

Content Curation for Consumer Health Search: Search and Misinformation Detection

Marcos Fernández-Pichel¹, Mario Ezra Aragón¹, Manuel Couto¹ and David E. Losada¹

¹*Centro Singular de Investigación en Tecnoloxías Intelixentes (CiTIUS), Universidade de Santiago de Compostela, Santiago de Compostela, 15782, Spain*

Abstract

The C3HS project focuses on developing technologies for the analysis, retrieval, and curation of health-related online content, with special attention to misinformation detection and early identification of health risks in digital environments. The research combines techniques from Information Retrieval, Natural Language Processing, and large-scale Data Processing to study how health information is produced, accessed, and interpreted on the web and social media. Our efforts are oriented toward developing models to assess the credibility and reliability of health-related content, evaluating search engines and conversational AI systems in answering medical questions, and designing new evaluation metrics that consider not only relevance and correctness but also the potential risks to users. Additionally, the project explores scalable architectures for large-scale data collection and processing, including focused crawling strategies and distributed pipelines that analyze large volumes of online data in real time.

Keywords

Consumer Health Search, Health Misinformation Detection, Information Retrieval, Natural Language Processing, Social Media Analysis, Credibility, Large-scale text analytics

1. Introduction

Recent advances in digital communication have transformed how people access and share health-related information. Search engines, online forums, social media platforms, and AI-generated content have become primary sources for users seeking advice about symptoms, treatments, or medical conditions. However, this abundance of information also poses significant challenges, including the rapid dissemination of misleading or unreliable content and the difficulty users face in assessing the credibility of online sources and information-access tools. These issues motivate the development of computational methods that improve access to trustworthy health information and detect potentially harmful or misleading content in large-scale online environments.

The coordinated research project C3HS addresses these challenges by integrating techniques from Information Retrieval (IR), Natural Language Processing (NLP), and large-scale data analysis. The project aims to study how health information is produced, disseminated, and consumed on the web, and to develop computational tools that facilitate safer, more reliable access to this information. In particular, the project investigates new evaluation frameworks and methodologies that go beyond traditional relevance-based metrics by incorporating dimensions such as credibility, reliability, and potential risk for end users.

Within this coordinated effort, the USC subproject focuses on analyzing and processing large collections of health-related texts from diverse online sources. Its work combines data-collection infrastructure, linguistic analysis, and machine learning methods to support several research objectives. These include assessing the credibility of health information retrieved from search systems, evaluating conversational

SEPLN 2026: 42nd International Conference of the Spanish Society for Natural Language Processing, León, Spain, 22-25 September 2026.

✉ marcos.fernandez.pichel@usc.gal (M. Fernández-Pichel); ezra.aragon@usc.gal (M. E. Aragón); manuel.couto@usc.gal (M. Couto); david.losada@usc.gal (D. E. Losada)

🌐 <https://citius.gal/> (M. Fernández-Pichel); <https://citius.gal/> (M. E. Aragón); <https://citius.gal/> (M. Couto); <https://citius.gal/> (D. E. Losada)

🆔 0000-0002-6560-9832 (M. Fernández-Pichel); 0000-0002-8213-957X (M. E. Aragón); 0000-0002-0593-7199 (M. Couto); 0000-0001-8823-7501 (D. E. Losada)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

AI systems' ability to answer medical questions, and studying how users interact with and interpret health-related information in online environments.

Another central line of research in the USC subproject concerns the large-scale monitoring and analysis of user-generated content in social media. By studying linguistic and behavioral patterns in online texts, the project explores methods to identify signals associated with health risks, including the early detection of psychological disorders. This work relies on scalable data processing architectures capable of collecting and analyzing large volumes of textual data, enabling the development and evaluation of models for risk prediction and health-related information analysis.

Overall, the USC subproject contributes new datasets, evaluation methodologies, and computational models to improve the quality and safety of health information in digital ecosystems. The project seeks to support the development of systems that not only retrieve relevant information but also help users identify trustworthy content and reduce exposure to potentially harmful or misleading health information.

