



Text Simplification for Scientific Information Access: CLEF 2021 SimpleText Workshop

Liana Ermakova, Patrice Bellot, Pavel Braslavski, Jaap Kamps, Josiane Mothe, Diana Nurbakova, Irina Ovchinnikova, Eric Sanjuan

► To cite this version:

Liana Ermakova, Patrice Bellot, Pavel Braslavski, Jaap Kamps, Josiane Mothe, et al.. Text Simplification for Scientific Information Access: CLEF 2021 SimpleText Workshop. 43rd edition of the annual BCS-IRSG European Conference on Information Retrieval : Advances in Information Retrieval (ECIR 2021), Mar 2021, Lucca (virtual), Italy. hal-03121986

HAL Id: hal-03121986

<https://hal.science/hal-03121986>

Submitted on 26 Jan 2021

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Text Simplification for Scientific Information Access

CLEF 2021 SimpleText Workshop

Liana Ermakova¹, Patrice Bellot², Pavel Braslavski³, Jaap Kamps⁴,
Josiane Mothe⁵, Diana Nurbakova⁶, Irina Ovchinnikova⁷, and Eric San-Juan⁸

¹ Université de Bretagne Occidentale, HCTI - EA 4249, France
liana.ermakova@univ-brest.fr

² Aix Marseille Univ, Université de Toulon, CNRS, LIS, Marseille, France

³ Ural Federal University, Yekaterinburg, Russia

⁴ University of Amsterdam, Amsterdam, The Netherlands

⁵ Université de Toulouse, IRIT, France

⁶ Institut National des Sciences Appliquées de Lyon, Lyon, France

⁷ Sechenov University, Moscow, Russia

⁸ Avignon Université, LIA, France

Abstract. Modern information access systems hold the promise to give users direct access to key information from authoritative primary sources such as scientific literature, but non-experts tend to avoid these sources due to their complex language, internal vernacular, or lacking prior background knowledge. Text simplification approaches can remove some of these barriers, thereby avoiding that users rely on shallow information in sources prioritizing commercial or political incentives rather than the correctness and informational value. The CLEF 2021 SimpleText track will address the opportunities and challenges of text simplification approaches to improve scientific information access head-on. We aim to provide appropriate data and benchmarks, starting with pilot tasks in 2021, and create a community of NLP and IR researchers working together to resolve one of the greatest challenges of today.

Keywords: Scientific text simplification · (Multi-document) summarization · Contextualization · Background knowledge

Everything should be made as simple
as possible, but no simpler

Albert Einstein

1 Introduction

Scientific literacy, including health related questions, is important for people to make right decisions, evaluate the information quality, maintain physiological and mental health, avoid spending money on useless items. For example, the stories the individuals find credible can determine their response to the COVID-19 pandemic, including

the application of social distancing, using dangerous fake medical treatments, or hoarding. Unfortunately, stories in social media are easier for lay people to understand than the research papers. Scientific texts such as scientific publications can also be difficult to understand for non domain-experts or scientists outside the publication domain. Improving text comprehensibility and its adaptation to different audience remains an unresolved problem. Although there are some attempts to tackle the issue of text comprehensibility, they are mainly based on readability formulas, which are not convincingly demonstrated the ability to reduce the difficulty of text [26].

To put a step forward to automatically reduce difficulty of text understanding, we propose a new workshop called SimpleText which aims to create a community interested in generating simplified summaries of scientific documents. Thus, the goal of this workshop is to connect researchers from different domains, such as Natural Language Processing, Information Retrieval, Linguistics, Scientific Journalism etc. in order to work together on automatic popularisation of science.

Improving text comprehensibility and its adaptation to different audience bring societal, technical, and evaluation challenges. There is a large range of important *societal challenges* SimpleText is linked to. Open science is one of them. Making the research really open and accessible for everyone implies providing it in a form that can be readable and understandable; referring to the “comprehensibility” of the research results, making science understandable [16]. Another example of those societal challenges is offering means to develop counter-speech to fake news based on scientific results. SimpleText also tackles *technical challenges* related to data (passage) selection and summarisation, comprehensibility and readability of texts.

To face these challenges, SimpleText provides an open forum aiming at answering questions like:

- **Information selection:** Which information should be simplified (e.g., in terms document and passage selection and summarisation)?
- **Comprehensibility:** What kind of background information should be provided (e.g., which terms should be contextualized by giving a definition and/or application)? What information is the most relevant or helpful?
- **Readability:** How to improve the readability of a given short text (e.g., by reducing vocabulary and syntactic complexity) without information distortion?

