

USING LEARNER CORPORA TO REDESIGN UNIVERSITY-LEVEL EFL GRAMMAR EDUCATION

MICHAEL O'DONNELL*

Universidad Autónoma de Madrid

ABSTRACT. *This paper outlines the developing work in the TREACLE project, which is using learner corpora to inform the redesign of English grammar curricula in Spanish University contexts. The paper outlines the two components of the annotation: manual error analysis and automatic syntactic analysis, which together provide information as to what syntactic structures require attention at each proficiency level, and with what degree of attention. The degree of usage of a syntactic feature compared to native usage is often used to judge the criticality of the syntactic feature for learners at each proficiency level, but we argue for an alternative metric: onset of use, which measures how many of the learners at each level use the feature at all. This measure provides a clearer measure of how critical the feature is to the particular group. We finish the paper with proposed extension of the project to complement classroom teaching with intelligent online learning informed by the learner corpora.*

KEYWORDS. *Learner corpora, error analysis, curriculum design, blended learning.*

RESUMEN. *Este artículo presenta el trabajo que se está realizando en el proyecto TREACLE, que utiliza un corpus de aprendices para informar el diseño curricular de gramática inglesa en el contexto de universidades españolas. En este artículo se describen los dos componentes de la anotación: análisis manual de errores y análisis sintáctico automático, que, juntos, proporcionan información sobre qué estructuras sintácticas requieren atención, y cuánta atención, en cada nivel de competencia. A menudo se utiliza la frecuencia de uso de una estructura sintáctica en comparación con el uso nativo para determinar hasta qué punto esa estructura es crítica en el nivel de competencia de los aprendices. Sin embargo, aquí mantenemos que este enfoque presenta deficiencias. En su lugar, se propone una medida que llamamos inicio de uso, que mide cuántos de los estudiantes de cada nivel utilizan esa estructura sintáctica en algún momento. Se argumenta que el inicio de uso constituye una medida más clara de la importancia de esa estructura para un grupo determinado de aprendices. Por último, proponemos una extensión del proyecto para complementar la enseñanza en el aula con un sistema inteligente de aprendizaje en línea informado por corpus de aprendices.*

PALABRAS CLAVE. *Corpus de aprendices, análisis de errores, diseño curricular, B-learning.*

1. INTRODUCTION

The way we teach grammar to learners of English has evolved over decades of experience: teachers providing feedback to developers of teaching materials and curriculum designers, or those developers using their own teaching experiences to improve the materials and curricula. However, most of the material for teaching English has evolved to be applicable to learners from any particular language background, equally applicable to students whose mother tongue (L1) is Spanish, German, Chinese or Hungarian (e.g., any of the “first certificate” or “advanced certificate” textbooks that abound).

This practice ignores the fact that the mother tongue greatly influences the ease or difficulty of acquiring skills in particular areas of English grammar. Germanic and Romance languages for instance share many grammatical correspondences with English (e.g., perfect and progressive aspects, Subject-Finite agreement. etc.), which should make the acquisition of these features fairly straightforward. Other aspects of English grammar may not correspond to the learner’s L1, and we can expect these aspects to be more difficult to acquire, for instance, the great difficulty learners whose L1 is not Germanic tend to have with English phrasal and prepositional verbs.

Close study of learner language offers a valid and verifiable alternative to teacher intuitions. Over the last 3 years, members of the TREACLE investigation group (based at the Universidad Autónoma de Madrid and Universitat Politècnica de València) have been pursuing this approach, analysing a large corpus of learner English produced by Spanish university students, to discover what vocabulary and grammatical features are most critical to learners at each level of proficiency, with the goal of informing a reformulation of the grammar teaching curriculum in Spanish University contexts. This paper will report on this work.

1.1. *Issues in Curriculum Design for EFL Grammar Teaching*

The higher level issues that need to be addressed in designing a grammar teaching curriculum include:

1. *What to teach?* Which grammatical structures should be taught, and how much attention given to each?
2. *When to teach?* How should topics be distributed over a course? Over a degree?
3. *How to teach?*
 - Should grammatical terminology be explicitly taught, or should the teaching take place without dependence of grammatical terminology?
 - Should the curriculum be shaped by a grammar framework (e.g., systematically working through aspects of clause grammar), or should grammar teaching be opportunistically distributed through a curriculum based on other criteria (e.g., as is done with situation-centred teaching approaches)?
 - Is teaching to be purely classroom-based, purely online, or some mix of the two (i.e., blended learning)?

In this paper, we will mainly address the first two issues, exploring learner data to see what needs to be taught to learners, and in what order.