2. Funding Institutions

This project has been supported by the Spanish Ministry of Science and Innovation and the State Research Agency (MCIN/AEI) through the “Proyectos de Generación de Conocimiento 2022” call. The project follows the “Investigación Orientada” modality within the thematic priority of health and information technologies. The project was initially proposed as a three-year project, from September 1, 2023, to August 31, 2026.

3. Participating Teams

The research is carried out as part of a coordinated project that brings together two complementary subprojects led by different institutions. The coordinating subproject (PID2022-137061OB-C21) is hosted at the University of A Coruña (UDC), while the second subproject (PID2022-137061OB-C22) is developed at the University of Santiago de Compostela (USC). Both teams collaborate to address challenges related to the search, selection, and organization of health-related information on the web. The coordinated structure allows the project to integrate complementary expertise in IR, NLP, large-scale data processing, and evaluation methodologies, enabling the development of new technologies for accessing and analyzing health-related online content.

The two subprojects pursue closely related but distinct objectives. The UDC subproject focuses on developing IR technologies for consumer health search, including search models, recommendation methods, and experimental evaluation frameworks to support reliable access to health information. In contrast, the USC subproject concentrates on large-scale analysis of online content and user-generated data, including methods for detecting health misinformation, assessing the credibility of medical information, and identifying early signals of health risks in social media environments.

The teams participating in the project bring together researchers with complementary backgrounds in computer science, IR, NLP, and large-scale data analytics. The UDC team includes researchers from the Center for Research in Information and Communication Technologies (CITIC), with extensive experience in search technologies, recommender systems, and experimental evaluation of information access systems. The USC team, based at the CiTIUS research center, includes specialists in IR, large-scale text mining, high-performance computing, and the analysis of health-related online behavior. Together, the participating researchers combine expertise in algorithmic design, large-scale experimentation, and interdisciplinary applications of language technologies to analyze health information and online user behavior.

4. Objectives and Tasks

In this paper, we focus on the USC Subproject's activities. Here, we describe the main organization of objectives and tasks, and in the next section, we report the main outcomes achieved so far.

The project is structured around a set of research objectives aimed at advancing the analysis and management of health-related information in large-scale online environments. These objectives focus on developing computational methods and infrastructures for collecting, processing, and analyzing large volumes of textual data from the web and social media. The work is organized into several interconnected tasks that address key challenges, including evaluating the credibility of health information, detecting misleading or harmful content, and identifying behavioral signals associated with health-related risks.

To achieve these goals, the project combines methodological developments with the creation of experimental resources and evaluation frameworks. Several tasks focus on designing and implementing large-scale data collection strategies, including targeted crawling techniques and pipelines to process heterogeneous online sources. These infrastructures enable the construction of datasets and experimental environments that support the study of health information access, the evaluation of search engines and conversational AI systems, and the analysis of user-generated content in social media.

In addition, the organization of the tasks reflects the project's interdisciplinary nature. The activities integrate techniques from various fields, enabling the team to study both the properties of health-related information and the behavior of users interacting with it. This structure facilitates the development of end-to-end approaches that span data acquisition and large-scale processing through to the design and evaluation of models for credibility assessment, misinformation detection, and early risk prediction. Specifically, the C3HS project was organized into the following work packages:

WP0: Project management and coordination:

- T0.1. Kick-off: Refinements are required based on the budget, personnel hiring, and the initial meeting.
- T0.2. Monitoring and control: risk control and contingency planning, change management, monitoring deviations from plans, ensuring compliance with objectives, quality assurance, writing monitoring reports, holding periodic meetings, and making decisions on corrective actions.
- T0.3. Close project. Final report, analysis of compliance with objectives, the study of deviations on planning, analysis of problems, errors, and lessons learned, analysis of results of technology transfer and dissemination.

WP1: Creation of test collections and evaluation methods:

- T1.1. Analysis of areas of consumer health search (e.g., related to mental health, psychological disorders, nutrition-related needs, etc.). Study of their suitability for the project.
- T1.2. Development of crawling software, content compilation, and information extraction.
- T1.3. Creation of test collections for content curation and evaluation of their quality.
- T1.4. Proposal of new metrics and evaluation methodologies for content curation oriented to consumer health search.