We will provide data and benchmarks, and address evaluation challenges underlying the technical challenges, including:

- How to evaluate information selection?
- How to evaluate background information?
- How to measure text simplification?

2 Information Selection, Comprehensibility, Readability

In order to simplify scientific texts, one have to (1) *select the information* to be included in a simplified summary, (2) decide whether the selected information is sufficient and *comprehensible* or he/she should provide some background knowledge, (3) improve the *readability* of the text. Our tasks are based on this pipeline.

2.1 Selecting the information to be included in a simplified summary

People have to manage the constantly growing amount of information. According to several estimates the number of scientific journals is around 30,000, with about two million articles published per year [3]. About 180,000 articles on Covid-19 were published from January 2020 to October 2020 [1]. To deal with this data volume, one should have a concise overview, i.e. a summary. People prefer to read a short document instead of a long one. Thus, even single-document summarization is already a step of text simplification. Notice, that the information in a summary designed for a scientist from a specific field should be different from that adapted for general public.

Automatic summarization can simplify access to primary scientific documents – the resulting concise text is expected to highlight the most important parts of the document and thus reduces the reader’s efforts. Evaluation initiatives in the 2000s such as Document Understanding Conference (DUC) and the Summarization track at the Text Analysis Conference (TAC) have focused primarily on the automatic summarization of news in various contexts and scenarios. Scientific articles are typically provided with a short abstract written by the authors. Thus, automatic generation of an abstract for a stand-alone article does not seem to be a practical task. However, if we consider a large collection of scientific articles and citations between them, we can come to a task of producing an abstract that would contain important aspects of a paper from the perspective of the community. Such a task has been offered to the participants of the TAC 2014 Biomedical Summarization Track⁹, as well as of the CL-SciSumm shared task series. In particular, the 2020 edition of CL-SciSumm features LaySummary subtask, where a participating system must produce a text summary of a scientific paper intended for non-technical audience¹⁰ without using technical jargon. However, in most cases, the names of the objects are not replaceable in the process of text transformation or simplification due to the risk of information distortion. In this case it is important to explain these complex concepts to a reader (see Section 2.2 Comprehensibility). Another close work is CLEF-IP 2012-2013: Retrieval in the Intellectual Property Domain¹¹ (novelty search). Given a claim, the task was to retrieve relevant passages from a document collection. However, CLEF-IP focused on extractive summarization only and did not consider text simplification.

Sentence compression can be seen as a middle ground between text simplification and summarization. The task is to remove redundant or less important parts of an input sentence, preserving its grammaticality and original meaning [18]. Thus, the main challenge is to *choose which information* should be included in a simplified text.

2.2 Comprehensibility

Comprehensibility of a simple text varies for different readership. Readers of popular science texts have a basic background, are able to process logical connections and recognize novelty [24]. In the popular science text, a reader looks for rationalization and

⁹ <https://tac.nist.gov/2014/BiomedSumm/>

¹⁰ <https://ornloda.github.io/SDProc/sharedtasks.html#laysumm>

¹¹ <http://www.ifs.tuwien.ac.at/clef-ip/tasks.shtml>

clear links between well known and new [28]. To adopt the novelty, readers need to include new concepts into their mental representation of the scientific domain.

According to The Free Dictionary, *background knowledge* is “information that is essential to understanding a situation or problem” [2]. Lack of basic knowledge can become a barrier to reading comprehension [30]. In [30], the authors suggested that there is a knowledge threshold allowing reading comprehension. Background knowledge, along with content, style, location, and some other dimension, are useful for personalised learning [35]. In contrast to newspapers limited by the size of the page, digital technologies provide essentially unbounded capabilities for hosting primary-source documents and background information. However, in many cases users do not read these additional texts. It is also important to remember, that the goal is to keep the text simple and short, not to make it indefinitely long to discourage potential readers.