Of the issues involved in *how to teach*, this paper will not address the first two. We believe that the findings of this research are equally applicable to any teaching methodology, whether grammar is explicitly taught or is left implicit, only visible in the choice of textual examples and exercises provided to the learner. And similarly, the results are applicable where a grammatical framework directs the curriculum (as in many university-level language courses), or where grammar and vocabulary are distributed across a curriculum shaped by a sequence of situations that are considered important to learners. In all these cases, grammar is still taught, whether it shapes the framework or not, and whether explicit grammar terminology is taught or not.

We will, however, in the final section, address issues of how our grammatical profiles could be used to facilitate a blended learning approach to grammar teaching.

1.2. The TREACLE project

The TREACLE project is a cooperation of English teachers at the Universidad Autónoma de Madrid (UAM) and the Universitat Politècnica de Valencia (UPV), interested in exploring learner corpora to better understand how their students learn English. TREACLE stands for *Teaching Resource Extraction from an Annotated Corpus of Learner English*. For more information on the project, see: <http://www.uam.es/treacle>.

We started working together in early 2009, and were awarded a national project funded by the Ministerio de Ciencia e Innovación, titled “Developing an annotated corpus of learner English for pedagogical applications” (FFI2009-14436/FILO), running from January 2010 to December 2012. The goals of the project are:

- To use learner corpora to produce profiles of grammatical competence at each proficiency level (we measure proficiency in terms of the 6 CEFR levels (Council of Europe 2001): A1, A2, B1, B2, C1 and C2).
- To use these profiles to redesign the teaching curriculum: determining which grammatical features need to be taught, in what order, and with what degree of emphasis.
- Extract teaching examples and exercises from the corpus.
- Provide a web-based language learning system which dynamically adapts exercises presented to the student by reference to the student’s current performance and the proficiency profiles derived above.

2. BUILDING PROFICIENCY PROFILES

The basis of the TREACLE approach is to analyse the language of learner essays so as to measure the proficiency of both individual learners, and of the learners of a proficiency level grouped together.

2.1. *Studying learner corpora*

Various approaches exist for the analysis of learner language. One of the most discussed is Error Analysis (Corder 1967), which has been used to explore the grammatical needs of students at each level (e.g. James 1998; Dagneaux et al. 1998). By finding systematic explanations behind errors, the researchers hope to better understand the process of learning a language, and to distinguish errors which are part of the general developmental process from those which derive from linguistic traits of the mother tongue.

However, Error Analysis may not by itself constitute a sound tool for assessing the proficiency of learners:

- Conservative learners make few errors, because they avoid structures they are not sure about;
- Adventurous learners take risks with more complex structures, and thus make more errors.

To overcome this limitation of Error Analysis as a means of assessing proficiency, we need to pay attention not only to what students do wrong, but also to what they are doing correctly. One approach which moves towards this end is called *Contrastive Interlanguage Analysis* (CIA) (Granger 1998:12), which aims to explore the *interlanguage* of learners, based on the hypothesis that the language produced by a learner has a grammar of its own, and thus the errors are systematic in relation to this interlanguage (Selinker 1972). In this approach, the differences between the interlanguage (IL) and the language being learned (L2) are charted, often attempting to explain these differences in terms of the influence of the mother tongue (L1). Such studies do not focus just on errors, but rather on the syntactic structures or word choices used, comparing the frequency of use in learners compared to native producers (e.g., Biber and Reppen 1998 on complement clauses; Aijmer 2002 on modal words; Römer 2005 on progressive forms, etc.).

2.2. *Towards learner profiling*

Our particular interest is in using the learner corpus for curriculum design. There have been uses of *native* corpora to inform pedagogical design (Biber *et al.* 1994; Grabowski & Mindt 1995). However, the application of *learner* corpora to pedagogical design is much rarer. Granger (1999) explores verb tense errors in high proficiency learners, and concludes that this can lead to more targeted teaching of this area for this proficiency level. Work charting the use of syntactic structures in relation to proficiency levels includes that of Díez Bedmar (2010), who compares the use of the article system in upper secondary and lower tertiary learners of English. Note however that these works either explore single structures (or topics) over several proficiency points, or look at a single proficiency point.

More ambitious studies target a wider range of syntactic structures at a number of distinct proficiency levels. However, most of these studies seem to get tied up on the construction of the corpus, and never reach the point of pedagogical application. For instance, the good intentions of Muehleisen (2006: 119) are clear when she says: “The corpus is being created to better understand the state of students’ writing as they enter SILS and as it develops through the course of their first few semesters. The corpus will be immediately useful for the SILS language program developers in creating course material for the writing classes”. However, the paper makes clear that this work had not been attempted at that point. Rankin (2010) also considers applications to the curriculum, proposing to compare the kinds of adverb errors found in student texts to those taught in the course, with the objective of including material for common errors where they are not already covered. However, this is currently just a proposal. As Meunier (2002: 123) said, “the actual implementation of corpus research results in curriculum design is timid, if not absent”.