WP2: Search technologies for curated health-related content:

- T2.1. Models of a search for evidence of the most relevant and highest quality health-related information.
- T2.2. Generation of queries and query refinement based on medical resources and reputed terminologies. Query expansion and relevance feedback.
- T2.3. Neural re-ranking models. Deep learning strategies for analyzing the quality of textual content.
- T2.4. Passage/Sentence retrieval. Fusion of rankings.
- T2.5. Extraction and analysis of topics. Temporal evolution models.

WP3: Misinformation detection and credibility estimation:

- T3.1. Automatic elaboration of resources for misinformation detection. Unsupervised analysis of texts and extraction of information. Identification and selection of multiword terms, named entities, and semantic relationships.
- T3.2. Adaptation of analysis and extraction tools for the detection of signs of misinformation. Models of credibility.
- T3.3. Distributional techniques for the elaboration of semantic models. Selection of relevant contexts. Transparent vectors. Embeddings. Transformers.
- T3.4. Algorithmic development and adaptation of semi-automatic and automatic discourse analysis techniques for estimating credibility and detecting signs of misinformation. Systematic discourse analysis for extracting discursive patterns. Discursive coherence in the detection of misinformation and credibility.
- T3.5. Definition of discursive models including factors for the detection of health-related misinformation and untruthfulness.

WP4: Massive processing technologies for extracting curated content in real-time:

- T4.1. New crawling technologies ("focused crawling") for extracting health-related curated content.
- T4.2. Big Data architectures with task-specific topologies tailored to the task pipeline (retrieval, processing, classification, etc.) are necessary for analyzing curated content in real time.
- T4.3. Load prediction and adaptive methods dynamically adjust the platform to the volume of available information.
- T4.4. Development of a new Big Data framework that supports the same multi-language code HPC tasks (using MPI) with Big Data tasks (using MapReduce operations).

WP5: Content personalization and recommendation of health-related information:

- T5.1. Definition of user profiles and collection of items to be recommended (articles, experts, individuals' information needs, expert responses). Search and construction of experimental collections with expert validation. Risk profiles of users and personality analysis.
- T5.2. Development and implementation of content-based models trained on expert-labeled data (golden truth pairs: user search request-good response item). Train semantic models using bi- and cross-encoders, as well as text-to-text transformer-based architectures.
- T5.3. Formalization of expert knowledge to guide search, web mining, and recommendation models (knowledge of psychologists, specialized medical questionnaires, formalized procedures in these domains, etc.). Selection of web sources and personalization strategy.
- T5.4. Definition of new evaluation models that measure not only the quality (precision, diversity, novelty) of the suggested items, but also the appropriateness of the provided results to help individuals in connection with their health condition/profile.
- T5.5. Evaluation of the models implemented with these evaluation methodologies and metrics.

WP6: Demonstrators and technology transfer:

- T6.1. Task-oriented integration of the new algorithms resulting from the gradual achievement of C3HS objectives.
- T6.2. Development of demonstrators tailored to different health-related needs (e.g., psychological issues, nutritional needs, eating disorders).

WP7: Scientific, industrial and social dissemination:

- T7.1. Scientific dissemination.
- T7.2. Dissemination of demonstrators in industrial and social forums.

5. Results

Over the course of the project, we have conducted extensive research at the intersection of information retrieval, natural language processing, and computational social science, with a particular focus on health information, misinformation, and mental health analysis in online environments. Our work has explored multiple complementary directions, including the detection of psychological signals in social media, the development of resources and models for mental health assessment, and the study of how users search for and evaluate health and news information online. In parallel, we have investigated methods to improve the safety and reliability of information access, addressing challenges such as harmful queries, misinformation annotation, and the credibility of online content. The following subsections summarize the main outcomes of these research efforts (see Figure 1 for a global overview of the results).