Entity linking (also known as Wikification) is the task of tying named entities from the text to the corresponding knowledge base items. A scientific text enriched with links to Wikipedia or Wikidata can potentially help mitigate the background knowledge problem, as these knowledge bases provide definitions, illustrations, examples, and related entities. However, the existing standard datasets for entity linking such as [23] are focused primarily on such entities as people, places, and organizations, while a lay reader of a scientific article needs rather assistance with new concepts, methods, etc. Wikification is close to the task of terminology and keyphrase extraction from scientific texts [4]. Searching for background knowledge is close to INEX/CLEF Tweet Contextualization track 2011-2014 [7] and CLEF Cultural micro-blog Contextualization 2016, 2017 Workshop [14], but SimpleText differs from them by making a focus on selection of notions to be explained and the helpfulness of the information provided rather than its relevance. The idea to contextualize news was further developed in Background Linking task at TREC 2020 News Track aiming at a list of links to the articles that a person should read next¹². In contrast to that, SimpleText try to determine terms to be contextualized. SimpleText is similar to the Wikification task at TREC 2020 News Track since it also aims to evaluate whether the critical context for understanding is missing but the types of background knowledge are different since our target is a scientific text. Besides, we will rank terms to be contextualized rather than passages.

Thus, the main challenge of the comprehensibility is to *provide relevant background knowledge* to help a reader to understand a complex scientific text.

2.3 Readability

Readability is the ease with which a reader can understand a written text. Readability is different from legibility, which measures how easily a reader can distinguish characters from each other. Readability indices have been widely used to evaluate teaching materials, news, and technical documents for about a century [45, 21]. For example, Gunning fog index, introduced in 1944, estimates the number of years in a scholar system required to understand a given text on the first reading. Similarly, the Flesch–Kincaid readability test shows the difficulty of a text in English based on word length and sentence length [19]. Although these two metrics are easy to compute, they are criticized

¹² <http://trec-news.org/guidelines-2020.pdf>

for the lack of reliability [36]. The very structure of the readability indices suggested to authors or editors how to simplify a text: organize shorter and more frequent words into short sentences. Later studies incorporate lexical, syntactic, and discourse-level features to predict text readability [33]. In NLP tasks, readability, coherence, conciseness, and grammar are usually assessed manually since it is difficult to express these parameters numerically [13]. However, several studies were carried out in the domain of automatic readability evaluation, including the application of language models [36, 10, 22, 17] and machine learning techniques [32, 17]. Traditional methods of readability evaluation are based on familiarity of terms [37, 9, 20] or their length [41] and syntax complexity (e.g. sentence length, the depth of a parse tree, omission of personal verb, rate of prepositional phrases, noun and verb groups etc.) [10, 29, 42, 46, 8]. Word complexity is usually evaluated by experts [38, 9, 20]. [6] computed average normalized number of words in valid coherent passages without syntactical errors, unresolved anaphora, and redundant information. Several researches argue also the importance of sentence ordering for text understanding [5, 15].

Automatic text simplification might be the next step after estimation of text complexity. Usually, text simplification task is performed and assessed on the level of individual sentences. To reduce the reading complexity, in [11], the authors introduced a task of sentence simplification through the use of more accessible vocabulary and sentence structure. They provided a new corpus that aligns English Wikipedia with Simple English Wikipedia and contains simplification operations such as rewording, reordering, insertion and deletion. Accurate lexical choice presupposes unambiguous reference to the particular object leading to actualization of its connections with other objects in the domain. Domain complexity concerns the number of objects and concepts in the domain, and connections among them described by the terminology system (see a survey: [25]). Names of the objects are not replaceable in the process of text transformation or simplification due to risk of information distortion [27, 12]. For example, ‘hydroxychloroquine’ represents a derivative of ‘chloroquine’, so the substances are connected thanks to belonging to a set ‘chloroquine derivatives’. However, it is impossible to substitute ‘hydroxychloroquine’ by ‘chloroquine’ while simplifying a medical text about a Covid-19 treatment because of the difference in their chemical composition. A hypernym ‘drugs’ can refer to the substances. The hypernym generalizes the information while omitting essential difference between the drugs; however, the generalization allows to avoid misinformation [40]. Science text simplification presupposes facilitation of readers’ understanding of complex content by establishing links to basic lexicon, avoiding distortion connections among objects within the domain.

Ideally, the results undergo a human evaluation, since traditional readability indices can be misleading [43]. Automatic evaluation metrics have been proposed for the task: SARI [44] targets lexical complexity, while SAMSA estimates structural complexity of a sentence [39]. Formality style transfer is a cognate task, where a system rewrites a text in a different style preserving its meaning [34]. These tasks are frequently evaluated with BLEU metrics [31] to compare system’s output against gold standard.

Thus, the main challenge of the readability improvement is to *reduce vocabulary and syntactic complexity* without information distortion while keeping the target genre.

