

Error Recognition Model: High-mathability End-user Text Management

**Sebestyén Katalin^a, Csapó Gábor^a, Mária Csernoch^b,
Bernadett Aradi^b**

^a Doctoral School of Informatics, University of Debrecen
Kassai út 26, 4028 Debrecen, Hungary
sebestyen.katalin@inf.unideb.hu, csapo.gabor@inf.unideb.hu

^b Faculty of Informatics, University of Debrecen
Kassai út 26, 4028 Debrecen, Hungary
csernoch.maria@inf.unideb.hu, aradi.bernadett@inf.unideb.hu

Abstract: Discussion, evaluation and error recognition, in natural language digital texts, is one of the most neglected areas in the digital world, despite the fact that text management is the most prevalent computer related activity. Millions of erroneous text-based documents of different types are in circulation, without us being aware of how fragile, damaged and harmful they are. It is well accepted in programming and even in other end-user activities, that error recognition plays a crucial role in teaching, learning and in real-world problem-solving processes. In the present paper, we introduce the High-mathability Error Recognition Model, which consists of the processes used in discussion and concept-based problem-solving and we also provide examples of the utilization of the model. We argue that error recognition and correction, and the assessment of problems in text management are as important as in other fields of informatics and computer sciences. In our study experimental groups – studying with the Error Recognition Model – and control groups – studying with low-mathability tool-centered approaches – were compared. It was found that the Error Recognition Model is more effective in digital text management, than for the tool-centered methodologies, in two error types: the typographic and layout-breaking error categories and a strong compensation effect was found in the syntax error category.

Keywords: teaching/learning strategies; lifelong learning; improving classroom teaching; data science applications in education; information literacy

1 Introduction

A well-accepted and effective method in teaching programming is to present erroneous source codes for students to detect and correct mistakes and errors [1] [2] [3]. This method is used to teach students program-formulation and language-constructs [4] and to prepare them for searching for errors, either in their own or in

others' programs. However, program codes are not the only source of errors in the digital world. Natural language texts are more error-filled [5] [6] [7] [8] [9] [10] [11] [12] [13], a problem which has various reasons: the wide variety of languages, the grammar of languages – which is more complex than that of artificial languages –, and the level of users' knowledge (none do programming on a daily basis, without any background knowledge). Beyond the special features of languages, a text should meet several other requirements. The contents of the texts – their creation with the attendant typing difficulties or from copying, not to mention any copyright issues, and their language settings and spelling and grammar checkers – also raise considerable issues. Among these requirements, typographic [14] [15] [16] [17] [18] rules play a crucial role, with an emphasis on the various interfaces used to display content, and how to make the content the most legible on the selected interface.

However, because word processors are available for all, we are faced with several misconceptions regarding end-user text management, which leads to widespread 'bricolage' [5] [6] [7] [8] [9] [10] [11] [12] [13]. The first problem is that end-users are untrained or self-trained, and, compared to programmers, most of them are never formally taught how to design a text, how to build an algorithm for the text-management process, and how to format properly [24] [28]. They do not even know what the requirements of a properly formatted text are [12] [13]. They are in the comfortable state of not knowing what they do not know, the Dunning-Kruger effect [19].

The second problem is that end-user text management, in most cases is identified as interface-navigation – wrapped up in the 'user-friendly' slogans of the software companies – which leads to low-mathability activities [7] [29] [30] [31]. Consequently, nothing is learned, nothing is stored in long term memory, and no schemata are constructed, which leads to the intensive use of attention mode thinking, which is more prone to produce errors than automatic mode thinking. It is not surprising that some researchers find text and spreadsheet management low level, boring, routine activities [20] [21], in spite of the warnings that end-user computing should be taken seriously [22] [32]. This discrimination against end-user computing can also be explained by teachers' belief in the "fixed" nature of the subject and their own low-level efficiency [23].

1.1 Text Management – a Problem Solving Approach

In end-user text management, the most common document types are digital texts (texts handled with word processors and text editors), presentations, and web pages. Within this framework, the text management process follows the four steps of Pólya's concept-based problem-solving approach [24] [12] [25] [27]. In our context, the focus is on the discussion element – i.e. on evaluation, error recognition, classification and correction, which is ever-present in natural language text management, and which therefore has a great impact on the whole process and the

output(s). In the present paper we introduce the Error Recognition Model and detail a test to compare the effectiveness of the traditional, surface approach methods and the Error Recognition Model.

The concept-based problem solving approach of Pólya [24], recently recognized as a high-mathability problem solving method [29] [30] [31] [33] [34] [35] [36] [37], has proved efficient both in maths and programming (and many other subjects); the different phases of this approach are recognized as the levels of mastery in the digital age [25], and have been found to be adaptable to the other popular end-user activity, such as spreadsheet management [38] [39].

Pedagogical Content Knowledge (PCK) [40] [41] and/or Technological Pedagogical Content Knowledge (TPCK) [39] are required for effective teaching in the digital era. Most teachers apply low mathability approaches to computer problem solving; teaching their students with a fixed belief in science – IT, ICT, CS – and with low self-efficacy. These approaches mainly focus on interface navigation, and passing exams, without providing students opportunities to fully understand and appreciate science [23] and can lead to erroneous end-user activities [5] [6] [7] [8] [9] [10] [11] [12] [14] [13] [22]. As mentioned above, effective end-user text management requires the cooperation of teachers of different subjects. Consequently, teacher education should be changed according to the requirements of high mathability problem solving; this would be preferable to teaching how to use a software environment in pre- and in-service courses. A similar approach was suggested in spreadsheet management by Angeli in 2013 [39], but the idea has not reached the wider public, and remains isolated. Crowd sourcing would help us reach a wide range of teachers who would understand that end-user text management involves lot more than typing, and copying – not infrequently without naming sources – and clicking on the buttons of the toolbars offered by the software companies [42].