In recent years, more attention has been turned to this issue. Of particular interest is the work of English Profile, a research group based in the U.K. The group aims “to provide a detailed set of Reference Level Descriptions for English. Linked to the Common European Framework of Reference for Languages (CEFR), these will provide specific criteria for describing what a learner knows at a particular level of English” (English Profile, 2012).

This group, particularly in the work of Hawkins and Buttery (e.g., Hawkins and Buttery 2009, 2010) has been increasingly exploring the use of learner corpora to chart grammatical development with increasing proficiency, using the notion of criterial features. Their work parallels in many ways what the TREACLE project is exploring, although our focus is particularly on the context of Spanish learners of English.

3. TREACLE APPROACH TO PROFICIENCY PROFILES

One problem with CIA is that it requires construction of a hypothetical interlanguage, using the evidence from the L1, from the native L2, and from learner texts. In the TREACLE project, we avoid the messiness of this approach by just analysing the evidence at hand: texts produced by learners. We do not enter into the issue of whether learners function with an interlanguage or not, nor what form it may take. We simply compare learner language use to that of proficient English usage. To do this, we take a two-pronged approach:

- *(Manual) Error Analysis*, to see which language features each learner is attempting, but getting wrong.
- *(Automatic) Syntactic analysis* of the corpus, to see what language features learners demonstrate, and which they do not demonstrate.

The rest of this section will outline our work in these two directions.

3.1. *The corpora*

The project uses two corpora:

- The *WriCLE* corpus (UAM) - *Written Corpus of Learner English*. 521 essays of around 1000 words each, written by Spanish learners of English at University level (about 500,000 words) (Rollinson and Mendikoetxea 2010).
- The *UPV Learner Corpus* (UPV) containing 779 essays (150,000 words) of shorter texts by English for Specific Purposes (ESP) students. (Andreu et al. 2010).

For each corpus, the Oxford Placement Test (UCLES 2001) was given within the month of writing, to estimate the learner's proficiency. Other metadata was also collected, including gender, academic year, degree, languages of parents, time spent abroad, resources used in writing, etc. All essays not by native Spanish speakers (of whichever variety) were eliminated, as we are interested in the learning needs of Spanish learners of English.

3.2. *Software*

All manual and automatic annotation of the corpus is performed using UAM CorpusTool (O'Donnell 2008; 2009). The software runs on both Windows and MacOSX, and is available for free from <http://www.wagsoft.com/CorpusTool/>.

3.3. *Error annotation*

Error annotation in TREACLE is covered in more depth in MacDonald et al. (2011). Here I will provide just a brief overview of this work in regards its role in curriculum planning.

In error analysis, each error is tagged with an error class, typically from a taxonomy of errors. There are several such schemes available, with probably the most widely used being that developed by the Centre for English Corpus Linguistics, Université Catholique de Louvain (Dagneaux et al. 1996).

However, we found that none of the existing error coding systems fitted our goals, which required the errors made by learners to be related to the grammar teaching curriculum (for the ease of communication between researchers on the project, we assumed a fairly well-known grammar framework as presented in, e.g., Greenbaum and Quirk (1990)). We thus developed an error classification scheme which organises errors in relation to grammar topics, e.g., one sub-tree for errors related to the noun phrase, another for errors in the verb phrase, etc. Within each sub-tree, errors are further organised into sub-topics, e.g., for the noun-phrase, errors related to the determiner slot, errors related to the pre-modifier slots, etc. See Figure 1.



Figure 1. *The np-error and determiner-error subtrees of the error scheme.*

Our problem with the Louvain system was that it tended to associate errors to word classes. Rather, we associate errors to the grammatical unit which provides the context for the error (phrase or clause). For example, “He runs quick” is not for us an adverb error, but rather an error at clause level, with an inappropriate filler for the Adjunct slot.

As of October 2012, we have error-coded 307 texts of some 113,000 words, and have identified 16,200 errors. These results provide some useful insights into grammar curriculum design. Figure 2 shows the percentage of each main type of grammar errors, one bar for each proficiency level.¹ The main thing to note here is that around 40% of all grammar errors occur within the NP, which suggests that a large proportion of grammar classes could concentrate on this unit. We also note that as proficiency rises, the importance of NP errors falls, while errors in clause construction increase, suggesting that issues of clause construction should be addressed later.

