

Identifying the Determinants of Successful and Unsuccessful STANAG 6001 Level 2 Writing in English

Robert HELÁN^{1*}

¹Testing Department, Language Centre, University of Defence, Kounicova 65, 662 10 Brno, Czech Republic

Correspondence: *robert.helan@unob.cz

Abstract

English serves as the primary operational language of NATO, with proficiency assessed through the STANAG 6001 examination. This study analyses 240 authentic texts produced during the STANAG 6001 Level 2 writing examination at the University of Defence using corpus-based error analysis. Results indicate that failure to achieve Level 2 is mainly caused by systematic morphosyntactic errors and omission of obligatory language elements. Level 2 texts demonstrate greater grammatical accuracy and complexity, with implications for targeted English language instruction in military higher education.

KEY WORDS: *corpus linguistics, error analysis and frequency, military higher education, proficiency level variations, STANAG 6001, level 2 writing*

Citation: Helán, R. Identifying the Determinants of Successful and Unsuccessful STANAG 6001 Level 2 Writing in English. In Proceedings of the Challenges to National Defence in Contemporary Geopolitical Situation, Brno, Czech Republic, 7-10 September 2026. ISSN 2538-8959, <https://doi.org/10.47459/cndcgs.2026.122>

1. Introduction

English plays a central role in the defence sector as the principal means of communication in multinational and coalition-based operations. Within the North Atlantic Treaty Organization (NATO) framework, effective communication in English is essential for planning, coordination, and execution of military activities, as well as for training, command, and logistical support [1]. The ability to understand and respond accurately in a shared operational language is therefore a key prerequisite for successful performance in international military environments [2,3].

Language proficiency in the defence sector is formally assessed through the standardized STANAG 6001 (standardization agreement 6001) examination, which evaluates four core skills: listening, reading, speaking, and writing [4]. This system provides a common benchmark for measuring language competence and ensures interoperability among allied forces. Achieving the required proficiency levels is essential for academic progression, professional deployment, and participation in multinational structures [2].

At the University of Defence in Brno, language training is systematically implemented through courses aimed at achieving standardized language profiles (SLPs) at levels SLP 1111, 2222, and 3333 [1]. Attaining Level 2 (SLP 2222) represents a mandatory requirement for successful completion of studies and for further professional engagement in international environments, including courses, missions, and allied assignments [1].

Despite these structured training conditions, a significant proportion of students repeatedly fail to meet the required standards, particularly in the writing component of the STANAG 6001 examination. Their written performance often does not correspond to the prescribed descriptors, suggesting the presence of specific linguistic deficiencies that hinder attainment of the target level. Identifying these deficiencies is essential for improving the effectiveness of language instruction within military higher education.

The present study addresses this issue by analyzing the linguistic determinants of failure in the writing component of the STANAG 6001 examination. Building on previous research [5], it focuses on identifying the types of errors that most frequently prevent students at the University of Defence from achieving Level 2. The analysis is based on a corpus of 240 authentic texts at proficiency levels 1, 1+, and 2 and combines corpus-linguistic [6] and error-analysis methods [7].

The findings may contribute to the refinement of instructional strategies, the increased effectiveness of preparatory courses and individual self-study, and a more targeted approach to developing writing competence required for STANAG 6001 Level 2 in the context of military higher education.

2. Theoretical Background

Within the field of second language acquisition (SLA) and error analysis (EA), considerable attention has been devoted to identifying and categorizing deviations from target language norms in learner production [7, 8, 9]. Such differentiation enables a more precise assessment of learners' language proficiency and provides a structured basis for evaluating their performance.

A fundamental distinction, originally formulated by Corder [10, 11], is drawn between *errors* and *mistakes*. Errors are systematic deviations reflecting incomplete acquisition of linguistic rules and thus indicate gaps in the learner's underlying competence. In contrast, mistakes are unsystematic performance lapses caused by factors such as fatigue, stress, or momentary inattention, rather than by a lack of knowledge. While mistakes are generally considered less relevant for pedagogical intervention, errors are viewed as a valuable source of information for both language assessment and instruction, as they reveal persistent deficiencies in learners' interlanguage systems. In line with this theoretical perspective, the present study focuses primarily on systematic errors, as these offer the most relevant information about the level of language development and the stability of learners' writing skills. Such deviations are not random; rather, they display a certain degree of regularity and recurrence, which allows for detailed analysis and subsequent interpretation for research purposes [7, 12]. This focus is consistent with long-standing research in second language acquisition, which shows that tracking and categorizing such phenomena can significantly contribute to understanding how learners actually use language and where persistent difficulties occur [10, 11].