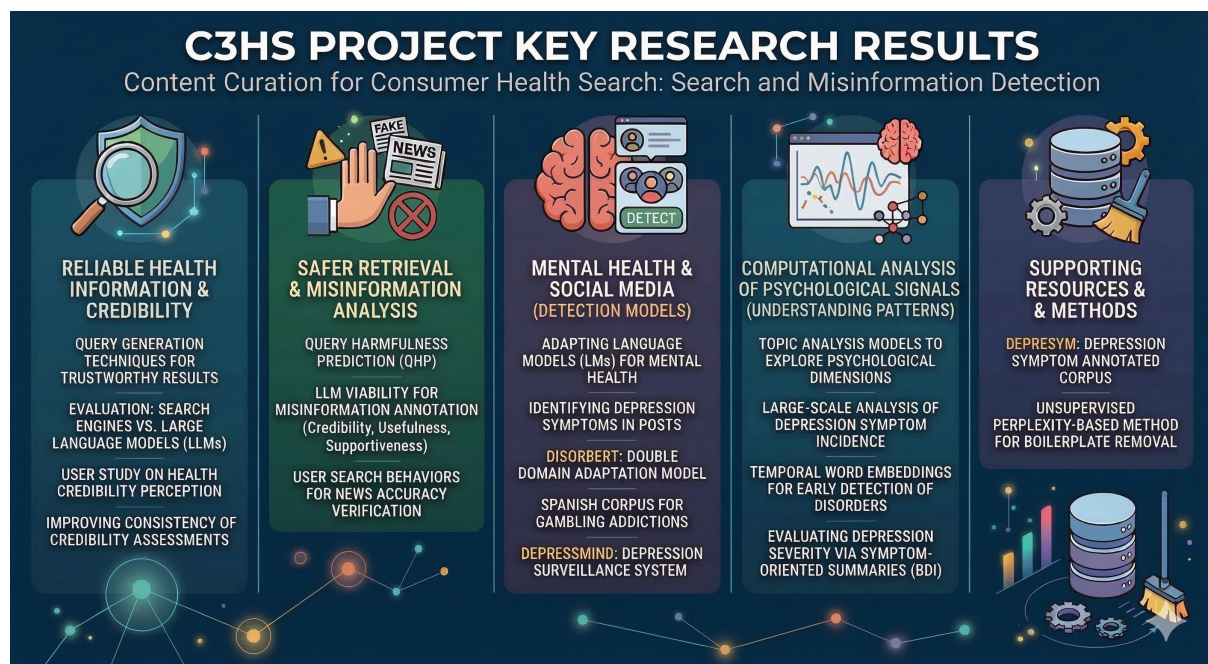


Figure 1: Summary of the main publications derived from the project.

5.1. Reliable information access and health information credibility

A central line of research within the project has examined how people search for and assess health information online, with the aim of improving the reliability and credibility of information access systems.

In “Generating Effective Health-Related Queries for Promoting Reliable Search Results” [1], we proposed new query generation techniques to retrieve more reliable health information. The approach reformulates health-related queries to increase the likelihood of retrieving trustworthy sources, improving the quality of search results.

In “Evaluating search engines and large language models for answering health questions” [2], we compared the performance of traditional web search engines and large language models in answering health-related questions. The study analyzed the quality, reliability, and potential risks of each information access approach.

In “A user study on people’s perception to the credibility of online health information” [3], we investigated how users perceive and evaluate the credibility of online health content. Through a comprehensive user study, we analyzed the factors that influence trust judgments when people assess health information sources.

In “Improving the Reliability of Health Information Credibility Assessments” [4], we studied methods to improve the consistency and reliability of credibility assessments in health information retrieval. The work analyzed different evaluation strategies and signals that can help produce more robust credibility judgments.

5.2. Safer retrieval and misinformation analysis

The project has also investigated methods to improve the safety of information retrieval systems and support research on misinformation and harmful content.

In “Query Harmfulness Prediction (QHP): a New Challenge for Safer Retrieval Systems” [5], we introduced a new task that aims to predict whether a search query is likely to retrieve harmful content despite being topically relevant. The paper defines the task and evaluates different approaches for estimating query harmfulness before executing the search.

In “On the Viability of Exploiting Large Language Models for Misinformation Annotation” [6], we analyzed whether large language models can reliably assist in annotating misinformation-related dimensions such as credibility, usefulness, and supportiveness. The study compared LLM-based annotations with human judgments and discussed their potential and limitations.

