2020) applied corpus tools to the lexical demands of secondary-level listening and reading texts; and Minh and Thuy (2021) examined lexical cohesion in a single reading-oriented series. None of these studies engages with primary-level dialogue as its central unit of analysis, and none integrates vocabulary, sentence structure, and tense-aspect within a single design.

Studies applying SFL to textbook dialogue provide the most directly comparable methodological precedent. Niman and Darong (2024) analysed dialogues from two elementary EFL textbooks through Halliday's three metafunctions and found relational processes and declarative mood strongly predominant, cautioning that reliance on structurally limited dialogue may constrain learners' exposure to varied interactional patterns; Djatmika et al. (2022) applied SFL to a curriculum-alignment question in Indonesian secondary materials, and Guna and Darong (2023) and Canggung Darong and Regus (2024) used SFL to reveal ideological and cognitive-accessibility dimensions of textbook grammar. These studies confirm that SFL is empirically productive for textbook dialogue analysis but applied no corpus-based frequency quantification or tense-aspect coding, gaps the present study is designed to close. Taken together, the literature reveals a persistent tendency to examine single linguistic dimensions in isolation and a near-total absence of direct, corpus-informed comparison between Global Success and Family and Friends at the primary level—the intersecting gaps that the present study addresses.

2. Literature Review

The analysis integrates two complementary frameworks. Halliday's Systemic Functional Linguistics (Halliday & Matthiessen, 2014) treats language as a social-semiotic resource organised around three simultaneous metafunctions: the ideational, realised through transitivity, concerning how language represents experience; the interpersonal, realised through mood and modality, concerning how language enacts relationships between speakers; and the textual, concerning how information is packaged into coherent messages. Because SFL treats meaning rather than form as the primary object of inquiry, it is well suited to interpreting the communicative-function distribution across the corpus, in particular the interpersonal dimension realised through the distribution of Asking, Answering, Describing, Expressing, Requesting, Greeting, Introducing, and Narrating utterances. Leech's Tense and Aspect Theory (Leech, 2014) supplies a functionally oriented account of how English verb forms encode temporal and aspectual meaning, distinguishing tense, which encodes deictic time reference, from aspect, which encodes the speaker's perspectivisation of an event as simple, progressive, or complete. The framework is well matched to the A1–A² level, where the operative distinctions, Present Simple, Present Progressive, Past Simple, Future, and Imperative, are few and functionally transparent, and it foregrounds a distinction bare frequency counts obscure: the difference between tense-as-form and tense-as-meaning, since Present Simple owes its instructional prominence to formal simplicity rather than semantic ease, while Past Simple requires learners to locate a completed event relative to the moment of utterance. Used together, the two frameworks allow formal linguistic properties, frequency distributions, structural patterns, and tense-aspect ratios, to be interpreted in relation to the communicative meanings they carry for young learners, with SFL guiding the interpretation of vocabulary and sentence structure and Leech's framework guiding the interpretation of the verbal system.

Methodology

3. Methodology

The study adopts a comparative, corpus-informed design combining quantitative corpus analysis with qualitative functional interpretation, consistent with the broader tradition of corpus-based textbook analysis (Nguyet & Long, 2020; Vitta, 2021). The design is descriptive rather than experimental: no variables are manipulated, and what is documented is naturally occurring language as it appears in published textbook dialogues. The unit of analysis is the individual utterance, defined as each speaking turn within a dialogue, whether a single word, a phrase, or a multi-clause sentence.

3.2. Data Collection

Data were drawn from the Student's Books of Family and Friends (Oxford University Press) and Global Success (Vietnam Education Publishing House, in collaboration with Macmillan Education), covering Grades 3 to 5, broadly the A1–A² CEFR band. Each Family and Friends level is contained in a single Student's Book (FF³, FF⁴, FF⁵), whereas each Global Success grade is divided into two semester volumes; for the comparison to be made at grade level, the two volumes of each Global Success grade were combined into a single sub-corpus (GS³, GS⁴, GS⁵). The corpus includes only texts in which two or more speakers exchange utterances in structured interaction—main dialogue sections, speaking-model exchanges, and communicative task dialogues—while reading passages, listening

scripts of non-conversational content, teachers' notes, and grammar boxes were excluded, so that the corpus represents the spoken-language model each series presents to learners rather than a mixture of genres. Where a dialogue appeared first in a presentation context and again in a practice context within the same unit, only the presentation version was retained.

All selected dialogues were manually transcribed and entered into a structured spreadsheet, with each utterance assigned a unique ID (encoding series, grade, unit, and utterance number), Series, Level, Unit, verbatim Utterance text, and Word_count. Two further fields, Tense_Aspect and Communicative_Function, were populated during the coding pass described in Section 3.4. The compiled corpus contained 1,143 utterances across six sub-corpora (Table 1): FF³ = 153, FF⁴ = 212, FF⁵ = 236 (FF_all = 601 utterances, 3,990 tokens); GS³ = 177, GS⁴ = 175, GS⁵ = 190 (GS_all = 542 utterances, 3,031 tokens). Token counts follow AntConc's tokenisation convention, in which contracted forms such as "it's" are counted as two tokens; the Word_count field records orthographic word counts and is used only for the utterance-length calculations reported in Section 4.2.

Table 1. Corpus Composition by Series and Grade Level

	FF3	FF4	FF5	FF_all	GS3	GS4	GS5	GS_all
Utterances	153	212	236	601	177	175	190	542
Tokens	847	1,337	1,806	3,990	761	914	1,356	3,031
Types	227	350	457	709	171	243	358	523
Units included	13	13	13	39	22	20	21	63

Note. FF_all and GS_all combine the three respective grade levels; GS figures combine both semester volumes of each grade. Source: author's coding spreadsheet.