From a linguistic perspective, errors can be classified according to both their origin and their formal characteristics. One widely used distinction is that between interlingual and intralingual errors [7, 13]. Interlingual errors arise from the influence of the learner's first language (L1), typically through the transfer of syntactic structures or word order patterns into the target language (L2). Intralingual errors, on the other hand, stem from the internal complexity of the target language itself and are often associated with processes such as overgeneralization or incomplete rule application.

Another commonly used classification focuses on the formal characteristics of errors, that is, on the specific form of the linguistic deviation. Errors may take the form of (a) *omission* of a linguistic element, (b) *addition* of an unnecessary element, (c) incorrect ordering of linguistic units (*misordering*), or (d) use of an incorrect form (*misformation*) [7]. These types of errors not only indicate the degree of mastery of specific linguistic structures but also reflect general developmental tendencies observable in learners as users of the target language. Identifying such errors makes it possible to better understand which linguistic features constitute the greatest obstacles to the development of written production and effective communication. This categorization also enables a targeted and systematic analysis of linguistic deficiencies across different proficiency levels.

In this context, the present analysis focuses on the occurrence and nature of linguistic errors in the written production of students at the University of Defence in relation to their performance in the STANAG 6001 Level 2 writing examination. Systematic errors recurring in the texts produced by students at different proficiency levels represent an important diagnostic tool—not only for assessing compliance with the level descriptors of STANAG 6001 [4], but also for identifying those areas of language that lead to unsuccessful performance in writing. In practical terms, this means that students fail to achieve Level 2 and are instead assessed at Level 1 or 1+. The analysis also applies the above-mentioned typology of errors based on formal characteristics, namely *omission*, *addition*, *misordering*, and *misformation*—which allows for a more detailed description of the nature of the identified linguistic deficiencies.

3. Corpus-Based Research on Error Analysis

Corpus-based research on errors in the written production of learners of English as a Foreign Language (EFL) is characterized by considerable diversity, both in terms of target groups and analytical approaches. One group of studies focuses on general secondary and tertiary students without a specific disciplinary orientation. For example, Mustafa, Kirana, and Bahri [14] examined grammatical errors among Indonesian secondary school students, while Pan and Wang [15] and Zhang [16] analyzed linguistic deficiencies in Chinese university students from non-linguistic fields. These studies primarily identified persistent difficulties in the use of verb tenses, verb patterns, and lexis, despite the fact that learners often demonstrated a relatively extensive vocabulary.

A second group comprises learners with higher levels of proficiency or with a more specific academic background. Divsar and Heydari [17], for instance, analyzed essays written by candidates taking the international IELTS examination and found the highest frequency of errors in lexical choice and verb forms. Doolan [18] compared linguistic errors across three groups: L1 students (native speakers of English), L2 learners (users of English as a foreign language), and so-called Generation 1.5 learners (individuals who immigrated to an English-speaking environment at an early age and display characteristics of both L1 and L2 users). The third group showed a higher incidence of errors in verb constructions and the use of prepositions. A specific case is represented by the longitudinal study conducted by Sasi and Lai [19] among Taiwanese university students, in which errors related to tense usage, prepositions, and subject–verb agreement were found to be most frequent. With regard to research methods, most studies employ a combination of quantitative and corpus-based approaches. Tools such as AntConc [20], along with other annotation and tagging software, enable the systematic analysis of error frequency, distribution, and typology. Díaz-Negrillo and Domínguez [21] further demonstrated that analyzing associations between different types of errors can contribute to more precise diagnostic procedures and subsequently to more consistent evaluation of written performance.

The findings of these studies confirm that precise classification and systematic analysis of errors provide valuable data for feedback as well as for a deeper understanding of the areas in which learners most frequently encounter difficulties. The present study builds on these insights by analyzing authentic written texts produced by students at the University of Defence as

part of the written component of the STANAG 6001 examination. Using corpus-based tools, it identifies the most frequent types of errors in relation to proficiency levels and examines their relationship to the overall examination outcome. This analysis is distinctive not only in its focus on a specific target group—students of a military higher education institution—but also in that the analyzed material originates from a highly formalized and standardized testing environment, the results of which have direct implications for the candidates’ professional careers.

4. Data Specification and Research Context

The dataset analyzed in this study consists of written texts produced by students at the University of Defence as part of the STANAG 6001 English language examination between 2020 and 2024. The sample includes students who achieved proficiency levels 1, 1+, and 2. A total of 120 participants were selected, evenly distributed by gender and proficiency level, with each level comprising 20 male and 20 female students. Each participant completed two tasks from the written component of the examination, resulting in a corpus of 240 authentic texts, which form the basis of the STANAG 6001 Corpus. Table 1 summarizes the distribution of participants according to proficiency level and gender.