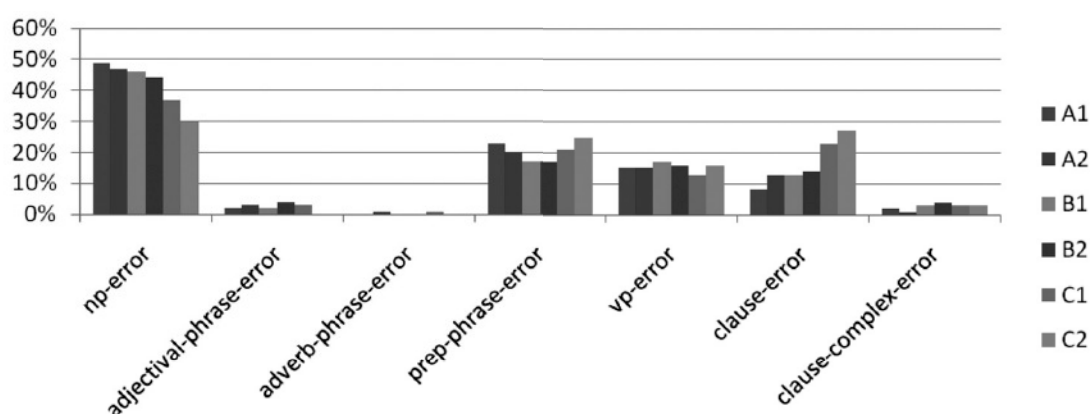


Figure 2. *Changing proportion of errors as students gain proficiency.*

3.4. Syntactic analysis

UAM CorpusTool also produces *automatic syntactic analysis* of the sentences in the text. These analyses can be used to explore which grammatical structures each

student uses in their essays. We can explore how often grammatical structures are used at each proficiency level. We can thus construct *grammatical profiles*: the degree to which each proficiency level uses each kind of structure. From these we can see when it is best to teach particular structures.

UAM CorpusTool uses the Stanford parser (Klein and Manning 2003) to parse the text, and then converts this into syntactic analyses used in our framework. Structurally, each clause is analysed in terms of Subject, Predictor, Object, etc. and each phrase is also structured. Each unit is also assigned a set of syntactic features representing the salient aspects we need to deal with. Table 1 shows the range of clause level features we extract. A later research phase will further develop the range of syntactic features for NPs, which our error annotation work has shown to warrant much attention.

TABLE 1. *Clause level features recognized by the parser.*

TENSE simple-present present-perfect present-progressive simple-past past-progressive past-progressive simple-modal modal-perfect modal-progressive	FINITENESS simple-finite finite-with-connector relative-clause that-clause wh-nominal-clause infinitive-clause pres-participle-clause past-participle-clause	VERB-TYPE intransitive-verb monotransitive-verb ditransitive-verb ergative-verb relational-verb verbal-verb mental-verb
MODALITY nonmodal-clause true-modal-clause future-clause	DO-INSERTION do-inserted no-do-inserted	POLARITY positive-polarity negative-polarity
PROCESS TYPE material-clause verbal-clause mental-clause relational-clause	VOICE active-clause passive-clause	MOOD declarative-clause imperative-clause interrogative-clause

After parsing, we have a corpus of 1300 texts, 660,000 words, 90,000 clauses, and 150,000 NPs. The next question is, given all this data, how do we use it to inform us about *what* students need to learn and *when*?

3.5. *Extracting Profiles from the Corpus*

1. *Degree of Usage (Simple Frequency) Approach*: Some researchers contrast the learner's degree of usage of a syntactic feature with the degree of usage of

natives. Where students under-use the feature, more emphasis is supposedly needed in teaching. Over-usage also needs to be corrected (perhaps by teaching alternative lexico-grammatical strategies, or teaching appropriate contexts of use). For instance, Figure 3 shows the degree of usage of the passive structure in the UAM section of our learner corpus, at each proficiency level (note this sub-corpus lacks A1 learners).⁴ Degree of usage is measured in terms of the percentage of clauses which are passive rather than active. At the A2 level (pre-intermediate), only 5% of clauses are passive, which rises to 9.5% at the C2 (advanced learner) level.

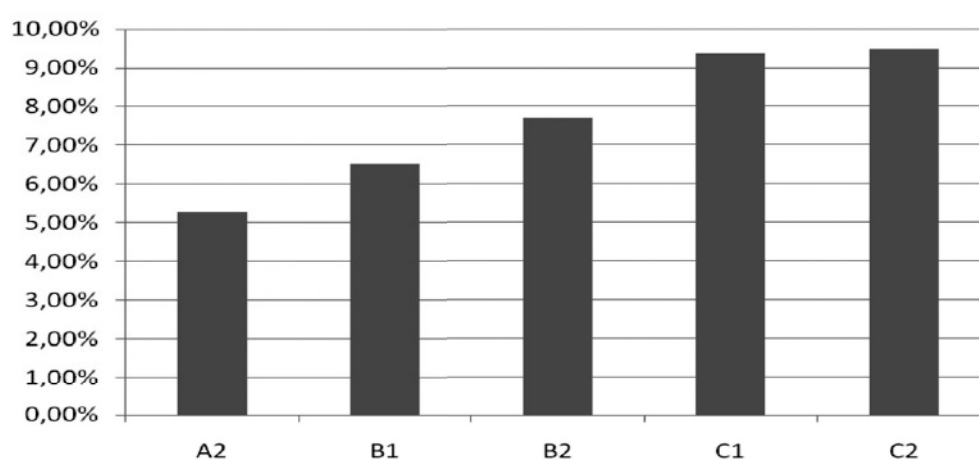


Figure 3. *Use of Passive voice increases with proficiency.*

There are however two problems with using average usage to measure proficiency:

1. The degree of usage of many features is register-dependent, so we cannot really compare with native corpus unless we have a register-matched native corpus.
2. Students in a given proficiency band are not all homogeneous: if we say that average usage of passives at a particular level is 7%, that ignores the fact that some students will over-use passives, and others will not use them at all. Grouping together the usage patterns of very different students may not give the best results.

To demonstrate the problem, Figure 4 shows the percentage of B1 learners who use the passive to different degrees. On the X axis, we have percent of usage, and on the Y axis, the percentage of learners with that usage level. For instance, the first column shows that 4.1% of B1 learners don't use a passive in any of their texts. The highest column shows that 12.5% of B1 learners use a passive in 6% of their clauses, and the final column shows that some (but very few) use a passive in 16% of their clauses.

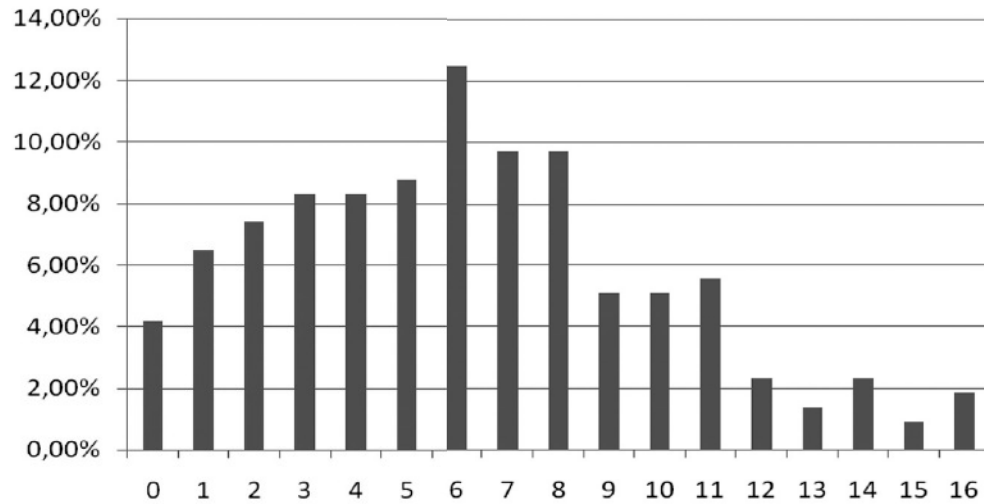


Figure 4. *Showing degree of variation of Passive use within a single proficiency band (B1).*

What this shows is that B1 learners are not homogeneous in their usage of passives. To treat this group of students as a homogeneous group which uses passive 6.4% of the time ignores the great range that exists in the group.

2. *Onset of Use Approach:* Our belief is that the first concern should be with whether a learner is capable of producing a structure at all. We thus look at each text individually, to see if the structure is present or not. We then measure the percentage of texts which use the feature at all (at each level). Figure 5 shows the percentage of texts at each proficiency level which do not use passive clauses (using only the WriCLE corpus with longer texts). While the graph is not perfectly regular, it does show that at lower levels (A2 and B1), around 3-4% of learners are still not using passives, while from B2, all learners are using at least one passive in their essay. Figure 6 shows similar results for present-participle clauses (e.g. the underlined clause in “Going to the shop, I lost my purse”).

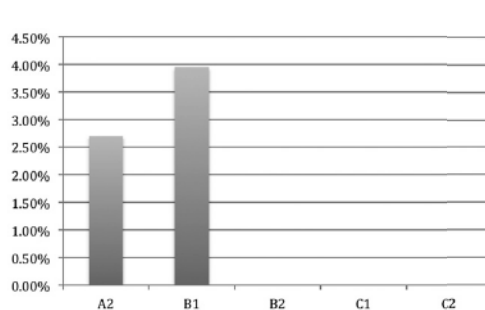


Figure 5. *Percentage of texts at each proficiency level which do not use passive clauses².*

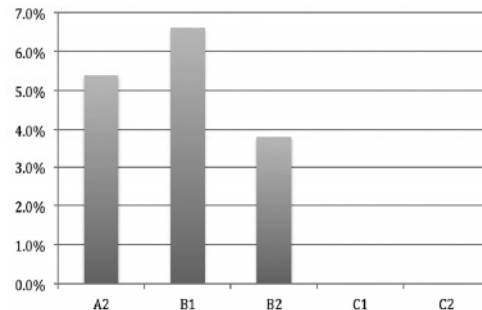


Figure 6. *Percentage of texts at each proficiency level which do not use present-participle clauses³.*

The difference between the data in Figures 5 and 6 suggest that learners develop the ability to form passive clauses before they start developing present-participle clauses.