Crowd sourcing would also help find IT professionals who at present, mostly ignore end-user computing [22] [43] and could recognize that High Mathability End-user Computing would be an effective introduction to serious computing.

1.2 The Error Recognition Model (ERM)

Due to the extremely high number of possible errors in a digital natural language text, it is necessary to classify them. The main classes of errors are constructed on the systems of rules which create the major guidelines for a correct natural language text. A digital text should fulfill the requirements of the language(s) and the content of the text, as well as the rules relating to displaying, breaking or layout, and formatting. These are the most common rule systems which should be taken into consideration in the process of end-user text management. Violation of these rules leads to errors of the following types: syntactic, semantic, typographic, layout or breaking, formatting, and style errors [12] [27] [28].

2 Hypotheses

Based on the theoretical background of the high-mathability Error Recognition Model in text-management, we launched a research project, whose aim was to prove the effectiveness of the approach compared to the interfaces centered, traditional methods. Five hypotheses are formulated to see how the error recognition abilities of students develop using the traditional and the newly introduced ERM model. To complete our study, a testing series was administered in primary and secondary schools within the topic of word processing and text-management.

- [H1] In the pre-test there is no difference between grade 9 experiment (9E) and control groups (9C).
- [H2] In the pre-test there is a difference between grade 7 and grade 9 students.
- [H3] In the post-test there is a difference between grade 9 experiment (9E) and control groups (9C).
- [H4] In the post-test there is a difference between grade 7 experiment (7E) and grade 9 experiment groups (9E).
- [H5] In the post-test there is a difference between grade 7 experiment (7E) and grade 9 control groups (9C).

3 Testing

To quantify and prove the efficiency of the ERM method, we tested experiment groups where this novel, high-mathability [28] [29] [30] [31] [33] [34] [35] [37] approach was introduced, and compared their results to control groups where the traditional, low-mathability, tool-focused methods are used.

3.1 Sample

The teaching and testing process took place in the academic year 2017/2018, in one of the high schools in Debrecen, Hungary. Students from grade 7 and 9, formed both the experimental and control groups (Table 1). Considering the background knowledge of the participating students, the selection of groups plays a crucial role, because all of the students had learned the basics of word processing – whatever the word “basic” means in this context – in primary education. In general, according to the Hungarian frame curricula [44] [45], these students are able to construct text documents based on a sample provided [45]. The methods with which students were taught in primary education are not documented; however, the structure of the frame curricula and the structure and content of the textbooks clearly indicate that primarily the traditional, tool-focused methods are applied.

During our experiment, students were tested in two rounds: in a pre- and a post-test. The pre-tests were administered in advance of the teaching-learning process, to register what knowledge the students bring with them from their previous studies, while the post-tests were completed at the end of the topic, both groups.

Table 1
The number of students who participated in both tests

	Experimental groups	Control groups	Total
Grade 7	26	–	26
Grade 9	66	38	104
Total	92	38	130

Considering all groups, 153 students completed the pre-test. 34 and 69 students participated in the experiment from grades 7 and 9, respectively. 146 students completed the post-test: 102 in the experiment and 44 in the control groups. Pairing the students, 130 students completed both the pre- and post-tests, 92 from the experiment groups and 38 from the control groups (Table 1).

3.2 Conducting the Measurements

The students had 45 minutes to complete the tests in both the pre- and post-tests. Both tests consist of three phases:

- (1) Unplugged
- (2) Semi-unplugged error recognition
- (3) Plugged-in error correction (not detailed in this paper)

During the unplugged phase, the computers were turned off and each student got a printed version of the text document and a blue pen. The students' task was to scan the printed document and mark (circle) and name the errors or error types which they identify. We did not expect the students to provide the terminologically correct name for each error, but to give a short explanation of why they marked that part of the document as erroneous.

In the next phase we collected the blue pens and the students opened the electronic version of the document and got a red pen to mark and explain on the paper the errors they discover in the digital version of the document.

The third and last phase of the test was to correct the errors and to save the document using a name and folder provided. In the present study the results of the recognition of errors and the analyses of these tasks are presented. We also must note that recognizing the semantic errors in the documents is beyond the scope of this analysis; consequently, any further details are not provided regarding this category of errors.

The process of testing and the evaluation of the pre- and post-tests are identical. However, a larger amount of text and more complex errors were presented in the post-test. The length and quality of the document used in the post-test adjusts to the acquired level of knowledge in the topic and is matched with the students' time management ability. The differences between the documents selected for the pre- and post-tests originated from our partially different goals. In the pre-test we were interested in documenting the knowledge the students brought with them from their previous studies, while in the post-test, we wanted to register the students' improvement in applying the different approaches in the teaching-learning process.

3.3 Errors in the Test Documents

In Table 2, the errors, the error types, and the place of recognition of the pre-test are listed. In Table 3 we present the errors of the post-test and also in the Error column we mark which errors occur in both tests (PE/PO) or only in the post-test (PO), where PE refers to the pre-test and PO to the post-test.