Table 1.

Distribution of students and written outputs by proficiency level and gender

STANAG 6001 Level	Gender	Number of students	Tasks per student	Total texts
Level 1	Male	20	2	40
	Female	20	2	40
Level 1+	Male	20	2	40
	Female	20	2	40
Level 2	Male	20	2	40
	Female	20	2	40
Total		120		240

(Source: adapted from [5], p.86)

The standardized language proficiency examination at Level 2 assesses functional writing skills through two tasks. The first task typically takes the form of a short notice or message (with a recommended minimum length of 70 words), while the second task requires a longer piece of written communication (recommended minimum of 150 words), most commonly a report, invitation, or complaint. The prompts for the second task are usually composed of several points and reflect real-life communicative situations that members of the Czech Armed Forces may encounter. For example, when writing a report on an incident, candidates are expected to specify what happened, when and where it occurred, describe the circumstances, and propose preventive measures. These tasks are designed to assess not only grammatical accuracy and lexical range, but above all the ability to organize information coherently and communicate effectively in a military and professional context. As such, they correspond to the requirements defined for the STANAG 6001 Level 2 writing in English by the Bureau for International Language Co-ordination [4].

Such a structured dataset provides a representative basis for analyzing errors in an authentic testing environment. The balanced distribution of participants by gender (male, female) and proficiency level (1, 1+, and 2), together with the total of 240 texts, allows for a reliable examination of the relationship between linguistic errors and achieved proficiency in the written component of the STANAG 6001 examination. All data were anonymized in accordance with established research ethics.

5. Methods of Investigation

This study adopts a corpus-based approach to the analysis of errors in the written production of students at the University of Defence who completed the STANAG 6001 English language examination at proficiency levels 1, 1+, and 2. A mixed-methods research design was applied, combining quantitative analysis of error frequency with qualitative examination of error typology [22, 23]. The aim was not only to identify the frequency of linguistic deviations, but also to examine their distribution and characteristic features across different proficiency levels.

Two corpus analysis tools were employed for the processing of the linguistic data: *TeiTok* [24] and *VersaText* [26]. *TeiTok* (Text Encoding Initiative for Tokenization) was used for the creation, management, and annotation of the text corpus. It is an advanced web-based platform built on TEI standards, enabling detailed linguistic annotation, including the tagging of orthographic, morphosyntactic, and lexical errors. Through its integration with XML technology, the platform allows for both machine-readable and human-readable annotations, facilitating systematic data retrieval and analysis.

Each text was uploaded into *TeiTok* [24] and subsequently annotated using the *SCOPE* and *SUBSTANCE* framework [25], which enables a structured representation of linguistic deviations. First, each instance of error was identified as a deviation from the standard language norm. It was then assigned a *SCOPE* label, indicating the level at which the error occurs (e.g. word, phrase, sentence), and a *SUBSTANCE* label, specifying the nature of the error (e.g. suffix error, punctuation error, spelling error).

Complementary analysis was conducted using *VersaText* [26], a web-based application designed for the visualization and statistical processing of textual data. The tool enables the generation of word clouds, in which the size of individual words reflects their frequency within the corpus. It also supports concordance analysis, allowing individual words to be examined in their immediate textual context, as well as the analysis of lexical distribution patterns. In addition, *VersaText* provides a range of statistical outputs, including word, sentence, and character counts, type–token ratio, the proportion of function words (e.g. articles, auxiliary verbs) and content words (e.g. nouns, adjectives), measures of lexical complexity, and readability indices (e.g. Flesch–Kincaid Reading Ease). These features made it possible not only to assess the range and distribution of linguistic resources, but also to examine the structural complexity of students’ written production across the three STANAG 6001 proficiency levels (i.e. 1, 1+, and 2).

6. Investigation Results

The following sections draw on three complementary analytical approaches which together contribute to a deeper understanding of the relationship between students’ linguistic error patterns and their success in achieving Level 2 in the STANAG 6001 examination. The first analytical framework is based on data annotated in the *TeiTok* [24] environment and focuses on the occurrence of orthographic, morphosyntactic, and lexical errors. This approach makes it possible to identify the linguistic domains in which deviations from the standard most frequently occur, as well as to assess their potential impact on overall performance.

The second layer of analysis builds on the same *TeiTok* [24] annotation but focuses on the formal characteristics of the identified deviations—specifically whether a linguistic element is omitted (*omission*), used incorrectly (*misformation*), or added unnecessarily (*addition*). This type of classification provides insight into the underlying nature of learners’ difficulties, indicating whether they stem primarily from insufficient acquisition, uncertainty in language use, or an excessive attempt at linguistic accuracy.