By looking at the degree to which learners use grammatical features at each level, we can see when it is most critical to teach them that feature. We believe that it is best to teach a structure when the more advanced learners in the group are starting to use the structure correctly, while others have not yet begun to use the structure, or are using it incorrectly. If no student have started to use the structure, it may be too difficult for them to acquire it. If all of the students are using it, then the teaching is a waste of the student's time. Exactly where in this range a structure is best taught is not clear, however, some flexibility is good, to allow adjustment of materials to fit into a structured grammar teaching environment.

3.6. Limitations of our approach

Length of text: Measuring “onset of use” of a feature requires a reasonable length of text per student. For the WriCLE corpus (UAM), we have approximately 1000 words per essay. This is reasonable for structures where native use is over, say, 3% of clauses, meaning that there is a reasonable expectation that the structure will occur within the text. However, for rarer structures (e.g., clefting), longer texts, or multiple texts by same student, would be needed to accurately indicate onset of use of the feature.

Effect of Task: It is clear that the writing task set to the student will affect issues of grammatical usage. For instance, a task such as “talk about your last holiday” will involve strong use of past tenses, while “what do you want to do after university?” will be responded to with lots of future tenses, modals, and conditional clauses. We have attempted to alleviate this problem by including in the corpus a wide range of different tasks at each level, such that the affect of each task on overall results is minimal.

3.7. Future Work

So, far, we have only applied this technique to a range of clause level structures. For the full study, we need to explore the full range of structures taught in grammar courses, particularly including the aspects of noun phrase construction that our error analysis show to be critical for university level learners of English.

In future work, we also need to combine the evidence provided by our error analysis with that of the syntactic analysis into a single proposal for curriculum redesign. This work is currently being finalised, and will be reported in a future publication.

4. ADAPTIVE ONLINE LEARNING ASSISTANT

One aspect of the TREACLE project has been outlined above: using a learner corpus to order grammatical concepts within a grammar teaching program. The other aspect that the project has been addressing is how the information gained from the TREACLE approach might be used in a blended learning environment.

One of the problems of the traditional classroom is that students in a foreign language classroom are rarely all at the same level, yet the teacher needs to direct their teaching at some particular proficiency level. The net effect is that those students below the target level may fail to learn because they are not ready for the material being taught, while those with proficiency above the target may have already mastered the material, and thus become bored and lose interest.

One solution, which we have been applying at the UAM over the past few years, is to stream the students into groups based on proficiency. All students are given the Oxford Placement test in the first week of semester, and students are assigned to one of four groups based on their test results. While each group receives basically the same material, the teachers can provide more attention to the basic material for the lower groups, and provide the higher groups more explanation to the advanced topics in the material. However, this approach will not be possible in many EFL teaching contexts.

A more general solution involves the use of blended learning: complementing traditional classroom teaching with out-of-class activities (cf. Singh 2003). Applying this approach to the current context:

- Teaching in each class is targeted at the median point.
- Out-of-class activities are assigned for each student, targeted at their particular weaknesses and strengths.

Out-of-class activities can include both traditional paper-based activities, and online learning systems:

- *Traditional Paper-based activities:* in the UAM, after completing the placement test at the start of semester, a computer program generates a report for each student, outlining their areas of weakness, and providing references into study materials which they are recommended to work through. We find that, as this report is customised for the individual rather than a generic recommendation, the student is more likely to follow through with the recommended work.
- *Computer-based activities:* many computer-assisted language learning (CALL) systems allow the student to choose a level of difficulty appropriate to their current learning needs. This means that students who are below the target level of the class can study the concepts they need to catch up with the rest of the class, while those above the target level can proceed in their learning at their own pace.

CALL activities can rewardingly include collaboration between learners over the internet, e.g., Computer-Supported Collaborative Learning (CSCL) or Computer-Mediated Communication (CMC) Exchanges. In the TREACLE project however we are focusing on building an intelligent learning environment for individual learners.

In terms of a learning environment, we are constructing a system that integrates descriptions of areas of grammar with exercises to measure the degree to which the student understands the material (as happens with most CALL systems).