The sub-corpora are modest by the standards of large-scale corpus research, ranging from 153 to 236 utterances per book, comparable in scale to other primary-level dialogue studies (Gilmore, 2004; Niman & Darong, 2024). Their principal strength is exhaustiveness: they contain the complete dialogue content of all units across the three grades of each series rather than a sample. Comparisons across grade levels, where sub-corpora are more closely matched in size, are treated as more reliable than aggregate cross-series comparisons throughout.

3.3. Instruments

Two instruments were used. AntConc 4.3.1 (Anthony, 2022) generated word-frequency lists and type-token statistics for each sub-corpus, performed keyword analysis using the Log-Likelihood (LL) statistic with each series as reference for the other, and extracted 2-, 3-, and 4-word N-gram clusters at a minimum frequency of three. A structured Excel spreadsheet served as the surface for manual coding of Tense_Aspect and Communicative_Function and for the cross-tabulations underlying the tables in Section 4. The pairing of an automatic frequency tool with manual functional coding is deliberate: AntConc cannot determine whether an interrogative seeks information or confirmation, or whether a Past Simple form narrates or answers a question, distinctions that depend on contextual reading and that carry most of the pedagogical interest.

3.4. Data Analysis

Vocabulary range was assessed through the Type-Token Ratio (TTR), frequency-ranked word lists with grammatical function words set aside, and Log-Likelihood keyword analysis using the $p < .01$ threshold ($LL \geq 6.63$) recommended by Rayson and Garside (2000). Because TTR is sensitive to corpus length, the Moving-Average Type-Token Ratio (MATTR) was additionally computed with a custom Python script following Covington and McFall (2010), using a 500-token window advanced one token at a time; level-by-level comparisons, where sub-corpora are more closely matched in size, and MATTR at the aggregate level, are treated as more reliable than raw aggregate TTR.

Sentence structure was examined through mean utterance length (word tokens per utterance, with standard deviation, median, minimum, and maximum), a four-way sentence-type classification (declarative, interrogative, imperative, exclamative/non-clausal) following the major independent-clause types identified by Biber et al. (2021) and derived from the existing coding, and N-gram cluster analysis in AntConc, with function-word-only clusters excluded from interpretation.

Each utterance was assigned one of six Tense_Aspect labels (Present Simple, Present Progressive, Past Simple, Future, Imperative, None) on the basis of its primary finite verb form, and one of eight Communicative_Function labels (Asking, Answering, Describing, Expressing, Requesting, Greeting, Introducing, Narrating) developed for the study and interpreted through Halliday's interpersonal metafunction; full definitions and corpus examples for both schemes are provided in Appendix A. Coding was carried out in two passes, and borderline cases were resolved against a written set of decision rules and reviewed with the study's academic supervisors. All coding was performed by a single coder, the principal methodological limitation of the study; to address this partially, a stratified 15% sample of utterances (172 of 1,143) was independently coded by a second rater following the same decision rules, yielding substantial agreement (Landis & Koch, 1977) of $\kappa = .82$ for tense-aspect and $\kappa = .79$ for communicative-function labels. Chi-square tests and Cramér's V were used to test and quantify the association between series and categorical distributions (sentence type, tense-aspect); because utterances are nested within dialogues, units, and volumes rather than fully independent, these tests are interpreted cautiously as descriptive evidence of distributional association, with Cramér's V, which is not sensitive to sample-size inflation in the way p-values are, weighted more heavily than significance levels alone.

Ethics Statement

The study involved no human participants; all data were drawn from published, commercially available textbook materials, and consultation with the supervising faculty confirmed that no formal ethics committee clearance was required for research of this kind. Transparency was maintained by stating explicitly which nine Student's Books supplied the data and by listing their full bibliographic details in the References. Quoted utterances are kept to the minimum length necessary to illustrate a coding category or structural pattern, consistent with the fair-use conventions governing non-commercial academic research.

4. Results

4.1. Vocabulary Range

Table 2 reports types, tokens, raw TTR, and MATTR for every sub-corpus. Both series grow steadily from Grade 3 to Grade 5 (FF from 227 to 457 types; GS from 171 to 358), indicating that neither series front-loads its vocabulary. At the aggregate level, raw TTR is misleading: FF_all (17.77%) and GS_all (17.26%) appear almost identical, an artefact of FF_all being the larger corpus. MATTR removes this distortion and reveals a real difference, FF_all reaching 39.89% against GS_all's 37.08%.

Table 2. Types, Tokens, TTR, and MATTR by Series and Level

	FF3	FF4	FF5	FF_all	GS3	GS4	GS5	GS_all
Tokens	847	1,337	1,806	3,990	761	914	1,356	3,031
Types	227	350	457	709	171	243	358	523
TTR (%)	26.80	26.18	25.30	17.77	22.47	26.59	26.40	17.26
MATTR (%)	35.78	39.81	40.37	39.89	31.43	34.77	41.30	37.08

Note. Raw TTR = Types/Tokens × 100. MATTR computed with a 500-token sliding window following Covington and McFall (2010); values for the smallest sub-corpora (FF3, GS3) are indicative given window size relative to corpus length.