3 Pilot tasks

To start with, we will develop three pilot tasks that will help to better understand the challenges as well to discuss these challenges and the way to evaluate solutions. Details on the tasks, guideline and call for contributions can be found at www.irit.fr/simpleText, in this paper we just briefly introduce the planned pilot tasks. Note that the pilot tasks are means to help the discussions and to develop a research community around text simplification. Contributions will not exclusively rely on the pilot tasks.

3.1 Task 1: Ranking the words/sentences to be included in a simplified summary

Participants will be provided with scientific articles. This pilot task aims at automatically deciding which passages of these scientific articles should be included in extractive summaries in order to get a simplified summary of the initial texts. Note, that the information in a summary designed for an expert should be different from those for the general audience. To evaluate these results, we will rely on manual annotation and automatic metrics.

3.2 Task 2: Searching for background knowledge

The goal of this pilot task is to provide relevant background knowledge to help a reader to understand a complex scientific text. Participants should keep the text simple and short, not to make it indefinitely long to discourage potential readers. The participants have to answer two questions: (1) What kind of background information should be provided (e.g. which terms should be contextualized by giving a definition and/or application)? (2) What information is the most relevant (passage retrieval from an external source, e.g. Wikipedia)? The evaluation will be a combination of manual assessment and automatic metrics.

3.3 Task 3: Scientific text simplification

In this pilot task, the participants will be provided with the abstract of scientific papers. The goal will be to provide a simplified version of these abstracts. In this pilot task, we thus consider that the summarization part is already solved and that the main science nuggets are in the provided summaries. We will thus use scientific paper summaries which consist on context, aims, methodology, findings and discussion. Some medical papers will be used in this task. The guideline will detail the targeted simplification. Evaluation will be a combination of manual and automatic evaluation, the results of which will also be discussed during the workshop.

4 Conclusion

The paper introduced the CLEF 2021 SimpleText track, consisting of a workshop and pilot tasks on text simplification for scientific information access. Full details about the tasks and how to participate in the track can be found in the detailed call for papers and guidelines at the SimpleText website: <https://www.irit.fr/simpleText/>. Please join this effort and contribute by working on one of the greatest challenges of today!

References

1. "2019-nCoV" OR ... Publication Year: 2020 in Publications - Dimensions, <https://covid-19.dimensions.ai/>
2. background knowledge, <https://www.thefreedictionary.com/background+knowledge>
3. Altbach, P.G., Wit, H.d.: Too much academic research is being published (Jul 2018), <https://www.universityworldnews.com/post.php?story=20180905095203579>
4. Augenstein, I., Das, M., Riedel, S., Vikraman, L., McCallum, A.: Semeval 2017 task 10: Scienceeie-extracting keyphrases and relations from scientific publications. arXiv preprint arXiv:1704.02853 (2017)
5. Barzilay, R., Elhadad, N., McKeown, K.R.: Inferring strategies for sentence ordering in multidocument news summarization. *Journal of Artificial Intelligence Research* pp. 35–55 (2002), 17
6. Bellot, P., Doucet, A., Geva, S., Gurajada, S., Kamps, J., Kazai, G., Koolen, M., Mishra, A., Moriceau, V., Mothe, J., Preminger, M., SanJuan, E., Schenkel, R., Tannier, X., Theobald, M., Trappett, M., Wang, Q.: Overview of INEX. In: *Information Access Evaluation. Multilinguality, Multimodality, and Visualization - 4th International Conference of the CLEF Initiative, CLEF 2013, Valencia, Spain, September 23-26, 2013. Proceedings.* pp. 269–281 (2013)
7. Bellot, P., Moriceau, V., Mothe, J., SanJuan, E., Tannier, X.: INEX Tweet Contextualization task: Evaluation, results and lesson learned. *Inf. Process. Manage.* **52**(5), 801–819 (2016). <https://doi.org/10.1016/j.ipm.2016.03.002>, <https://doi.org/10.1016/j.ipm.2016.03.002>
8. Chae, J., Nenkova, A.: Predicting the fluency of text with shallow structural features: case studies of machine translation and human-written text. *Proc. of the 12th Conference of the European Chapter of the ACL* pp. 139–147 (2009)
9. Chall, J.S., Dale, E.: *Readability revisited: The new Dale–Chall readability.* MA: Brookline Books, Cambridge (1995)
10. Collins-Thompson, K., Callan, J.: A Language Modeling Approach to Predicting Reading Difficulty. *Proceedings of HLT/NAACL* **4** (2004)
11. Coster, W., Kauchak, D.: Simple english wikipedia: a new text simplification task. In: *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies.* pp. 665–669 (2011)
12. Cram, D., Daille, B.: Terminology Extraction with Term Variant Detection. In: *Proceedings of ACL-2016 System Demonstrations.* pp. 13–18. Association for Computational Linguistics, Berlin, Germany (Aug 2016). <https://doi.org/10.18653/v1/P16-4003>, <https://www.aclweb.org/anthology/P16-4003>
13. Ermakova, L., Cossu, J.V., Mothe, J.: A survey on evaluation of summarization methods. *Information Processing & Management* **56**(5), 1794–1814 (Sep 2019). <https://doi.org/10.1016/j.ipm.2019.04.001>, <http://www.sciencedirect.com/science/article/pii/S0306457318306241>
14. Ermakova, L., Goeuriot, L., Mothe, J., Mulhem, P., Nie, J.Y., SanJuan, E.: CLEF 2017 Microblog Cultural Contextualization Lab Overview. In: *Experimental IR Meets Multilinguality, Multimodality, and Interaction - 8th International Conference of the CLEF Association, CLEF 2017, Dublin, Ireland, September 11-14, 2017, Proceedings.* pp. 304–314 (2017). https://doi.org/10.1007/978-3-319-65813-1_27, https://doi.org/10.1007/978-3-319-65813-1_27
15. Ermakova, L., Mothe, J., Firsov, A.: A Metric for Sentence Ordering Assessment Based on Topic-Comment Structure (short paper). In: *ACM SIGIR Special Interest Group on Information Retrieval (SIGIR), Tokyo, Japan, 07/08/2017-11/08/2017 (août 2017), selection rate 30*