The third analytical dimension employs the *VersaText* [26] tool to examine selected linguistic indicators, such as the average length of words and sentences, their overall frequency across the three proficiency levels, and text readability. These indicators reflect the degree of linguistic elaboration and structural complexity in written production and may serve as important predictors of success in the STANAG 6001 examination. Taken together, these approaches capture different aspects of language proficiency and provide a comprehensive picture of the factors that may either facilitate or hinder the attainment of Level 2 in a standardized testing environment.

6.1 Orthographic, Morphosyntactic, and Lexical Errors as Indicators of Language Proficiency

Based on detailed annotations carried out in the *TeiTok* [24] environment, it was possible to quantify language errors across individual categories and proficiency levels. Following the refinement and verification of the annotation scheme, the data were aggregated into three principal groups—orthographic, morphosyntactic, and lexical errors. The results of this quantitative analysis indicate a clear relationship between both the type and frequency of errors and the language proficiency level achieved in the STANAG 6001 examination. The lowest incidence of errors was observed among students who attained Level 2, whereas students at Level 1 exhibited the highest error rates across all categories. This finding confirms that certain types of language errors may be regarded as predictors of failure, that is, as factors that significantly contribute to not achieving the required level of language proficiency.

6.1.1 Orthographic Errors

Orthographic errors constitute a substantial proportion of the overall error load (1,341 instances), with spelling errors representing the dominant category at 62.4% of all recorded cases (Table 2).

Table 2.
Overview of orthographic errors

Error type	Number of errors	Share (%)
Spelling	837	62.4 %
Punctuation	283	21.1 %
Capitalization	221	16.5 %
Total	1,341	100 %

This high proportion indicates that mastery of the written form of the language remains a major challenge for learners. An analysis of error distribution by proficiency level further revealed that students classified at Level 1 accounted for nearly half of all orthographic errors (42.6%), suggesting that this level is associated with the greatest difficulties in adhering to spelling and punctuation conventions (Fig. 1). In contrast, students at Level 2 contributed only 21.4% of the total number of orthographic errors, confirming that increasing language proficiency is accompanied by a marked reduction in error rates. These findings suggest that the ability to adhere to orthographic norms represents an important indicator of examination success and may serve as a reliable predictor of overall language proficiency.

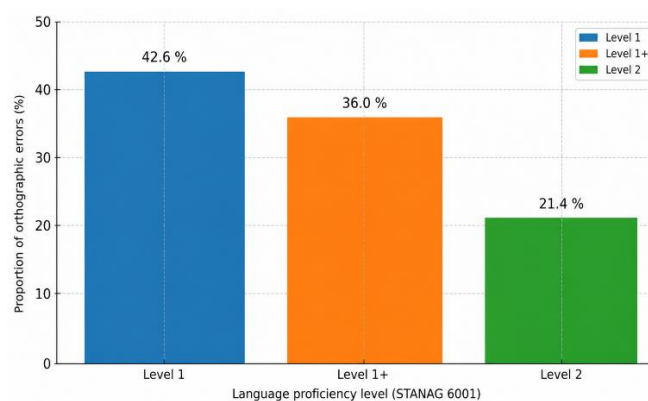


Fig. 1. Distribution of orthographic errors by proficiency level (created with ChatGPT [35])

6.1.2 Morphosyntactic Errors

From a morpho-syntactic perspective, errors in the use of articles emerge as the most problematic category, accounting for 42.4% of all recorded cases (Table 3).

Table 3. Overview of morphosyntactic errors

Error type	Number of errors	Share (%)
Articles	907	42.4 %
Suffixes	588	27.5 %
Auxiliaries	318	14.9 %
Pronouns	254	11.9 %
To-infinitives	68	3.2 %
Prefixes	3	0.1 %
Total	2,138	100 %

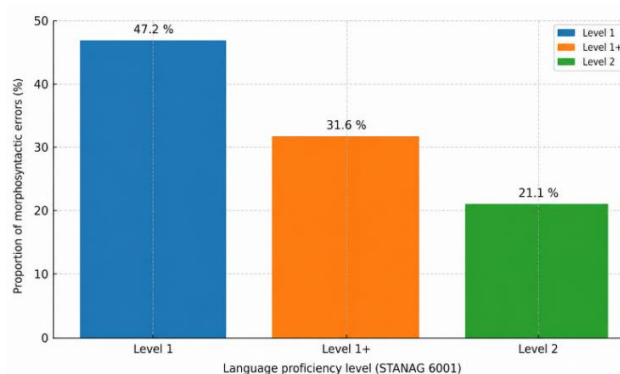


Fig. 2. Distribution of morphosyntactic errors by proficiency level (created with ChatGPT [35])