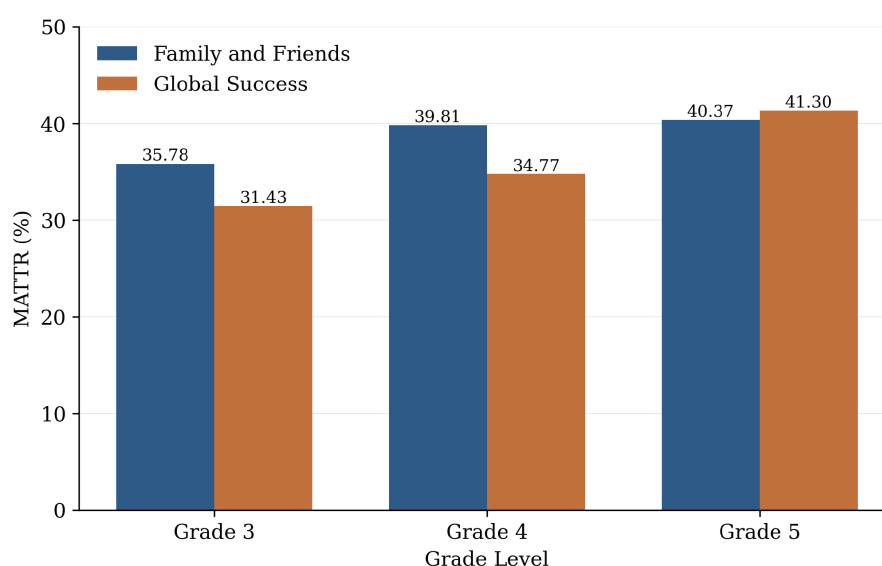


Figure 1. Moving-Average Type-Token Ratio (MATTR) by Grade Level

The level-by-level MATTR trajectory (Figure 1) is more informative than the aggregate figures. At Grade 3 the gap between series is 4.35 percentage points (FF3 35.78%, GS3 31.43%); at Grade 4 it widens to 5.04 points (FF4 39.81%, GS4 34.77%); at Grade 5 it closes and slightly reverses, GS5 reaching 41.30% against FF5's 40.37%. FF thus offers a wider vocabulary range from the outset and holds it fairly steady, while GS starts narrower and widens, converging with FF only by the end of primary school. The pattern is consistent with two defensible but different accounts of vocabulary acquisition: Cameron's (2001) view that young learners benefit from meeting words in many contexts early on, which favours FF's breadth, and Nation's (2001) view that high-frequency vocabulary requires repeated, spaced encounters, which is closer to GS's narrower, more recycled base. A

comparable conservatism appears in Nguyet and Long's (2020) finding that a newer national secondary series delivered essentially the same lexical coverage as its predecessor despite being larger, and in To's (2018) finding that textbook complexity in a Vietnamese series does not always peak where level structure implies it should—a bend that recurs here, where GS5 overtakes FF5 rather than continuing to trail.

Table 3 lists the twenty most frequent content words in each corpus (grammatical function words excluded). The single shared top item is “like” (61 occurrences in FF, 41 in GS), though it performs several distinct functions (lexical verb, comparison marker, formulaic element) that raw frequency cannot disambiguate. Beyond this, the two lists diverge sharply: eight of FF's top twenty items are character or family-role names (Billy, Mom, Dad, Max, Leo, Grandma, Holly, Amy), reflecting its design around a recurring cast, while its remaining lexical items (look, doing, going, play, swimming, favorite, because, bag, book) describe activity, possessions, and preference. GS's list mixes Vietnamese and international character names (Mai, Minh, Linh, Nam, Ben, Lucy) with curricular topic words (school, time, day, summer, sports, live, old); its most telling item is “does” (3rd rank), reflecting the recurrence of third-person question frames such as “What does he do?” that run through GS as a structural spine.

Table 3. Top Content Words in FF and GS Dialogues

Rank	FF_all	Freq.	GS_all	Freq.
1	like	61	like	41
2	look	37	school	17
3	billy*	26	does	16
4	mom*	21	mai*	15
5	dad*	15	play	14
6	doing	15	last	13
7	max*	15	time	13
8	leo*	12	want	13
9	bag	11	ben*	12
10	favorite	11	day	12

*Note. Items marked * are proper nouns / character names. Top 10 of 20 shown; full list available from the author. Source: AntConc 4.3.1.*

A keyword analysis was then run in AntConc using the Log-Likelihood statistic, with each series as target and the other as reference, retaining only items above the $p < .01$ threshold ($LL \geq 6.63$; Rayson & Garside, 2000). Table 4 gives the strongest linguistically meaningful keywords for each series (proper nouns omitted for brevity). FF's strongest keywords are discourse and interactional markers,

oh (LL = 31.14), no, sorry, but, and favorite, the vocabulary of interaction: marking surprise, softening refusal, apologising, drawing contrast. This is precisely the kind of vocabulary Gilmore (2004) found routinely thinned out of textbook English; its retention in FF, and its near-total absence from GS (oh occurs once in the entire GS corpus), marks a clear editorial divergence. GS's strongest keywords are interrogatives and curricular nouns, what (LL = 32.53), where, do, sports, maths, day, summer, restating from another angle that the question-answer exchange is GS's structural core.

Table 4. Selected Keywords by Series ($p < .01$, $LL \geq 6.63$), Proper Nouns Omitted

FF keyword	Freq. FF	Freq. GS	LL	GS keyword	Freq. GS	Freq. FF	LL
oh	34	1	31.14	what	98	50	32.53
we	72	18	21.80	where	34	11	19.62
and	72	21	17.52	do	76	45	19.17
no	58	15	16.76	sports	9	0	15.14
but	19	1	15.25	day	12	1	14.26
look	37	7	15.11	maths	8	0	13.45
than	13	0	14.71	hi	24	8	13.43
sorry	13	0	14.71	summer	11	1	12.75
favorite	11	0	12.45	last	13	2	12.35
now	11	0	12.45	how	24	9	11.88

Note. LL = Log-Likelihood; threshold $p < .01$ (Rayson & Garside, 2000). Proper nouns (e.g., Billy, Mary, Lucy, Mai, Ben) accounted for a further seven items in each list and are omitted here as they largely restate cast composition rather than lexical preference; full lists available from the author. Source: AntConc 4.3.1.

Together, the frequency and keyword evidence converge on the same picture: FF's distinctiveness lies in interactional and evaluative language, while GS's lies in the apparatus of the topic-based question and curriculum-anchored vocabulary.