Table 2
Errors in the pre-test, categorized by error types

Error type	Error	Place
syntactic	– spelling mistakes – improper use of parentheses with Space characters	printed
typographic	– underline – whole text italic	printed
layout-breaking	– empty paragraphs – paragraph marks at the end of each line – indentation with Space characters – alignment with Space characters – manual hyphenation	digital

Table 3
Errors in the post-test, categorized by error types. Errors marked PE/PO are present in both the pre-test and the post-test, while errors marked PO are only present in the post-test.

Error type	Error	Place
syntactic	– spelling mistakes (PE/PO) – improper use of parentheses with Space characters (PE/PO) – missing Space characters (PO)	printed
typographic	– underline (PE/PO) – italic (PE/PO) – bold (PO) – all capitals (PO)	printed

layout-breaking	– empty paragraphs (PE/PO) – paragraph marks at the end of each line (PE/PO) – indentation with Space characters (PE/PO) – alignment with Space characters (PE/PO) – manual hyphenation (PE/PO) – manual numbering (PO) – multiple Space characters (PO) – character spacing, expanded (PO)	digital
-----------------	--	---------

Students got points if they recognized – marked –, named, and categorized the errors. In the pre-test, there are 7 syntactic, 2 typographic, and 5 different layout-breaking errors of varying appearance. In the post-test we can recognize 10 syntactic, 4 typographic, and 8 layout-breaking errors, where in all three categories multiple occurrences are detectable. We must note here that in the case of repeated instances of the same errors, on recognizing more than half of the same errors, students were awarded additional points.

4 Methods

4.1 The Evaluation Process

Following the administration of the tests, the evaluation process took place. In advance of the actual evaluation process an evaluation table was set up in Excel, where all the items were listed.

The items of the tests were decided based on the nature and the frequency of the errors (Table 4). The following items of the evaluation table served as the basis for the statistical analyses:

- Primarily the errors were grouped on the basis of the three error categories (Table 2 and Table 3): syntactic, typographic, and layout-breaking.
- Within the categories all the errors, were checked according to the previously established smaller items.
 - The errors had to be marked on the paper and the also had to be named (markers without any written notes were not accepted as correct answers).
 - The color of the markers and/or the notes were also checked to reveal the place of recognition. The syntactic and the typographic errors are recognizable on the hard copy, while the layout-breaking errors only in the digital form of the document.
 - In the case of multiple errors, we also checked and recorded how many of the same error were recognized. If more than half of the same errors

were marked, an additional point was added to the sum. For example, in the pretest, three additional points were available for marking multiple space characters, paragraph marks imitating vertical spacing, and paragraph marks at the end of lines.

Table 4

The items of the three errors types and their relative frequency in the three error categories in the pre- and post-tests

	Pre-test		Post-test	
	number of items	rel. frequency	number of items	rel. frequency
syntactic	13	43.33%	37	53.62%
typographic	4	13.33%	12	17.39%
layout	13	43.33%	20	28.99%
TOTAL	30	100.00%	69	100.00%

We must note here that in the pre-test there is a formatting error, where one of the paragraphs is formatted with Keep with next; however, none of the students recognized it. Furthermore, this type of error is not included in the post-test; consequently, we cannot examine how students developed. In general, this error is omitted from the analyses.

The Statistical Analyses

Considering our hypotheses, we focused on the following test groups / comparisons:

- Grade 9 experimental groups (9E) versus grade 9 control groups (9C)
- Grade 9 experimental groups (9E) versus grade 7 experimental groups (7E)
- Grade 9 control groups (9C) versus grade 7 experimental groups (7E)
- Pre-tests versus post-tests

The statistical analyses were carried out with the following methods.

First, to check whether the samples follow normal distribution or not, we used the 1-sample Kolmogorov-Smirnov test. Here the null hypothesis is that the sample is drawn from the reference distribution (normal, in our case).

Because in many cases we found that the samples we want to compare do not follow normal distributions, we used the nonparametric Mann-Whitney U test to compare the samples. The null hypothesis in this case is that the medians of the two samples are the same. We used this test for independent samples.

However, in some cases the normality assumption of Student's t-test was satisfied. In these situations, we used the latter test to compare the means of independent samples. Sometimes we could assume that the variances are equal, sometimes not; we checked this condition with Levene's test. To perform the statistical tests, we used SPSS.

5 Results

5.1 Pre-test

The comparison of the students' results was based on the three error types – syntactic, typographic, and layout-breaking categories – and the errors listed in section 3.3 (Table 2 and Table 3). In the pre-test, considering the total points, the students achieved 20.29%, 31.52%, and 23.67% in the 7E, 9E, and 9C groups, respectively (Table 5). Furthermore, we checked for significant differences between three groups comparing the three error categories. The grades were compared in all the possible variations:

The Kolmogorov-Smirnov Test proved that the results of the pre-tests in the three error categories do not necessarily follow normal distribution. All but one of the three error categories and the three groups of students showed a non-normal distribution:

- 7E: syntax, typography, layout-breaking
- 9E: syntax, typography, layout-breaking
- 9C: typography, layout-breaking

Consequently, we used the nonparametric Mann-Whitney U test to compare the pretest results.

In the typographic errors category all students in all three groups scored 0 points; therefore, there is no difference, so we only checked the other two error categories.