We however believe such systems need to be intelligent:

1. The system needs to be able to construct a ‘learner model’: a representation of which features of the target language the learner has already mastered, which they are currently struggling with, and which they have not yet begun to work with.
2. This learner model should be built up from multiple sources of evidence. For instance, in our current system, an initial learner model is derived by analyzing the student’s responses to the Oxford Placement Test (grammar component), where each correct response to a question has been associated with the concept(s) that the answer indicates the student understands, and each wrong response indicates which concepts they lack. Another source of information will be student works submitted via the system, automatically parsed to recognize linguistic forms that indicate the student does or doesn’t understand some concept. For instance, appearance of “much” with a count noun in “I have much books” indicates that the student lacks the concept: ‘much goes with mass nouns’. Error correction by the teacher of the works will offer additional evidence not picked up by automatic analysis (see Wibble et al (2003) on bootstrapping learner models using teacher annotation of errors).
3. The system should offer material and exercises related to the learner’s current abilities: students learn more efficiently when confronted with material that is neither too easy nor too difficult. Exercises should be selected by the system to develop those concepts that the student model indicates the student is ready to learn but has not yet mastered. Ideally, the system should know which of these concepts are most critical for the student to learn at this point of time, e.g., those that they need to learn to proceed to the next level of proficiency (e.g., the ‘critical features’ of Hawkins and Buttery (2008)).
4. The learner model should evolve as the student works with the system. As the student gets questions right (or wrong) the system should update its learner model.

We are currently working on such an online quiz system, with a grant from the Universidad Autónoma de Madrid (FyL-L2-7). The system is still in an early phase, although we plan to trail the system with our first year English students in 2012-13, with further trails in other courses the year after. Our plan is to initially restrict the system to noun phrase concepts, as this is the grammatical area with the most errors for elementary language learners.

5. CONCLUSIONS

This paper has presented the ongoing work of the TREACLE project, with particular emphasis on our work towards using a learner corpus to inform curriculum redesign for Spanish learners of English at University level.

We follow a two-pronged approach: error analysis to see what learners do wrong, and syntactic analysis to see what they do right. Both forms of annotation of the corpus, when projected over the various proficiency levels, provide insight as to when particular syntactic structures should be taught, and to how much attention each area should be given. We concluded that it is better to look at when individual students start to use a structure rather than to look at the usage of a group of students as a whole.

The information gained from our error- and syntactically-annotated corpus will also be input to our online learning system, with the questions being asked of the learner being tailored to the learner's particular needs.

NOTES

* Correspondence to: Mick O'Donnell. C/ Aurora 3, 5º. 28760. Tres Cantos. Madrid. E-mail: michael.odonnell@uam.es.

1 Chi-square = 86.9 with 30 DF, which is well above the critical value for $P=0.001$: 59.70.

2 Chi-square = 11.5 with 4 DF, producing a p-value of 0.021.

3 Chi-square = 8.42 with 4 DF, producing a p-value of 0.077, which is not significant at a 5% level, suggesting more data may be needed to verify this pattern.

4 Chi-square = 123.0 with 4 DF, which is well above the critical value for $P=0.001$: 18.47.

REFERENCES

- Andreu, M., A. Astor, M. Boquera, P. MacDonald, B. Montero and C. Pérez. 2010. "Analysing EFL learner output in the MiLC project: An error it's*, but which tag?" *Corpus-Based Approaches to English Language Teaching*. Eds. M. C. Campoy, B. Belles-Fortuno and M. L. Gea-Valor. London: Continuum. 167-179.
- Aijmer, K. 2002. "Modality in advanced Swedish learners' written interlanguage". *Computer Learner Corpora, Second Language Acquisition and Foreign Language Teaching*. Eds. S. Granger, J. Hung and S. Petch-Tyson. Amsterdam & Philadelphia: Benjamins. 55-76.
- Biber, D., S. Conrad and R. Reppen 1994. "Corpus-Based Approaches and Issues in Applied Linguistics". *Applied Linguistics* 15: 169-189.
- Biber, D. and R. Reppen 1998. "Comparing native and learner perspectives on English grammar: A study of complement clauses". *Learner English on Computer*. Ed. S. Granger. London: Addison Wesley Longman. 145-158.
- Corder, S. P. 1967. "The Significance of Learners' Errors". *International Review of Applied Linguistics* 5 (4): 161-170.
- Council of Europe. 2001. *The Common European Framework of Reference for Languages: Learning, Teaching, Assessment*. Cambridge: Cambridge University Press.
- Dagneaux, E., S. Denness, S. Granger, and F. Meunier. 1996. *Error Tagging Manual Version 1.1*. Louvain-la-Neuve: Centre for English Corpus Linguistics, Université Catholique de Louvain.