4.2. Sentence Structure

Table 5 reports utterance-length statistics. Means rise across grades in both series (FF: 5.31–7.89 words; GS: 3.95–6.89 words), but FF utterances are longer than GS utterances at every level, most widely at Grade 3 (5.31 vs. 3.95). Standard deviation reveals a further contrast: GS3 has the lowest SD in the dataset (1.85), indicating uniformly short, template-built utterances, while FF5 has the highest (4.33), reflecting a mix of short turns and long multi-clause constructions (up to 33 words) combining description, narration, and explanation. Both series arrive at longer utterances by Grade 5, but FF does so through clause combination and GS through repetition of short patterns, a divergence consistent

with the nominal-group simplicity that Niman and Darong (2024) and Canggung Darong and Regus (2024) associate with beginner-level materials in general, but which the two series here realise to markedly different degrees.

Table 5. Utterance Length Statistics by Series and Level (Words)

	FF3	FF4	FF5	FF_all	GS3	GS4	GS5	GS_all
Mean	5.31	6.03	7.89	6.58	3.95	4.96	6.89	5.31
SD	3.40	3.42	4.33	3.95	1.85	2.30	3.60	2.98
Median	4.0	5.5	7.0	6.0	4.0	4.0	6.0	4.5
Min–Max	1–23	1–19	1–33	1–33	1–11	1–13	1–28	1–28

Note. SD = standard deviation. Utterance lengths counted in orthographic words (differ slightly from AntConc token counts in Table 2, which split contractions).

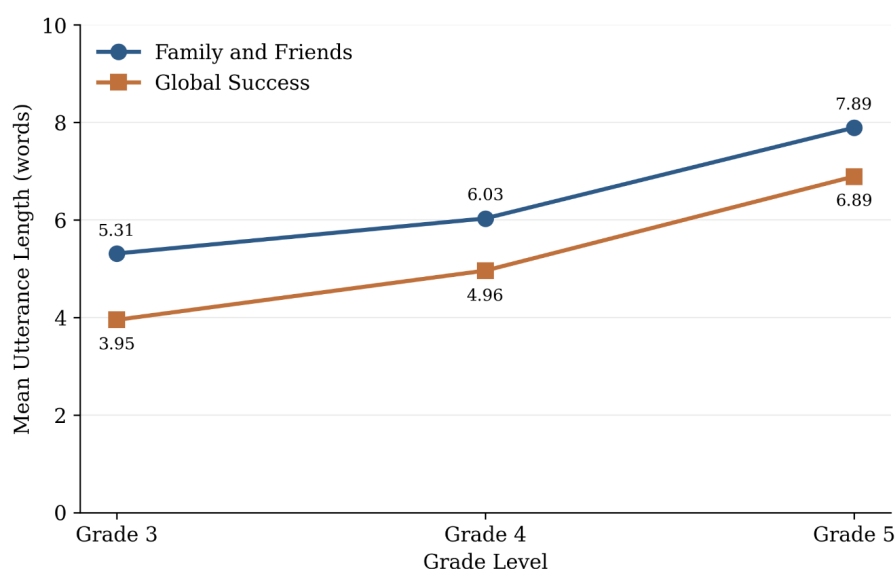


Figure 2. Mean Utterance Length by Grade Level

A four-way sentence-type classification (declarative, interrogative, imperative, exclamative/non-clausal) showed a significant association between series and sentence type, $\chi^2(3) = 39.61$, $p < .001$, Cramér's $V = .186$ (small effect). As Table 6 shows, GS places 39.7% of its utterances into questions against FF's 24.1%, a gap widest at Grade 4 (GS 44.0% vs. FF 22.2%); FF compensates with more declaratives (58.9% vs. 51.7%), more exclamative/non-clausal utterances (12.0% vs. 6.1%), and slightly more imperatives (5.0% vs. 2.6%). Notably, FF contains nine Narrating utterances, all in FF5, while GS contains none. This pattern both confirms and complicates the SFL literature: Niman and Darong (2024) and Canggung Darong and Regus (2024) found declarative mood predominant in elementary EFL dialogue generally, which holds here for both series, but their characterisation of interrogative and imperative moods as secondary does not describe GS, whose interrogative share (39.7%) nearly rivals its declaratives—evidence that a curriculum-tied national series can sit at the demanding extreme of a range that prior single-series studies could not reveal.

Table 6. Distribution of Sentence Types by Series (n and %)

Sentence type	FF (n)	FF (%)	GS (n)	GS (%)
Declarative	354	58.9	280	51.7
Interrogative	145	24.1	215	39.7
Imperative	30	5.0	14	2.6
Exclamative / non-clausal	72	12.0	33	6.1
Total	601	100	542	100

Note. $\chi^2(3) = 39.61$, $p < .001$, Cramér's $V = .186$.

N-gram analysis reinforces the same pattern at the formulaic level. “Do you” heads both two-gram lists but occurs nearly twice as often in GS (45) as in FF (23); GS produces 78 recurrent three-word clusters against FF's 54, and 33 four-word clusters against FF's 16, with the sharpest contrast at Grade 4 (GS4 = 20 clusters vs. FF4 = 3). FF's clusters (do you like, would you like, could I have, nice to meet you) span several speech acts—preference, request, social formula—while GS's clusters (do you have, what do you, where are you) stay close to information-seeking. The contrast is consistent with Cunningsworth's (1995) framing of controlled practice and communicative variety as a trade-off rather than a matter of better and worse: GS buys reliability at the cost of range, FF buys range at the cost of reinforcement.

4.3. Tense and Aspect

Six tense-aspect categories were coded for every utterance. Table 7 reports the overall distribution; a chi-square test confirmed a significant association between series and tense-aspect category, $\chi^2(5) = 33.93$, $p < .001$, Cramér's $V = .172$. Present Simple dominates both series, as expected at the A1–A2 level, but GS relies on it far more heavily (72.9% vs. 57.1%); the room FF gives to other forms goes mainly to Past Simple (13.5% vs. 8.9%), Present Progressive (9.0% vs. 6.5%), and Imperative (5.0% vs. 2.6%). Read through Leech's (2014) account of tense and aspect as communicative choices rather than formal labels, GS's concentration on Present Simple represents a concentration on a single temporal stance, the habitual and the general, while FF exposes learners to a larger portion of the choices the English verb system offers.