According to the Hungarian National Base Curriculum [44] and the frame curricula [45], grade 9 students had studied word processing in their previous studies; consequently, we assumed that in the pre-test no significant difference would be recognized between the two groups [H1]. Considering the total results, in the comparison of grade 9 students, no difference was found between the experiment and control groups ($p=0.103$). However, in the layout-breaking category the result of the experiment group was significantly higher than the control group ($p=0.014$). Considering the syntactical errors, the 9E groups scored higher, with no significant difference between the two groups ($p=0.879$) (Table 5). This latter result is in accordance with the 9E and 9C groups acceptance results for high school, however we must emphasize that in the totals there is no significant difference between the grade 9 groups [H1].

In the comparison of grade 7 and grade 9 students, based on their previous studies according to the NAT [44] and the frame curricula [45], we expected significant differences between the age groups [H2]. When comparing grade 9E with 7E, in accordance with our [H2] hypothesis, we found a significant difference in the total results ($p=0.003$) and in the layout-breaking error category ($p=0.000$). However, no significant difference was found in the syntax category ($p=0.482$). On the contrary,

between the 9C group and the 7E groups, considering their total results, no significant difference was found ($p=0.401$). In the comparison of the two error groups, neither the syntax ($p=0.505$) nor the layout-breaking (0.083) error groups showed a significant difference. Based on the results of the pre-test, we can neither confirm nor reject our [H2] hypothesis. This leads to the conclusion that previous studies in word processing and text-management do not necessarily build up long lasting, firm, reliable knowledge.

Table 5

The results (%) of the experiment and the control groups in the pre-test regarding the three error categories

	7E	9E	9C
Syntactic	40.72%	43.59%	41.08%
Typographic	0.00%	0.00%	0.00%
Layout-breaking	6.11%	29.14%	13.54%
TOTAL	20.29%	31.52%	23.67%

Considering syntactical errors, there is no significant difference between the two age groups; therefore, the results are independent of age and the previous experiment. In the case of syntactic errors, we can conclude that previous studies either in informatics or native language did not help students in recognizing printed errors of this type. As mentioned above, in the typographic category all students scored 0, which proves that previous studies did not pay attention to the typographic rules regarding the printed version of the documents.

5.2 Post-test

The comparison of the results of the experiment and control groups in the post-test was conducted to reveal the differences between the two teaching-learning approaches: the traditional vs. the ERM method. The average scores of the three groups of students in the three categories are presented in Table 6.

Table 6

The average results (%) of the three groups of students in the post-test

	7E	9E	9C
Syntactic	19.01%	26.40%	35.75%
Typographic	67.53%	60.05%	1.14%
Layout	41.72%	49.11%	35.68%
TOTAL	34.03%	38.83%	29.71%

In the case of the syntactic errors the 9C groups produced better results than both the experiment groups. The 7E and 9E groups reached 19.01% and 26.40%, respectively, while 9C achieved 35.75%. The difference between the groups are significant in all cases (9E vs. 9C $p=0.002$; 7E vs. 9E $p=0.048$; 7E vs. 9C $p=0.000$).

Considering typographic errors in the post-test, the experiment groups (9E vs. 9C $p=0.000$; 7E vs. 9E $p=0.172$; 7E vs 9C $p=0.000$) proved to be significantly better than the control groups ($p=0.000$). This result clearly demonstrates that the typographic problems were not handled in the control groups, while in the experiment groups the subject was covered. The results of the post-test clearly reveal that students of the experiment groups learned the fundamental typographic rules, which play an important role in the presentation of any text-based document. In the comparison of the two age groups of the experiment groups, the results of grade 7 students (67.53%) were somewhat higher than those of grade 9 students (60.05%) (Table 6), but there is no significant difference between the two age groups ($p=0.172$).

In the layout-breaking error category, it was found that grade 9E groups achieved the highest results (49.11%), followed by the 7E groups (41.72%); the 9C groups scored the lowest (35.68%) (Table 6). A significant difference was found between the 9E and 9C groups ($p=0.000$). However, between the 7E and 9E groups ($p=0.450$) and the 7E and 9C groups ($p=0.145$) no significant difference was found.

Summarizing the results of the post-test, we found that in two error types – typographic and layout – the experiment groups, studying with the Error Recognition Model, provided better results. In both cases we have found that the ERM model applied in the teaching-learning process gives at least a two year-advantage compared to the traditional, interface-focused, low-mathability approaches.

In the category of syntactic errors, we have revealed a different pattern. Most of the items in the syntactic category were grammatical errors, i.e. knowledge which is, officially, acquired in native language classes. In the process of handling text-based documents knowledge built up in another school subject has to be transferred to the digital environment. As mentioned above, there was significant difference between the 9E and 9C groups. Considering these results, it seems that the 9C group compensates for their lack of knowledge with syntax, in other word, they primarily focus on this error category.

However, at this point we must call attention to the low percentage of students recognizing grammatical errors, mainly spelling errors. This knowledge comes from native language classes, where students are supposed to write without grammatical errors at this age. We can conclude that the knowledge transfer rate between of native language and informatics subjects is extremely low. Consequently, we must develop a higher level of cooperation between the teachers of the two subjects, so that students can apply their grammatical knowledge in digital environments.

These findings prove hypothesis [H3], but we must evaluate the results by error categories. In the typographic and layout-breaking errors the ERM model is clearly more effective than the traditional method. However, the compensation in the syntax category, despite its low scores, resulted in higher scores in group 9C.

In general, the ERM group recognize more types of errors, but these new errors distract them from the syntactical errors. It seems that the students' working memory is flooded and they need more schemata to be able to handle all the errors effectively.