16. Fecher, B., Friesike, S.: Open science: one term, five schools of thought. In: *Opening science*, pp. 17–47. Springer, Cham (2014)
17. Feng, L., Jansche, M., Huenerfauth, M., Elhadad, N.: A comparison of features for automatic readability assessment. In: *Proceedings of the 23rd International Conference on Computational Linguistics: Posters*. pp. 276–284. COLING '10, Association for Computational Linguistics, Stroudsburg, PA, USA (2010), <http://dl.acm.org/citation.cfm?id=1944566.1944598>
18. Filippova, K., Altun, Y.: Overcoming the lack of parallel data in sentence compression. In: *Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing*. pp. 1481–1491 (2013)
19. Flesch, R.: A new readability yardstick. *Journal of Applied Psychology* **32**(3), p221 – 233 (Jun 1948)
20. Fry, E.: A readability formula for short passages. *Journal of Reading* **8**, 594–597 (1990), 33
21. Fry, E.: *The Varied Uses of Readability Measurement* (Apr 1986)
22. Heilman, M., Collins-Thompson, K., Eskenazi, M.: An analysis of statistical models and features for reading difficulty prediction. In: *Proceedings of the Third Workshop on Innovative Use of NLP for Building Educational Applications*. pp. 71–79. EANL '08, Association for Computational Linguistics, Stroudsburg, PA, USA (2008), <http://dl.acm.org/citation.cfm?id=1631836.1631845>
23. Hoffart, J., Yosef, M.A., Bordino, I., Fürstenau, H., Pinkal, M., Spaniol, M., Taneva, B., Thater, S., Weikum, G.: Robust disambiguation of named entities in text. In: *Proceedings of the 2011 Conference on Empirical Methods in Natural Language Processing*. pp. 782–792 (2011)
24. Jarreau, P.B., Porter, L.: Science in the Social Media Age: Profiles of Science Blog Readers. *Journalism & Mass Communication Quarterly* **95**(1), 142–168 (Mar 2018). <https://doi.org/10.1177/1077699016685558>, <https://doi.org/10.1177/1077699016685558>, publisher: SAGE Publications Inc
25. Ladyman, J., Lambert, J., Wiesner, K.: What is a complex system? *European Journal for Philosophy of Science* **3**(1), 33–67 (Jan 2013). <https://doi.org/10.1007/s13194-012-0056-8>, <https://doi.org/10.1007/s13194-012-0056-8>
26. Leroy, G., Endicott, J.E., Kauchak, D., Mouradi, O., Just, M.: User evaluation of the effects of a text simplification algorithm using term familiarity on perception, understanding, learning, and information retention. *Journal of medical Internet research* **15**(7), e144 (2013)
27. McCarthy, P.M., Guess, R.H., McNamara, D.S.: The components of paraphrase evaluations. *Behavior Research Methods* **41**(3), 682–690 (Aug 2009). <https://doi.org/10.3758/BRM.41.3.682>, <https://doi.org/10.3758/BRM.41.3.682>
28. Molek-Kozakowska, K.: Communicating environmental science beyond academia: Stylistic patterns of newsworthiness in popular science journalism. *Discourse & Communication* **11**(1), 69–88 (Feb 2017). <https://doi.org/10.1177/1750481316683294>, <https://doi.org/10.1177/1750481316683294>
29. Mutton, A., Dras, M., Wan, S., Dale, R.: Gleu: Automatic evaluation of sentence-level fluency. *ACL-07* pp. 344–351 (2007)
30. O'Reilly, T., Wang, Z., Sabatini, J.: How Much Knowledge Is Too Little? When a Lack of Knowledge Becomes a Barrier to Comprehension. *Psychological Science* (Jul 2019). <https://doi.org/10.1177/0956797619862276>, <https://journals.sagepub.com/doi/10.1177/0956797619862276>, publisher: SAGE PublicationsSage CA: Los Angeles, CA
31. Papineni, K., Roukos, S., Ward, T., Zhu, W.J.: Bleu: a method for automatic evaluation of machine translation. In: *Proceedings of the 40th annual meeting of the Association for Computational Linguistics*. pp. 311–318 (2002)
32. Petersen, S.E., Ostendorf, M.: A machine learning approach to reading level assessment. *Comput. Speech Lang.* **23**(1), 89–106 (Jan 2009). <https://doi.org/10.1016/j.csl.2008.04.003>, <http://dx.doi.org/10.1016/j.csl.2008.04.003>