Table 7. Overall Tense and Aspect Distribution by Series

Tense / aspect	FF (n)	FF (%)	GS (n)	GS (%)
Present Simple	343	57.1	395	72.9
Present Progressive	54	9.0	35	6.5
Past Simple	81	13.5	48	8.9
Future	21	3.5	17	3.1

Tense / aspect	FF (n)	FF (%)	GS (n)	GS (%)
Imperative	30	5.0	14	2.6
None	72	12.0	33	6.1
Total	601	100	542	100

Note. $\chi^2(5) = 33.93$, $p < .001$, Cramér's $V = .172$. Imperative and None rows share counts with the corresponding cells in Table 6, as the same utterances are classified under both schemes; the two tests are complementary descriptions of the same data rather than fully independent.

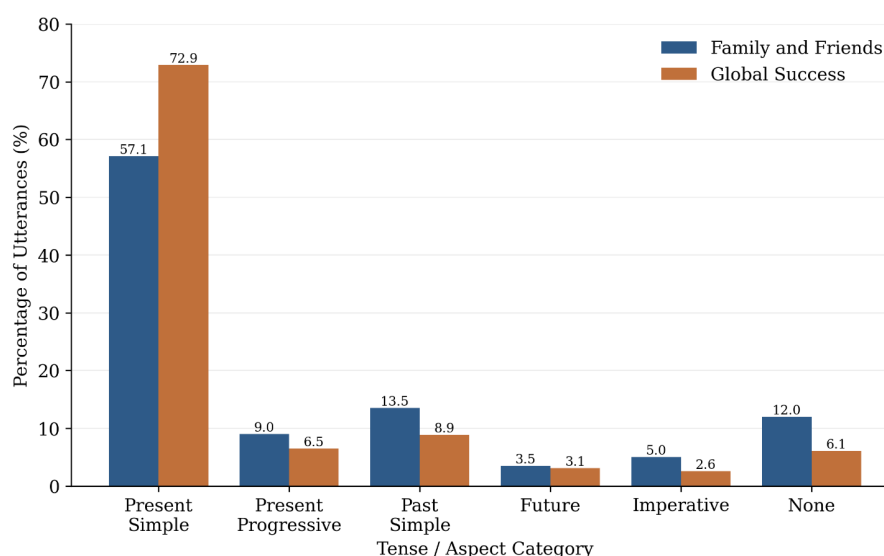


Figure 3. Distribution of Tense and Aspect Categories by Series

Breaking the distribution down by grade (Table 8) shows GS4 pushing Present Simple to 84.0%, the highest concentration of any form in any sub-corpus in the study, consistent with a consolidation year of intensive single-form practice, while FF4 (58.0% Present Simple, 14.2% Present Progressive) has already begun diversifying. Past Simple is absent from both series at Grade 3 but reaches 30.5% in FF5 against only 19.5% in GS5, the largest single Grade 5 gap in the study.

Table 8. Tense and Aspect Distribution by Series and Level (%)

Tense / aspect	FF3	FF4	FF5	FF_all	GS3	GS4	GS5	GS_all
Present Simple	71.2	58.0	47.0	57.1	79.7	84.0	56.3	72.9
Present Progressive	12.4	14.2	2.1	9.0	7.3	3.4	8.4	6.5
Past Simple	0.0	4.2	30.5	13.5	0.0	6.3	19.5	8.9
Future	0.0	0.0	8.9	3.5	0.0	0.0	8.9	3.1
Imperative	2.0	8.5	3.8	5.0	5.1	1.7	1.1	2.6

Tense / aspect	FF3	FF4	FF5	FF_all	GS3	GS4	GS5	GS_all
None	14.4	15.1	7.6	12.0	7.9	4.6	5.8	6.1

Note. Column percentages within each sub-corpus.

Because a frequency count alone cannot show what a tense form is used to do, Past Simple, which showed the largest Grade 5 gap, was cross-tabulated against communicative function (Table 9). At Grade 5, 97.3% of all GS5 Past Simple utterances fall into just two functions, Asking and Answering, with none used to describe or narrate; the aggregate GS figure (95.8%) confirms this is a stable design feature rather than a local accident. FF5 spreads Past Simple across five functions—Answering (40.3%), Asking (36.1%), Narrating (9.7%), Describing (8.3%), Expressing (5.6%)—with Narrating and Describing requiring learners to string several past-tense clauses into connected recount, a competence FF5 dialogues model and GS5 dialogues do not. Figure 4 shows that this restriction is one expression of a wider pattern: GS channels the great majority of its utterances, across all tenses, into Asking and Answering, while FF spreads its utterances more evenly and is the only series to use the Narrating function at all.

Table 9. Past Simple by Communicative Function, FF5 vs. GS5 (n and %)

	Asking	Answering	Describing	Narrating	Other*	Total
FF5 Past Simple	26 (36.1%)	29 (40.3%)	6 (8.3%)	7 (9.7%)	4 (5.6%)	72
GS5 Past Simple	17 (45.9%)	19 (51.4%)	0 (0.0%)	0 (0.0%)	1 (2.7%)	37

*Note. *Other includes Requesting, Greeting, and Introducing combined. Aggregate figures across all grades ($n(\text{FF_all}) = 81$, $n(\text{GS_all}) = 48$) show the same pattern: 95.8% of GS Past Simple utterances fall in Asking/Answering, versus 67.9% for FF, with Narrating absent from GS entirely.*