These are followed by errors in word formation through suffixation (27.5%) and incorrect use of auxiliary verbs (14.9%). Together, these three categories constitute the vast majority of morpho-syntactic errors, highlighting areas that require increased attention in both instruction and targeted corrective feedback. The distribution of errors across proficiency levels is again highly indicative: nearly half of these errors (47.2%) were produced by students at Level 1, whereas students at Level 2 accounted for only 21.1% of the total (Fig. 2). This marked disparity confirms that accurate command of morpho-syntactic structures—particularly the systematic use of articles (*the, a, an*), precise formation of word forms (e.g. regular verb endings *-d/-ed*), and appropriate selection of auxiliary verbs (e.g. *do/did*)—represents a key factor in achieving higher levels of language proficiency and, consequently, in successfully passing a standardized examination such as STANAG 6001.

6.1.3 Lexical Errors

In the lexical domain (see Table4), errors in the use of prepositions (50.7%) and content words (44.3%) predominate, while errors in conjunctions account for only a minor proportion (5.0%). This indicates that, in vocabulary acquisition, learners most frequently encounter difficulties related to selecting the appropriate preposition or using semantically precise lexical items. With regard to distribution across proficiency levels, students at Level 1 exhibit the highest proportion of these

lexical errors (46.80%), confirming that lower proficiency levels are characterized by a greater degree of uncertainty in handling lexical resources. At Level 2, the proportion of lexical errors decreases markedly to 18.65%, indicating that increasing language competence leads to improvement not only in morphosyntactic accuracy but also in lexical precision (Fig. 3).

Table 4. Overview of lexical errors

Error type	Number of errors	Share (%)
Prepositions	459	50.7 %
Content words	401	44.3 %
Conjunctions	46	5.0 %
Total	906	100 %

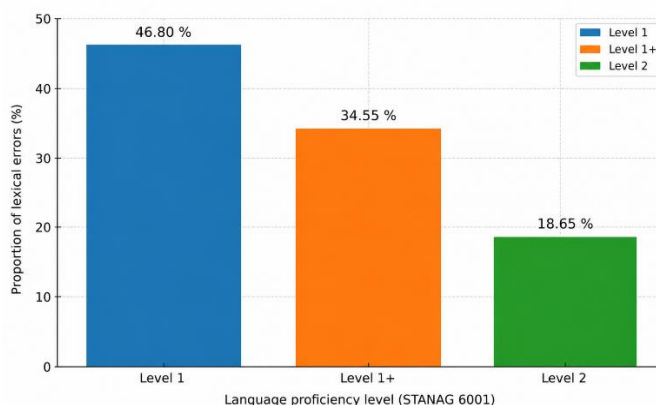


Fig. 3. Distribution of lexical errors by proficiency level (created with ChatGPT [35])

This trend confirms that the acquisition and consolidation of vocabulary, including the correct use of prepositions and content words, constitutes a key factor in achieving higher levels of language proficiency.

6.2 Errors by Type of Deviation: Omission, Addition and Misformation

Building on the previous analysis, a further examination focuses on the formal characteristics of language errors, drawing on the classification presented in the scholarly literature (e.g. [7], p. 37). This approach distinguishes errors according to their manifestation, namely *omission*, *addition*, *misformation*, and *misordering* of linguistic units. This typology was applied during the annotation of texts in the *TeiTok* [24] environment, whose flexible tagging system made it possible to assign each error an appropriate formal category.

Within this study, primary attention was given to the three most frequent categories—*omission*, *addition*, and *misformation*—which demonstrated a systematic relationship with the students' language proficiency level. In contrast, *misordering* errors were only minimally represented in the analyzed corpus and, due to their low frequency, were not included as a separate category in the quantitative evaluation.

Table 5. Types of errors by proficiency level and their percentage distribution

Omission			Misformation			Addition		
level	count	percent	level	count	percent	level	count	percent
1	784	48.48%	1	465	45.32%	1	322	47.08%
1+	486	30.06%	1+	349	34.02%	1+	225	32.89%
2	347	21.46%	2	212	20.66%	2	137	20.03%

This form of classification complements the previous perspective on errors according to linguistic level (orthographic, morphosyntactic, lexical) and enables a more detailed understanding of how students work with the language—what is missing in the text, what is redundant, and what is used incorrectly. It appears that this structural aspect of errors is closely related to the achieved language proficiency level and may represent one of the key factors influencing success in the STANAG 6001 examination.