- Dagneaux, E., S. Denness and S. Granger. 1998. "Computer-aided error analysis". *System* 26: 163–174.
- Díez-Bedmar, M. B. 2010. "From secondary school to university: The use of the English article system by Spanish learners". *Exploring corpus-based research in English language teaching*. Eds. B. Bellés-Fortuño, M. C. Campoy-Cubillo & M. L. Gea-Valor. Castelló de la Plana: Publicacions de la Universitat Jaume I. 45-55.
- English Profile. 2012. "What is English Profile?" [Available at: <http://www.englishprofile.org/>].
- Grabowski, E. and D. Mindt. 1995. "A corpus-based learning list of irregular verbs in English". *ICAME Journal* 19: 5-22.
- Granger, S. 1998. "The computer learner corpus: a versatile new source of data for SLA research". *Learner English on Computer*. Ed. S. Granger. London: Addison Wesley Longman. 3-18.
- Granger, S. 1999. "Use of tenses by advanced EFL learners: evidence from an error-tagged computer corpus". *Out of Corpora- Studies in Honour of Stig Johansson*. Eds. H. Hasselgard and S. Oksefjell. Amsterdam: Rodopi. 191-202.
- Greenbaum, S. and R. Quirk. 1990. *A Student's Grammar of the English Language*. London: Longman.
- Hawkins J. A. and P. Buttery. 2009. "Using learner language from corpora to profile levels of proficiency: Insights from the English Profile Programme". *Studies in Language Testing: The Social and Educational Impact of Language Assessment*. Cambridge: Cambridge University Press.
- Hawkins J. A. and P. Buttery. 2010. "Critical Features in Learner Corpora: Theory and Illustrations". *English Profile Journal* 1 (1): 1-23.
- James, C. 1998. *Errors in Language Learning and Use: Exploring error analysis*. Harlow: Longman.
- Klein, D. and C. Manning 2003. "Fast Exact Inference with a Factored Model for Natural Language Parsing". *Advances in Neural Information Processing Systems* 15: 3-10.
- MacDonald, P., S. Murcia, M. Boquera, A. Botella, L. Cardona, R. García, E. Mediero, M. O'Donnell, A. Robles and K. Stuart. 2011. "Error Coding in the TREACLE project". *Las tecnologías de la información y las comunicaciones: Presente y futuro en el análisis de corpora. Actas del III Congreso Internacional de Lingüística de Corpus*. Eds. M. L. Carrió Pastor and M. A. Candel Mora. Valencia: Universitat Politècnica de València. 725-740.
- Meunier, F. 2002. "The pedagogical value of native and learner corpora in EFL grammar teaching". *Computer learner corpora, second language acquisition and foreign language teaching*. Eds. S. Granger, J. Hung & S. Tyson. Amsterdam/Philadelphia, Benjamins. 119-142.
- Muehleisen, V. 2006. "Introducing the SILS Learners' Corpus: A Tool for Writing Curriculum Development". *Waseda Global Forum* 3: 119-125.

- O'Donnell, M. 2008. "Demonstration of the UAM CorpusTool for text and image annotation". *Proceedings of the ACL-08:HLT Demo Session* (Companion Volume), Columbus, Ohio: Association for Computational Linguistics. 13–16.
- O'Donnell, M. 2009. "The UAM CorpusTool: Software for corpus annotation and exploration". *Applied Linguistics Now: Understanding Language and Mind / La Lingüística Aplicada Hoy: Comprendiendo el Lenguaje y la Mente*. Eds. C. M. Bretones Callejas, et al. Almería: Universidad de Almería. 1433-1447.
- O'Donnell, M, S. Murcia, R. García, C. Molina, P. Rollinson, P. MacDonald, K. Stuart, M. Boquera. 2009. "Exploring the proficiency of English learners: The TREACLE project". *Proceedings of the Fifth Corpus Linguistics*. Liverpool.
- Rankin, T. 2010. "Advanced learner corpora data and grammar teaching: adverb placement". *Corpus-Based Approaches to English Language Teaching*. Eds. M. C. Campoy, B. Belles-Fortuno and M. L. Gea-Valor. London: Continuum. 205-215.
- Rollinson, P. and A. Mendikoetxea. 2010. "Learner corpora and second language acquisition: Introducing WriCLE". *Analizar datos: Describir variación/Analysing data: Describing variation*. Eds. J. L. Bueno Alonso, D. Gonzáliz Álvarez, U. Kirsten Torrado, A. E. Martínez Insua, J. Pérez-Guerra, E. Rama Martínez and R. Rodríguez Vázquez. Vigo: Universidade de Vigo (Servizo de Publicacións). 1-12.
- Römer, U. 2005. *Progressives, patterns, pedagogy: A corpus-driven approach to English progressive forms, functions, contexts and didactics*. Amsterdam: Benjamins.
- Selinker, L. 1972. "Interlanguage". *International Review of Applied Linguistics* 10: 209-231.
- Singh H. 2003. "Building Effective Blended Learning Programs". *Educational Technology* 43 (6): 51-54.
- UCLES 2001. *Quick Placement Test (Paper and pencil version)*. Oxford: Oxford University Press.
- Wible, D., C. Kuo, N. Tsao, A. Liu and H. Lin. 2003. "Bootstrapping in a Language Learning Environment". *Journal of Computer-Assisted Learning* 19 (1): 90-102.