The comparison of the grade 7 and 9 experiment groups revealed that in typographic and layout-breaking categories, the two groups scored similarly, with no significant difference between them. Considering the syntactical errors, the older students' results were significantly higher. This proves that the capacity of the working memory of the younger students is less than that of the older students, who might bring in schemata built up in previous studies, either in informatics or native language (Hungarian) classes. Consequently, hypothesis [H4] is rejected when the ERM method is used. The results of the pre-test prove that the selection of method is crucial, because hypothesis [H2] was only partially accepted.

In hypothesis [H5], we assumed that there is difference between groups 7E and 9C. Our analysis proved that apart from the syntax compensation, the younger students scored higher in both typography and layout-breaking errors, with significant and no significant differences, respectively. In the comparison of hypotheses [H3] and [H5] the ERM is more effective than the traditional method, even overriding the age differences.

5.3 The Comparison of the Results of the Pre- and Post-Tests

In Table 1 the number of students participating in both tests, and the way their numbers are distributed between the experiment and control groups and age groups are presented. We worked with, and used for comparison, the results expressed in percentage format due to the different number of items in the test (Table 4).

Table 7
The comparison of the results of the pre- and post-tests

	7E		9E		9C	
	Pre-test	Post-test	Pre-test	Post-test	Pre-test	Post-test
Syntactic	41.12%	19.75%	44.87%	27.02%	42.11%	36.27%
Typographic	0.00%	67.95%	0.00%	60.61%	0.00%	0.88%
Layout	4.14%	41.73%	30.65%	49.32%	12.55%	36.58%
TOTAL	19.62%	39.68%	32.73%	45.23%	23.68%	34.74%

In the pre- and post-tests both the experiment and control groups improved in recognizing and marking two of the error categories: typographic and layout-breaking errors (Table 7, Figure 2 and Figure 3). Considering syntactic errors, all the groups' results were lower in the post-test than in the pre-test. However, the difference in the 7E ($p=0.000$) and 9E ($p=0.000$) groups was significant, while in the 9C group it was not ($p=0.102$). This results also proved that the control group focuses only on syntax errors, leaving no place for typography in the unplugged phase.

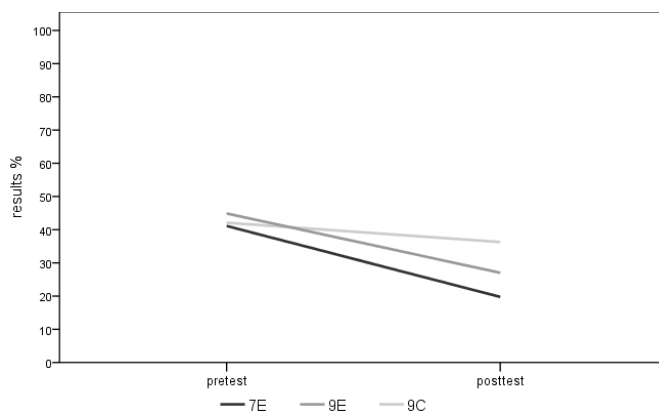


Figure 1

The results in the pre- and post-test, regarding syntax errors

The comparison of the pre- and post-test proved that the recognition of syntactic errors greatly depends on the knowledge built up in native language classes and brought into informatics. The short text of the pre-test, with a relative frequency of 43.33% for syntactical errors, suited 7E students better than the longer text of the post-test. One further explanation would be that while in the pre-test the students focused primarily on syntactical errors, in the post-test they divided their attention between the different error groups. However, this assumption requires further research. In general, we can conclude that the time spent on text management in informatics classes is not enough to transfer grammatical knowledge from other classes. One solution would be that language classes apply digital tools for integrating language knowledge into digital texts.

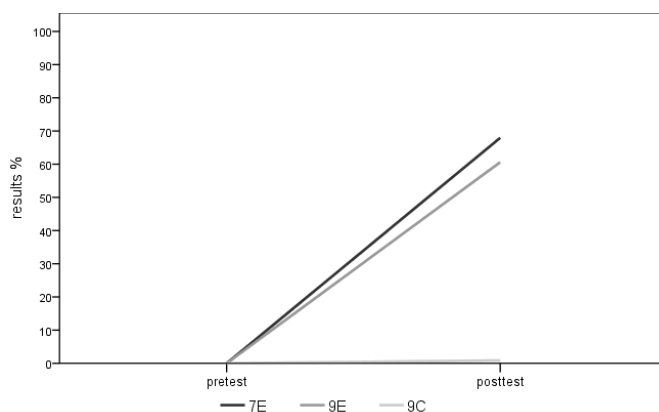


Figure 2

The results in the pre- and post-test, regarding typographic errors

In the typographic error category the experiment groups underwent a strong and significant improvement (7E, 9E: $p=0.000$), while the control group stagnated and produced similar results in both tests with no significant difference between the two tests ($p=0.160$) (Figure 2).