Table 5 and Fig. 4 indicate that students most frequently make errors caused by the *omission* of linguistic elements, and this occurs across all the levels examined. Nearly half of these errors are produced by students at Level 1, whereas their occurrence is noticeably lower among students at Level 2. This type of error often reflects an insufficient command of basic

grammatical structures—such as the omission of articles, auxiliary verbs, or prepositions—and appears to be one of the most significant factors preventing the attainment of a higher level of language proficiency.

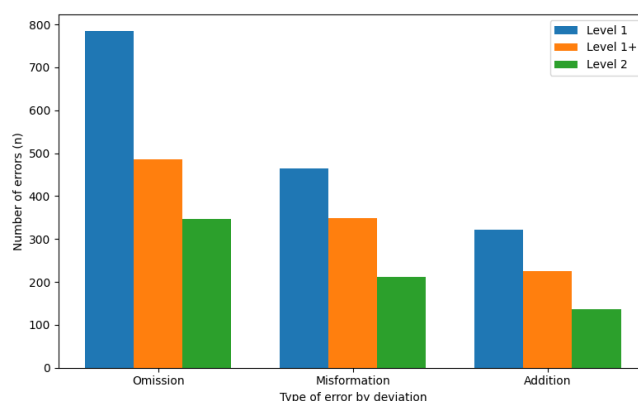


Fig. 4. Number of errors by language proficiency level and type of deviation (created with ChatGPT [35])

Misformation errors are also more frequent at lower proficiency levels; however, the differences are not as pronounced as in the case of *omission* errors. By contrast, *addition* represents the least frequent category overall, yet a consistent trend can still be observed: students at Level 2 produce substantially fewer errors than their less advanced counterparts.

Overall, this analysis confirms that not only the quantity of errors, but also their type and linguistic function, constitute an important indicator of language proficiency. Students who are able not only to form linguistic structures correctly but also to use them appropriately within a text demonstrate a higher degree of language control, which is a key prerequisite for success in standardized examinations such as STANAG 6001.

6.3 Textual Complexity as a Predictor of a Higher Proficiency Level

In addition to analyzing errors in terms of their type and frequency, the study also examined quantitative features of students' written production at individual STANAG 6001 language levels. For this purpose, the *VersaText* [26] tool was employed. Besides lexical analysis, this tool enables the automated processing of linguistic data, such as the number of words, sentences, and characters, average sentence length, word complexity, and readability index. Such a statistical framework provides an important complementary perspective on students' language proficiency. It also makes it possible to compare not only how many errors students make, but also how developed their written production is.

Table 6. Comparison of quantitative linguistic indicators across STANAG 6001 levels

Indicator	Level 1	Level 1+	Level 2
Total number of words	10,283	10,844	11,557
Total number of sentences	797	820	762
Average number of words per sentence	12.9	13.22	15.17
Average number of characters per word	5.24	5.28	5.31
Readability index (Flesch–Kincaid)	69.81	68.83	65.09
Flesch–Kincaid CEFR estimate	A2–B1	A2–B1	B1

The average number of words per sentence increases systematically with language proficiency (see Table 6) — from 12.9 (Level 1) through 13.22 (Level 1+) to 15.17 words at Level 2. At the same time, however, the absolute number of sentences decreases between Levels 1+ and 2 (820 → 762). This combination suggests that learners at higher proficiency levels tend to condense information into longer sentence structures, producing a greater proportion of complex or elaborated sentences, which serves as a direct indicator of increasing syntactic complexity.

The average number of characters per word also shows a slight increase (from 5.24 to 5.31), which may be interpreted as evidence of the use of longer and potentially more advanced lexical items. Together, these two indicators point to a growing level of linguistic complexity and lexical diversity in relation to the achieved proficiency level.

The Flesch–Kincaid readability index [27] decreases with increasing proficiency (69.81 → 65.09), which corresponds to a generally higher level of text difficulty—the lower the index, the more complex the text is to comprehend. This phenomenon reflects the fact that students at higher STANAG 6001 levels produce linguistically more complex texts that require greater flexibility and more advanced use of linguistic resources.

The estimated CEFR (Common European Framework of Reference for Languages) [28] level falls within the A2–B1 range in all cases, with Level 2 consistently corresponding to B1. It should be noted, however, that the descriptors of STANAG 6001 and the CEFR are not fully equivalent, and therefore cannot be directly mapped onto each other.

Nevertheless, this estimate provides a useful point of reference in relation to the widely recognized European framework of language proficiency.

7. Discussion and Conclusions

The present study focused on identifying factors influencing the success or failure of students at the University of Defence in the writing component of the STANAG 6001 English language examination, based on a corpus analysis of 240 authentic texts. The findings indicate that differences between students who achieved Level 2 and those assessed at Levels 1 or 1+ are not determined solely by the absolute number of errors, but primarily by their type, structure, and linguistic function.