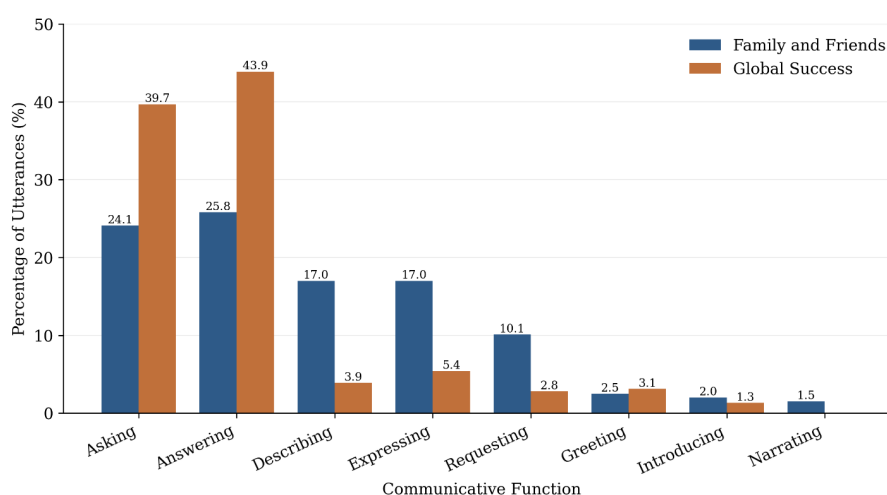


Figure 4. Distribution of Communicative Functions by Series

This functional restriction matters pedagogically: a learner completing GS5 will have produced many correct past-tense answers to direct questions but very few past-tense narratives, while a learner completing FF5 will have met both. The finding illustrates the value of cross-tabulating tense with function rather than counting tense alone, the methodological move Vitta (2021) identifies as what shifts textbook analysis from inventory toward function, and echoes Khojasteh and Mukundan's (2025) observation that multi-dimensional analysis tends to yield findings of more practical use than single-feature counts.

5. Discussion

The three dimensions examined tell a consistent and internally coherent story. Global Success's reliance on Present Simple sustains its question-answer routine; that routine generates its topic-anchored vocabulary and repetitive cluster profile; and the same routine confines its past tense to Asking and Answering. Family and Friends's greater tense variety, its narrative function, and its higher proportion of declaratives are the same communicative orientation surfacing in three separate measures. Neither series is simply richer or poorer than the other; each is a coherent design built on a different account of what early English instruction should prioritise, one favouring communicative breadth, the other patterned consolidation and curricular fit.

Because this study analysed textbook input rather than learner uptake, the implications offered here in response to RQ2 are input-based hypotheses rather than empirically verified claims about acquisition, and should be read in that spirit. For vocabulary, FF's wider MATTR at Grades 3 and 4 favours the broad receptive base that Cameron (2001) and Nation (2001) both regard as foundational to later growth, while GS's narrower, more repetitive vocabulary supports the productive mastery of a core set of high-frequency, curriculum-relevant items; the two serve complementary rather than competing vocabulary goals. For sentence structure, GS's intensive drilling of a few frames suits curriculum-aligned assessment, where learners must produce reliable instances of target patterns, while FF's structural variety, including narrative sequences absent from GS, suits extended communicative use, where a learner must choose among structures rather than reproduce one, consistent with Tomlinson's (2011) call for materials that are both rich and manageable. The tense findings carry the most concrete implication: a learner completing GS5 will have had almost no exposure to past-tense narration, practised only as an answer within an exchange, so a teacher relying on GS would do well to supplement narrative and descriptive past-tense work deliberately; conversely, a teacher relying on FF, whose interrogative practice is comparatively thin, may need to add the controlled question-answer work that GS supplies in abundance.

None of this supports ranking one series above the other. It points instead toward complementarity and informed combination in the Vietnamese primary context, with Family and Friends contributing lexical and structural range and the modelling of extended discourse, and Global Success contributing pattern security, curricular fit, and cultural familiarity, consistent with Cunningsworth's (1995) long-standing argument that no single coursebook meets every need of a teaching context and that teacher mediation is part of how any coursebook is properly used. The convergence of the two series by Grade 5 across MATTR, utterance length, and tense variety suggests that the divide between them is one of developmental route rather than ultimate destination.

Several limitations qualify these conclusions. Tense-aspect and communicative-function coding was performed primarily by a single coder; although a stratified 15% sample was independently double-coded with substantial agreement ($\kappa = .82$ and $\kappa = .79$), a full double-coding exercise remains the appropriate standard for future replications. The corpus, while exhaustive for the dialogue content of these nine books, is modest in absolute size (1,143 utterances), so level-by-level figures, especially at Grade 3, should be read as indicative; the chi-square tests, moreover, treat utterances as independent observations despite their nesting within dialogues, units, and volumes, which may inflate significance, though the small-to-moderate Cramér's V effect sizes reported (.172 to .375 across the three tests) are less sensitive to this clustering than the p -values. The study also examined dialogue alone, leaving the reading passages, listening scripts, and exercises that surround it outside its scope, and covered two specific series at one educational level in one country, so the patterns identified cannot be assumed to generalise without separate investigation. Most importantly, the study analysed input, not uptake: it can show what language each series places in front of learners but not what learners acquire from either book, so the pedagogical implications above remain hypotheses grounded in the input data and the wider literature rather than direct evidence of learning outcomes.

These limitations point to clear directions for further research: a replication incorporating full double-coding with reported inter-rater reliability; extension of the corpus to additional 2018-curriculum series or into lower-secondary grades, to test whether the contrast identified here is stable beyond these two books and three levels; application of the same framework to the reading and listening components of both series, to give a fuller picture of input across modes; and, most valuably, a classroom-based or longitudinal study tracking learner outcomes under the two series, or under the combined use recommended here, to test directly the prediction that learners completing *Global Success* will be less prepared to produce past-tense narrative than learners completing *Family and Friends*.

6. Conclusion

This study set out to supply a corpus-informed comparison of the dialogue content of two textbook series in simultaneous use across Vietnamese primary classrooms, examining vocabulary range, sentence structure complexity, and tense-aspect usage across Grades 3 to 5. The similarities between the series proved real but largely developmental: both expand vocabulary and utterance length steadily across the grades and lean heavily on Present Simple before introducing Past Simple and Future at Grade 5. The differences proved deeper and more consequential: once corpus-size effects were controlled through MATTR, *Family and Friends* showed a clearly wider lexical range at Grades 3 and 4; its utterances were longer and more variable in structure; and it distributed tense-aspect forms more widely, using Past Simple for narration and description in ways *Global Success*, which confined the form almost entirely to question-and-answer exchanges, did not. These differences did not sit independently of one another but formed two internally coherent communicative designs, one favouring breadth, the other favouring patterned consolidation and curricular alignment.