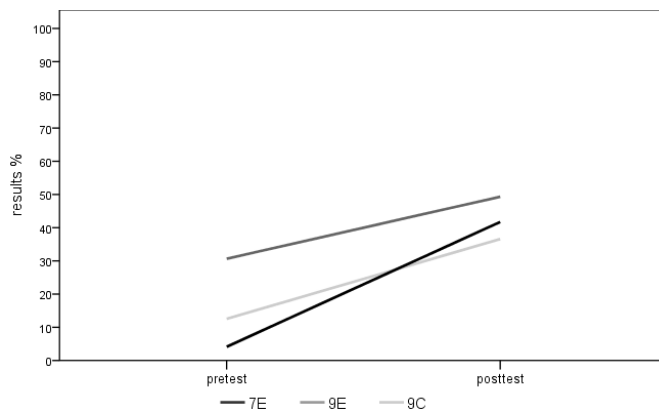


Figure 3

The results in the pre- and post-test, regarding layout-breaking errors

All three groups show a significant improvement in the layout-breaking error type (7E, 9E: $p=0.000$, 9C: $p=0.000$). Not only was the development between the pre- and post-tests significant, but the rate of improvement between the participating groups (experiment and control) was also noteworthy ($p=0.000$) (Figure 3). In text management, recognizing and naming errors all students showed an improvement; however, students working with the ERM achieved significantly better results compared to the students who learned with traditional approaches.

5.4 The Rate of Development

We were interested to see how students improved during this special teaching-learning period. Considering the syntax error group, it was found that all the groups' results were lower in the post-test than in the pre-test. On the other hand, the recognition of the layout-breaking errors improved in all the three groups, as did the recognition of the typographic errors in the two experiment groups. It seems that simultaneously, with the development of the students, a switch in focus is recognizable; while in the pre-test their knowledge was restricted to syntax, in the post-test their knowledge space was widened. In the experiment groups the two other error categories are taken care of, while in the control group only one category is, which explains the relatively better syntax results of the control group.

In syntax, the greatest drop is seen in the 7E group. There is no significant difference between the two experiment groups in terms of this drop ($p=0.545$), although there is a significant difference between the experiment and control groups: groups 7E and 9C ($p=0.009$) and groups 9E and 9C ($p=0.016$) (Table 7). Figure 3 clearly shows

these results with the corresponding parallels for the two experiment groups. Again, proof was found for the compensation for the lack of other knowledge in group 9C.

The development of the three groups in terms of typography showed different patterns. In this respect, the development of the 9C group is significantly lower than that of the other two groups ($p=0.000$, $p=0.000$). However, there is no significant difference between the development of the two experiment groups ($p=0.264$). In general, the greatest development was registered in the 7E group.

The pattern of development of the layout-breaking errors is different from the other two error categories. It was found that there was no difference between groups 9E and 9C in the development of this error type ($p=0.173$), while there is significant difference between the two experiment groups ($p=0.000$), and groups 7E and 9C ($p=0.003$). This result is also clearly shown in Figure 4 which shows the parallel results for the two grade 9 groups.

6 Summary

The results clearly show that either the ERM or the traditional method do not have a direct influence on the development of the ability to recognize syntactical errors. Furthermore, we have found that the perspective of the experiment groups widens; consequently, they can divide their attention between the three error categories, while the control group can only handle two of the categories. In this case, the 9C group in the unplugged phase compensates for the typographic errors with the syntax. The greatest development was recorded in the typographic errors in the experiment groups. In the layout-breaking errors, all the groups developed significantly.

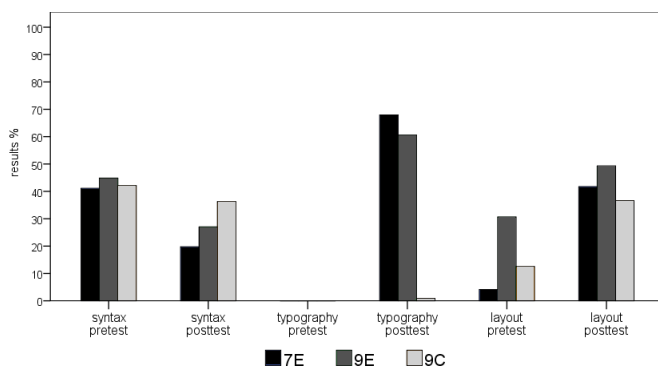


Figure 4

The results of the pre- and post-tests in groups 7E, 9E, and 9C

In general, (Figure 6), the 9E group's result was the highest, but the 7E group's development was the greatest. Here, based on our results, we can conclude that both subjects – recognition of typography and layout-breaking errors – can be taught in middle schools as effectively as in high schools. We would suggest starting these subjects as soon as possible, as students are ready for, and receptive to, this knowledge.

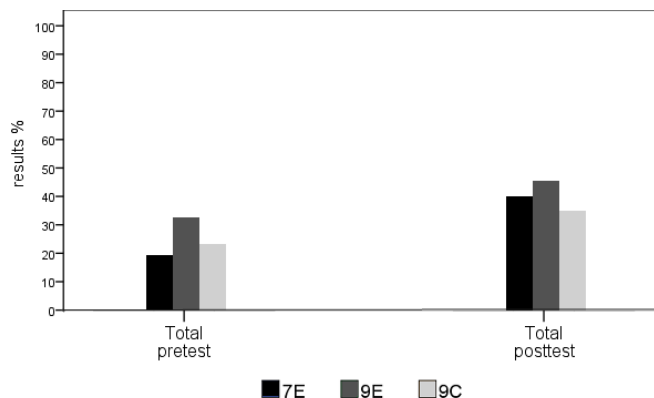


Figure 5

The total results of the pre- and post-tests in groups 7E, 9E, and 9C

Conclusions

In the present research, the effectiveness of two teaching methods, connected to word processing and digital text management, were analyzed. The traditional, widely accepted and supported surface approach methods, focusing on the features of word processors were compared to the newly introduced Error Recognition Model (ERM).