The most significant factor appears to be the occurrence of systematic errors, particularly in the morphosyntactic domain (e.g. the use of articles, suffixes, and auxiliary verbs), which were more frequent among less successful students. In addition, omission errors proved to be especially influential in relation to proficiency level, with nearly half of such errors produced by Level 1 students. Their presence suggests an insufficient command of basic grammatical structures and a limited ability to express linguistic relations effectively in formal written discourse.

Attention was also paid to quantitative indicators of written production, such as sentence length, word length, and readability index. Students who achieved Level 2 produced longer and more complex sentences, used longer lexical items, and exhibited lower readability index values, indicating a higher degree of stylistic complexity. These indicators may therefore be understood as complementary to error analysis, offering an additional perspective on language proficiency.

The findings suggest that success in the writing component of the STANAG 6001 examination depends on a combination of accuracy, completeness, and linguistic complexity. Students who demonstrate not only a low number of errors but also the ability to structure texts effectively and appropriately employ the expected linguistic features, in accordance with STANAG 6001 descriptors [4], are more likely to achieve the required level. The results may thus inform adjustments in teaching strategies, strengthen the emphasis on targeted feedback, and support effective intervention in the most problematic areas, particularly in the accurate and complete expression of grammatical relationships.

Corpus-based analysis grounded in authentic language data thus provides valuable insights for both pedagogical practice and further research in military language education. The methodology and findings presented here may serve as a basis for targeted support of students aiming to meet language proficiency requirements in a standardized examination such as STANAG 6001 Level 2. In view of the international character of this examination, a similar research approach may offer useful insights for linguists and educators working in military higher education institutions across NATO member states when analyzing written production and language proficiency within their respective national contexts. The role of English as the principal language of international military cooperation, command, and operational communication further underscores that writing competence is essential not only for assessment purposes but also for effective professional performance in multinational defence environments. Moreover, the findings contribute to the growing body of research in corpus linguistics [29, 30], learner corpus research [31, 32, 6], and data-driven learning approaches [33, 34, 36], particularly within the context of military education, where emphasis is placed on mastering functional English for both formal and operational communication. The study also addresses a relatively underexplored area in pedagogical linguistics by linking the analysis of student written production at the University of Defence with data derived from a highly standardized international examination, namely STANAG 6001.

Limitations. Several limitations of the present study should be acknowledged. First, the corpus comprises written responses produced by military students from a single institution, which may limit the generalizability of the findings to other military or civilian contexts. Second, the analysis focuses exclusively on the written component of the STANAG 6001 examination and therefore does not capture potential interactions between writing proficiency and performance in other skills, particularly speaking. Third, although corpus-based error annotation allows for systematic and replicable analysis, the categorization of errors inevitably involves a degree of interpretative judgment and may not fully reflect learners' communicative intent or strategic language use. Finally, the cross-sectional design provides a snapshot of proficiency-related differences but does not allow for tracking individual developmental trajectories over time.

Acknowledgements. The research was conducted within the framework of the DZRO (Long-Term Development Plan of Organization) Language II project under the auspices of the Ministry of Defence of the Czech Republic.

References

1. Šikolová M., Holcner V. Efficiency of Language Education at the Language Center of the University of Defence. *Czech Military Review*, 2020, 29 (2), p.37–48.
2. Green R., Wall D. Language Testing in the Military: Problems, Politics and Progress. *Language Testing*, 2005, 22 (3), p.379-98.
3. Monaghan R. Language and interoperability in NATO: the bureau for international language co-ordination (BILC). *Canadian Military Journal*, 2012, 13 (1), p.23-32.