Recognising this complementarity, rather than treating one series as simply superior, is what allows the two to be used well together in practice, with teacher mediation directed at whichever dimension each book leaves comparatively thin. The contribution of this article is a modest one, bounded by the limitations discussed above, but it brings corpus evidence to bear on a comparison that had previously rested on impression, and offers teachers, materials developers, and curriculum designers a clearer,

replicable account of the linguistic input the two series provide than was previously available. Extending the framework developed here to other materials, other levels, and eventually to learner outcomes would allow the empirical picture drawn in this study to inform a larger and more directly useful body of evidence.

Acknowledgements

The author thanks Dr. Trần Tín Nghị and Dr. Huỳnh Quang Minh for their supervision and guidance during the master's thesis research at Dalat University on which this article is based.

Funding

This research received no external funding.

Author Contributions

Dương Minh Hoàng: Conceptualization, Methodology, Data Curation, Formal Analysis, Writing—Original Draft, Writing—Review & Editing.

Conflicts of Interest

The author declares no conflict of interest.

Ethics Statement

This study did not involve human or animal participants. All data were drawn from published, commercially available textbook materials (see References), and consultation with the author's supervising faculty at Dalat University confirmed that no formal research ethics committee approval was required for a study of this kind.

Declaration of Generative AI and AI-Assisted Technologies

The author used Claude (Anthropic) to assist in reformatting and condensing this article from the author's master's thesis into the journal's manuscript template. The tool was not used to generate original data, analysis, or findings; all statistical results reported are drawn directly from the author's thesis. All AI-assisted text was reviewed, verified, and edited by the author, who takes full responsibility for the content of this article.

References

- Anjum, R. Y., & Javed, M. (2019). A Corpus-Based Halliday's Transitivity Analysis of 'To the Lighthouse'. *Linguistics and Literature Review*, 05(02), 139–162. <https://doi.org/10.32350/llr.52.05>
- Biber, D., Johansson, S., Leech, G. N., Conrad, S., & Finegan, E. (2021). *Grammar of Spoken and Written English*. John Benjamins Publishing Company. <https://doi.org/10.1075/z.232>
- Cameron, L. (2001). *Teaching Languages to Young Learners*. Cambridge University Press. <https://doi.org/10.1017/cbo9780511733109>
- Canggung Darong, H., & Regus, M. (2024). Unpacking linguistic features in EFL textbooks using systemic functional linguistics: Transitivity, Mood, and nominal group structure analysis. *East European Journal of Psycholinguistics*, 11(2), 8–32. <https://doi.org/10.29038/eejpl.2024.11.2.dar>
- Covington, M. A., & McFall, J. D. (2010). Cutting the Gordian Knot: The Moving-Average Type-Token Ratio

- (MATTR). *Journal of Quantitative Linguistics*, 17(2), 94–100. <https://doi.org/10.1080/09296171003643098>
- Cunningsworth, A. (1995). *Choosing Your Coursebook*. <https://openalex.org/W23944679>
- Dang, T. C. T., & Seals, C. (2016). An Evaluation of Primary English Textbooks in Vietnam: A Sociolinguistic Perspective. *TESOL Journal*, 9(1), 93–113. <https://doi.org/10.1002/tesj.309>
- Siregar, R. A., Sukyadi, D., & Yusuf, F. N. (2024). A critical content analysis of writing materials covered in Indonesian high school English textbooks. *Studies in English Language and Education*, 11(1), 205–227. <https://doi.org/10.24815/siele.v11i1.30169>
- Klein, M. W., Decker, S. H., & Winkle, B. V. (1997). Life in the Gang: Family, Friends, and Violence. *Social Forces*, 75(4), 1476. <https://doi.org/10.2307/2580685>
- Gilmore, A. (2004). A comparison of textbook and authentic interactions. *ELT Journal*, 58(4), 363–374. <https://doi.org/10.1093/elt/58.4.363>
- Guna, S., & Darong, H. C. (2023). The Representation of Social Actors in EFL Textbooks: Systemic Functional Linguistics Perspective. *Indonesian Journal of EFL and Linguistics*, 221–232. <https://doi.org/10.21462/ijefl.v8i2.649>
- Halliday, M., & Matthiessen, C. M. (2013). *Halliday's Introduction to Functional Grammar*. Routledge. <https://doi.org/10.4324/9780203431269>
- Thi, L. N. B., & Hằng, N. T. M. (2026). Đánh giá tính chính xác và lưu loát của các hoạt động sản sinh ngôn ngữ trong sách giáo khoa “tiếng Anh 10 – global success”. *The University of Danang - Journal of Science and Technology*, 24–29. [https://doi.org/10.31130/ud-jst.2026.24\(5c\).293](https://doi.org/10.31130/ud-jst.2026.24(5c).293)
- Thi, L. N. B., & Hằng, N. T. M. (2026). Đánh giá tính chính xác và lưu loát của các hoạt động sản sinh ngôn ngữ trong sách giáo khoa “tiếng Anh 10 – global success”. *The University of Danang - Journal of Science and Technology*, 24–29. [https://doi.org/10.31130/ud-jst.2026.24\(5c\).293](https://doi.org/10.31130/ud-jst.2026.24(5c).293)
- Thi, L. N. B., & Hằng, N. T. M. (2026). Đánh giá tính chính xác và lưu loát của các hoạt động sản sinh ngôn ngữ trong sách giáo khoa “tiếng Anh 10 – global success”. *The University of Danang - Journal of Science and Technology*, 24–29. [https://doi.org/10.31130/ud-jst.2026.24\(5c\).293](https://doi.org/10.31130/ud-jst.2026.24(5c).293)
- Thi, L. N. B., & Hằng, N. T. M. (2026). Đánh giá tính chính xác và lưu loát của các hoạt động sản sinh ngôn ngữ trong sách giáo khoa “tiếng Anh 10 – global success”. *The University of Danang - Journal of Science and Technology*, 24–29. [https://doi.org/10.31130/ud-jst.2026.24\(5c\).293](https://doi.org/10.31130/ud-jst.2026.24(5c).293)
- Doan, K. T., Huynh, B. G., Hoang, Dung T., Pham, T. D., Pham, N. H., Nguyen, Q. T. M., Vo, B. Q., & Hoang, S. N. (2024a). Vintern-1B: An Efficient Multimodal Large Language Model for Vietnamese. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2408.12480>
- Doan, K. T., Huynh, B. G., Hoang, Dung T., Pham, T. D., Pham, N. H., Nguyen, Q. T. M., Vo, B. Q., & Hoang, S. N. (2024b). Vintern-1B: An Efficient Multimodal Large Language Model for Vietnamese. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2408.12480>
- Khojasteh, L., & Mukundan, J. (2025). A Systematic Review of Corpus-Based Methodologies in Textbook Analyses and Evaluation. *Language Teaching Research Quarterly*, 47, 175–195. <https://doi.org/10.32038/ltrq.2025.47.10>
- Landis, J. R., & Koch, G. G. (1977). The Measurement of Observer Agreement for Categorical Data. *Biometrics*, 33(1), 159. <https://doi.org/10.2307/2529310>
- Leech, G. N. (2014). *Meaning and the English Verb*. Routledge. <https://doi.org/10.4324/9781315835464>
- Mai, N. H., Lien, N. N., & Trang, N. H. (2024). Lexical Demands and Features of English Textbooks for Vietnamese 10th Graders: An In-depth Comparison of Listening Sections. *Sage Open*, 14(4). <https://doi.org/10.1177/21582440241299247>
- Minh, N. H., & Thuy, D. T. (2021). A PRELIMINARY STUDY ON LEXICAL COHESIVE DEVICES IN READING TEXTS IN THE TEXTBOOK TIẾNG ANH 10 PUBLISHED BY VIETNAM EDUCATION PUBLISHING HOUSE. *VNU Journal of Foreign Studies*, 37(1). <https://doi.org/10.25073/2525-2445/vnufs.4665>