We tested three groups of students – grade 7 experiment (7E), grade 9 experiment (9E), and grade 9 control (9C) – with a pre- and a post-test.

We found that the Error Recognition Model, introduced in the experiment groups proved to be more effective in digital text management than the traditional methodologies in two error types: the typographic and layout-breaking error categories, while a strong compensation effect was found in the syntax error category. The compensation in this context means that if students are familiar with several error groups, they divide their attention among them. However, the fewer error categories they know, the more their primary focus is on syntactical errors. This finding was proved twice in our experiment: (1) In the pre-test, when the typographic and layout-breaking errors scored low, but the syntactical errors relatively high. This result was reversed in the experiment group in the post-test. (2) The control group could not improve significantly in the recognition of either the typographic or layout-breaking errors, but only in the detection of the syntactical errors.

In our opinion, this compensation effect is strongly related to the capacity of working memory and the schemata built up in previous studies. However, this latter finding requires more research and comparison.

The method based on the Error Recognition Model is a complex, time-consuming process, which assumes cooperative work from teachers and periodical assessment from both teachers and students, in order to reach a state which fulfills the requirements of digital text documents.

Acknowledgements

This work was supported by the construction EFOP-3.6.3-VEKOP-16-2017-00002. The project was supported by the European Union, co-financed by the European Social Fund.

References

- [1] Gould, J., Drongowski, P.: An exploratory study of computer program debugging. *Human Factors*, 16, 1974, pp. 258-277
- [2] Gould, J.: Some psychological evidence on how people debug computer programs. *International Journal of Man-Machine Studies*, (7) 1, 1974, pp. 151-182
- [3] Jerinic, L.: Teaching Introductory Programming Agent-based Approach with Pedagogical Patterns for Learning by Mistake. (IJACSA) *International Journal of Advanced Computer Science and Applications*, 2014
- [4] Chan Mow, I. T.: Analyses of Student Programming Errors In Java Programming Courses. *Journal of Emerging Trends in Computing and Information Sciences*, 2012, (3)5
- [5] Ben-Ari, M.: Bricolage Forever! PPIG 1999. 11th Annual Workshop. 5-7 January 1999. Computer-Based Learning Unit, University of Leeds, UK. Retrieved: 01/03/2020 from http://www.ppig.org/sites/ppig.org/files/1999-PPIG-11th-benari_0.pdf
- [6] Ben-Ari, M., Yeshno, T.: Conceptual models of software artifacts. *Interacting with Computers*, 18(6), 2006, pp. 1336-1350
- [7] Biró, P., Csernoch, M.: The mathability of computer problem solving approaches. 6th CogInfoCom, Győr, 2015, pp. 111-114
- [8] Csernoch, M.: Methodological Questions of Teaching Word Processing. 3rd International Conference on Applied Informatics, Eger-Noszvaj, Hungary, 1997, pp. 375-382
- [9] Csernoch, M.: Teaching word processing – the theory behind. *Teaching Mathematics and Computer Science*, 2009, 2009/1, pp. 119-137
- [10] Csernoch, M.: Teaching word processing – the practice. *Teaching Mathematics and Computer Science*, (8)2, 2010, pp. 247-262

- [11] Csernoch, M.: Clearing Up Misconceptions About Teaching Text Editing. In: I Candel Torres L Gómez Chova A López Martínez (ed.) ICERI2011: 4th International Conference of Education, Research and Innovation. Madrid, Spain, (IATED), 2011, pp. 407-415
- [12] Csernoch, M.: Do You Speak and Write in Informatics? (2019) The 10th International Multi-Conference on Complexity, Informatics and Cybernetics, March 12-15, 2019, Orlando, Florida, USA, 2019, pp. 147-152
- [13] Csernoch, M., Bujdosó, Gy.: Errors of exams and competitions in informatics and their consequences. In Hungarian: Vizsga- és versenyfeladatok szövegbeviteli hibái és ezek következményei, Új Pedagógia Szemle 1, 2009, pp. 19-40
- [14] Virágvölgyi, P.: The mastery of typography with computers. In Hungarian: A tipográfia mestersége számítógéppel, Osiris Kiadó Kft., 2004
- [15] Jury, D.: About Face: Reviving The Rules Of Typography, RotoVision SA, Switzerland, 2004
- [16] Jury, D.: What is Typography?, RotoVision SA, Switzerland, 2006
- [17] Reynolds, G.: Presentation Zen: Simple Ideas on Presentation Design and Delivery, Pearson Education Inc., 2008
- [18] Bujdosó, Gy., Csernoch, M.: Digital literacy, digital grammar. In Hungarian: Digitális írástudás, digitális nyelvhelyesség. Tudományos és Műszaki Tájékoztatás, 61(10), 2014, pp. 1-10
- [19] Kruger, J., Dunning, D.: Unskilled and Unaware of It: How Difficulties in Recognizing One's Own Incompetence Lead to Inflated Self-Assessments. *Journal of Personality and Social Psychology* 77(6), 1999, pp. 1121-34
- [20] Bell, T., Newton, H.: Unplugging Computer Science. Improving Computer Science Education. (Eds.) Kadijevich, D. M. Angeli, C. and Schulte, C., Routledge, 2013
- [21] Gove, M.: Michael Gove speech at the BETT Show 2012. Published 13 January 2012. Digital literacy campaign. Retrieved 12/02/2020 from <http://www.theguardian.com/education/2012/jan/11/digital-literacy-michael-gove-speech>
- [22] Panko, R., Port, D.: End User Computing: The Dark Matter (and Dark Energy) of Corporate It. *Journal of Organizational and End User Computing*, 25 (3), 2013, pp. 1-19
- [23] Chen, J. A., Morris, D. B., Mansour, N.: Science Teachers' Beliefs. Perceptions of Efficacy and the Nature of Scientific Knowledge and Knowing. In *International Handbook of Research on Teachers' Beliefs*. (Eds.) Fives, H. & Gill, M. G. Routledge, 2015