4. BILC (Bureau for International Language Co-ordination). STANAG 6001. Available online: <https://natobilc.org/wp-content/uploads/2024/11/STANAG-6001-Overview-Feb-2019.pdf> (accessed on 15 March 2026).
5. **Helán R., Bednářová R.** Error Analysis and Frequency in the Writing Section of the STANAG 6001 English Examination among Military Students. *CASALC (Czech and Slovak Association of Language Centres) Review*, 2024, 14 (2), p.79-101.
6. **Götz S., Granger S.** Learner corpus research for pedagogical purposes: An overview and some research perspectives. *International Journal of Learner Corpus Research*, 2024, 10 (1), p.1-38.
7. **James C.** Errors in language learning and use: Exploring error analysis. 2013, Routledge.
8. **Richards J. C.** Error Analysis: Perspectives on Second Language Acquisition. 2014, London: Routledge.
9. **Ferris D.** Treatment of Error in Second Language Student Writing. 2011, Michigan: University of Michigan Press.
10. **Corder S. P.** The Significance of Learner's Errors. *International Review of Applied Linguistics in Language Teaching*, 1967, 5 (4), p.161-70.
11. **Corder S. P.** Error Analysis and Interlanguage. 1982, Oxford: Oxford University Press.
12. **Ellis, R.** Second Language Acquisition. 1997, Oxford: Oxford University Press.
13. **Freiermuth M. R.** L2 Error Correction: Criteria and Techniques. *Language Teacher (JALT)*, 1997, 21 (1), p.13-18.
14. **Mustafa F., Kirana M., Bahri S.** Errors in EFL Writing by Junior High Students in Indonesia. *International Journal of Research Studies in Language Learning*, 2017, 6 (1), p.38-52.
15. **Pan J., Wang J.** A Corpus-based Study on Chinese Students' Writing Errors. In *International Conference on Social Science, Education Management and Sports Education*. Atlantis Press, 2015, p.113-116.
16. **Zhang R.** A Corpus-based Error Analysis of Students' Writing in Graded Teaching Classes. *Journal of Convergence Information Technology*, 2013, 8 (10), p.551.
17. **Divsar H., Heydari R.** A corpus-based study of EFL learners' errors in IELTS essay writing. *International Journal of Applied Linguistics & English Literature*, 2017, 6 (3), p.143.
18. **Doolan S. M.** Generation 1.5 writing compared to L1 and L2 writing in first-year composition. *Written Communication*, 2013, 30 (2), p.135-163.
19. **Sasi A. S., Lai J. C. M.** Error analysis of Taiwanese university students' English essay writing: A longitudinal corpus study. *International Journal of Research in English Education*, 2021, 6 (4), p.57-74.
20. **Aqeel M., Rubbani A., Hamid A.** Error analysis in learners' essay writing: A corpus-based study. *International Journal of Special Education*, 2022, 37 (3), p.18331-41.
21. **Negrillo A. D., Domínguez J. F.** Error tagging systems for learner corpora. *Revista española de lingüística aplicada*, 2006, 19 (1), p.83-102.
22. **Creswell J. W., Creswell J. D.** Research design: Qualitative, quantitative, and mixed methods approaches. 2017, Thousand Oaks, California: SAGE Publications.
23. **Johnson R. B., Christensen L. B.** Educational research: Quantitative, qualitative, and mixed approaches. 2024, Los Angeles, London, New Delhi: Sage publications.
24. TeiTok. Available online: <http://teitok.corpuswiki.org/> (accessed on 5 March 2026).
25. **Dobrić N.** (2023). Identifying errors in a learner corpus—the two stages of error location vs. error description and consequences for measuring and reporting inter-annotator agreement. *Applied Corpus Linguistics*, 2023, 3 (1), p.100039.
26. VersaText. Available online: <https://versatext.versatile.pub/> (accessed on 10 March 2026).
27. Flesch–Kincaid readability index. Available online: <https://readable.com/readability/flesch-reading-ease-flesch-kincaid-grade-level/> (accessed on 2 February 2026).
28. The CEFR Levels. Available online: <https://www.coe.int/en/web/common-european-framework-reference-languages/level-descriptions> (accessed on 20 April 2026)
29. **Paquot M., Gries S. T.** A practical handbook of corpus linguistics. 2022, Springer Nature.
30. **Egbert J., Biber D., Gray B.** Designing and evaluating language corpora: A practical framework for corpus representativeness. 2022, Cambridge: Cambridge University Press.
31. **Götz S., Callies M.** Learner corpora in language testing and assessment. 2015, Amsterdam, Philadelphia: John Benjamins Publishing Company.
32. **Götz S.** Learner corpora to inform testing and assessment. In *The Routledge handbook of corpora and English language teaching and learning*. 2022, Routledge, p.311-326.
33. **Poole, R.** Guide to Using Corpora for English Language Learners. 2018, Edinburgh: Edinburgh University Press.
34. **Boulton A., Leńko-Szymańska A.** Multiple affordances of language corpora for data-driven learning. 2015, Amsterdam, The Netherlands: John Benjamins Publishing Company.
35. ChatGPT. Available online: <https://chatgpt.com/> (accessed on 25 April)
36. **Helán R., Pořízková K.** Corpus tools in language for specific purposes: Technology-enhanced learning and emerging AI support. In *Holcner V. et al. Current challenges of language education in contemporary security environment*. 2025, Praha: powerprint s.r.o. – Bookla, p.104-111.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of CNDCGS 2026 and/or the editor(s). CNDCGS 2026 and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.