- Nation, I. S. P. (2001). *Learning Vocabulary in Another Language*. Cambridge University Press. <https://doi.org/10.1017/cbo9781139524759>
- Nguyet, H. T. T., & Long, N. V. (2020). APPLYING CORPUS LINGUISTICS TO ENGLISH TEXTBOOK EVALUATION: A CASE IN VIET NAM. *VNU Journal of Foreign Studies*, 36(6). <https://doi.org/10.25073/2525-2445/vnufs.4632>
- University of Economics - The University of Danang, Viet Nam, Tin, P. Q., Y, T. T. M., FPT Greenwich DaNang Centre, FPT University, Viet Nam, Hieu, T. T., Dong A University, Viet Nam, Thu Trang, H. T., & Huynh Trang Education Company Limited, Viet Nam (2023). Employee Retention in Luxury Accommodation Business in Da Nang, Viet Nam. *RA JOURNAL OF APPLIED RESEARCH*, 09(05). <https://doi.org/10.47191/rajar/v9i5.05>
- Niman, E. M., & Darong, H. C. (2024). An SFL Analysis of Dialogues in Elementary EFL Textbooks. *Celt: A Journal of Culture, English Language Teaching & Literature*, 24(1), 52–68. <https://doi.org/10.24167/celt.v24i1.10809>
- Nordlund, M. (2016). EFL textbooks for young learners: a comparative analysis of vocabulary. *Education Inquiry*, 7(1), 27764. <https://doi.org/10.3402/edui.v7.27764>
- Rayson, P., & Garside, R. (2000). Comparing corpora using frequency profiling. Proceedings of the workshop on Comparing corpora -, 9, 1–6. <https://doi.org/10.3115/1117729.1117730>
- Kasneci, E., Seßler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Richards, J. C., & Rodgers, T. S. (2001). *Approaches and Methods in Language Teaching*. Cambridge University Press. <https://doi.org/10.1017/cbo9780511667305>
- Römer, U., O'Donnell, M. B., & Ellis, N. C. (2014). Second Language Learner Knowledge of Verb–Argument Constructions: Effects of Language Transfer and Typology. *Modern Language Journal*, 98(4), 952–975. <https://doi.org/10.1111/modl.12149>
- Benson, P. (2007). Autonomy in language teaching and learning. *Language Teaching*, 40(1), 21–40. <https://doi.org/10.1017/s0261444806003958>
- To, V. (2018). Linguistic Complexity Analysis: A Case Study of Commonly-Used Textbooks in Vietnam. *Sage Open*, 8(3). <https://doi.org/10.1177/2158244018787586>
- Pham, G. (2011). Dual Language Development among Vietnamese-English Bilingual Children: Modeling Trajectories and Cross-Linguistic Associations within a Dynamic Systems Framework. <https://openalex.org/W2167746893>
- Pham, G. (2011). Dual Language Development among Vietnamese-English Bilingual Children: Modeling Trajectories and Cross-Linguistic Associations within a Dynamic Systems Framework. <https://openalex.org/W2167746893>
- Pham, G. (2011). Dual Language Development among Vietnamese-English Bilingual Children: Modeling Trajectories and Cross-Linguistic Associations within a Dynamic Systems Framework. <https://openalex.org/W2167746893>
- Vitta, J. P. (2021). The Functions and Features of ELT Textbooks and Textbook Analysis: A Concise Review. *RELC Journal*, 54(3), 856–863. <https://doi.org/10.1177/00336882211035826>