- [24] Pólya, G.: How To Solve It. A New Aspect of Mathematical Method. Second edition (1957) Princeton University Press, Princeton, New Jersey, 1954
- [25] IEEE&ACM Report 2013: Computer Science Curricula 2013. Curriculum Guidelines for Undergraduate Degree Programs in Computer Science. December 20, 2013. The Joint Task Force on Computing Curricula Association for Computing Machinery (ACM) IEEE Computer Society. Retrieved 25/05/2019 https://www.acm.org/binaries/content/assets/education/cs2013_web_final.pdf
- [26] Csernoch, M., Biró, P.: Teaching methods are erroneous: approaches which lead to erroneous end-user computing. EuSpRIG2016. London. Retrieved 02/03/2020 from <http://www.eusprig.org/mcsernoch-2016.pdf>
- [27] Csernoch, M.: Algorithms and Schemata in Teaching Informatics. In Hungarian: Algoritmusok és sémák az informatika oktatásában II., 2016 Retrieved 25/01/2020 from http://tanarkepzes.unideb.hu/szaktarnet/kiadvanyok/algoritmusok_es_semak_2.pdf
- [28] Csernoch, M., Biró, P.: Wasting Human and Computer Resources. International Journal of Social, Education, Economics and Management Engineering, 9(2), 2015, pp. 573-581
- [29] Baranyi, P., Gilányi, A.: Mathability: Emulating and Enhancing Human Mathematical Capabilities. 4th CogInfoCom, 2013, pp. 555-558
- [30] Baranyi, P., Csapo, A., Sallai, G.: Cognitive Infocommunications (CogInfoCom), Springer International Publishing Switzerland, 2015, p. 191, (978-3-319-19607-7, <http://www.springer.com/us/book/9783319196077#aboutBook>)
- [31] Baranyi, P., Csapo, A.: Definition and Synergies of Cognitive Infocommunication, Acta Polytechnica Hungarica, 2012, Vol. 9, No. 1, pp. 67-83 (ISSN 1785-8860)
- [32] Soloway, E.: Should we teach students to program? Communications of the ACM, 1993, 35(10) pp. 21-24
- [33] Chmielewska, K., Gilányi, A.: Computer assisted activating methods in education. 10th CogInfoCom, Nápoly, 2019, pp. 241-246
- [34] Chmielewska K., Gilányi, A.: Mathability and computer aided mathematical education. 6th CogInfoCom, Győr, 2015
- [35] Chmielewska, K., Gilányi, A., Łukasiewicz, A.: Mathability and Mathematical Cognition. 7th CogInfoCom, Wrocław, 2016
- [36] Chmielewska, K., Matuszak, D.: Mathability and coaching. 8th CogInfoCom, Debrecen, 2017
- [37] Csernoch, M., Dani, E.: Data-structure validator: an application of the HY-DE model, 8th CogInfoCom, Debrecen 2018, pp. 197-202

- [38] Hubwieser, P.: Functional Modeling in Secondary Schools using Spreadsheets. In *Education and Information Technologies*, 9(2), 2004, pp. 175-183
- [39] Angeli, C.: Teaching Spreadsheets: A TPACK Perspective. *Improving Computer Science Education*. (Eds.) Djordje M. Kadijevich, Charoula Angeli, and Carsten Schulte. Routledge, 2013, pp. 132-145
- [40] Shulman, L. S.: Those Who Understand: Knowledge Growth in Teaching. *Educational Researcher*, 15 (2), 1986, pp. 4-14
- [41] Shulman, L. S.: Knowledge and Teaching. *Foundations of the New Reform*. *Harvard Educational Review*, 57, 1987, pp. 1-22
- [42] Schulte, C., Saile, M.: Applying Standards to Computer Science Education. (Eds.) Djordje M. Kadijevich, Charoula Angeli, and Carsten Schulte. Routledge, 2013, pp. 117-131
- [43] Kadijevich, D. M.: Learning about Spreadsheets. (Eds.) Djordje M. Kadijevich, Charoula Angeli, and Carsten Schulte. Routledge, 2013, pp. 19-33
- [44] NAT 2012: National Base Curriculum. In Hungarian 110/2012. (VI. 4.) Korm. rendelete a Nemzeti alaptanterv kiadásáról, bevezetéséről és alkalmazásáról. Retrieved 12/02/2020 from http://ofi.hu/sites/default/files/attachments/mk_nat_20121.pdf
- [45] OFI: Frame Curricula. In Hungarian: Kerettanterv. 51/2012. (XII. 21.) számú EMMI rendelet – a kerettantervek kiadásának és jóváhagyásának rendjéről. 2012, Retrieved 12/01/2020 from https://www.oktatas.hu/kozneveles/kerettantervek/2012_nat