33. Pitler, E., Nenkova, A.: Revisiting readability: A unified framework for predicting text quality (2008)
34. Rao, S., Tetreault, J.: Dear sir or madam, may i introduce the gyafc dataset: Corpus, benchmarks and metrics for formality style transfer. In: Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers). pp. 129–140 (2018)
35. Shi, H., Revithis, S., Chen, S.S.: An agent enabling personalized learning in e-learning environments. In: Proceedings of the first international joint conference on Autonomous agents and multiagent systems: part 2. pp. 847–848. AAMAS '02, Association for Computing Machinery, New York, NY, USA (Jul 2002). <https://doi.org/10.1145/544862.544941>, <https://doi.org/10.1145/544862.544941>
36. Si, L., Callan, J.: A statistical model for scientific readability. In: Proceedings of the Tenth International Conference on Information and Knowledge Management. pp. 574–576. CIKM '01, ACM, New York, NY, USA (2001). <https://doi.org/10.1145/502585.502695>, <http://doi.acm.org/10.1145/502585.502695>
37. Stenner, A.J., Horablin, I., Smith, D.R., Smith, M.: The Lexile Framework. Durham, NC: MetaMetrics (1988)
38. Stenner, A., Horabin, I., Smith, D.R., Smith, M.: The lexile framework. Durham, NC: MetaMetrics (1988)
39. Sulem, E., Abend, O., Rappoport, A.: Semantic structural evaluation for text simplification. In: Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers). pp. 685–696 (2018)
40. S  e, S.O.: Algorithmic detection of misinformation and disinformation: Gricean perspectives. *Journal of Documentation* **74**(2), 309–332 (Jan 2018). <https://doi.org/10.1108/JD-05-2017-0075>, <https://doi.org/10.1108/JD-05-2017-0075>, publisher: Emerald Publishing Limited
41. Tavernier, J., Bellot, P.: Combining relevance and readability for INEX 2011 question–answering track pp. 185–195 (2011)
42. Wan, S., Dale, R., Dras, M.: Searching for grammaticality: Propagating dependencies in the viterbi algorithm. Proc. of the Tenth European Workshop on Natural Language Generation (2005)
43. Wubben, S., van den Bosch, A., Krahmer, E.: Sentence simplification by monolingual machine translation. In: Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers). pp. 1015–1024 (2012)
44. Xu, W., Napoles, C., Pavlick, E., Chen, Q., Callison-Burch, C.: Optimizing statistical machine translation for text simplification. *Transactions of the Association for Computational Linguistics* **4**, 401–415 (2016)
45. Zakaluk, B.L., Samuels, S.J.: Readability: Its Past, Present, and Future. International Reading Association, 800 Barksdale Rd (1988), <https://eric.ed.gov/?id=ED292058>
46. Zwarts, S., Dras, M.: Choosing the right translation: A syntactically informed classification approach. Proc. of the 22nd International Conference on Computational Linguistics pp. 1153–1160 (2008)