In “Query Smarter, Trust Better? Exploring Search Behaviours for Verifying News Accuracy” [7], we examined how users formulate search queries when attempting to verify the accuracy of news content. Through behavioral analysis and user studies, the work identified search strategies that support effective fact-checking.

5.3. Mental health and social media

Another part of the project has focused on developing computational methods and linguistic resources to detect and analyze mental health signals in social media. This included the design of specialized datasets, the adaptation of language models, and the development of systems for large-scale monitoring of mental health indicators.

In “Adapting Language Models for Mental Health Analysis on Social Media” [8], we adapted pretrained language models for mental health detection on social media. The work evaluated various domain-adaptation models and demonstrated that targeted adaptation can significantly improve the identification of mental health indicators in social media text.

In “Online expressions, offline struggles: Using social media to identify depression-related symptoms” [9], we analyzed how computational linguistic methods can identify depression-related symptoms in social media posts. The study explored linguistic patterns associated with clinical symptoms and evaluated models that detect these signals in large-scale online data.

In “Analyzing Gambling Addictions: A Spanish Corpus for Understanding Pathological Behavior” [10], we created a Spanish-language corpus for studying gambling addiction in online environments. The dataset contains annotated content reflecting pathological gambling behaviors and supports the development of computational models for addiction-related behavior analysis.

In “DepressMind: A Depression Surveillance System for Social Media Analysis” [11], we presented a system for large-scale monitoring and analysis of depression-related signals in social media. The platform integrates data collection, processing, and visualization components to support researchers in studying mental health trends over time.

In “DisorBERT: A Double Domain Adaptation Model for Detecting Signs of Mental Disorders in Social Media” [12], we proposed a double-domain adaptation approach that adapts pretrained language models to both mental health content and the linguistic characteristics of social media. The model improves the detection of signals of mental disorders in user-generated content and outperforms standard strategies across multiple datasets.

5.4. Computational analysis of psychological signals

Beyond detection models, the project has also explored computational approaches to better understand psychological patterns and behavioral signals expressed in online communication.

In “Exploiting topic analysis models to explore psychological dimensions in social media data” [13], we investigated how topic modeling can uncover latent psychological patterns in social media discussions. The study links the thematic structures identified by topic models to psychological dimensions relevant to behavioral analysis.

In “On the incidence of depression symptoms on social media” [14], we analyzed the prevalence and distribution of depression-related symptoms expressed in social media platforms. The study provided large-scale evidence of how different symptoms appear in online discourse and examined their relative frequency and patterns.

In “Temporal Word Embeddings for Early Detection of Psychological Disorders on Social Media” [15], we proposed the use of temporal word embeddings to model linguistic changes over time. By capturing evolving semantic patterns in user language, the approach enables the early detection of signals associated with emerging psychological disorders.

In “Delving into the Depths: Evaluating Depression Severity through BDI-biased Summaries” [16], we explored the use of summaries guided by the Beck Depression Inventory (BDI) symptom framework to estimate depression severity. The study analyzed how symptom-oriented summarization can support the automated assessment of mental health severity levels.

5.5. Additional supporting resources and methods

The project has also produced methodological contributions and new resources that support several of the research directions described above.

In “DepreSym: A Depression Symptom Annotated Corpus and the Role of Large Language Models as Assessors of Psychological Markers” [17], we introduced DepreSym, a dataset of sentences annotated according to the depression symptoms defined in the BDI-II clinical inventory. The work also analyzed the potential of large language models to assist in annotating psychological markers.

In “An unsupervised perplexity-based method for boilerplate removal” [18], we proposed an unsupervised method that uses the notion of perplexity to identify and remove boilerplate content from web pages. The approach improves the quality of textual data used in downstream natural language processing and information retrieval tasks.

6. Conclusions

The work carried out in the C3HS project demonstrates the potential of combining Information Retrieval, Natural Language Processing, and large-scale data analytics to address critical challenges in consumer health search. Through the development of novel datasets, evaluation frameworks, and computational models, the project has improved the reliability, safety, and interpretability of health-related information in online environments. In particular, the results highlight the importance of moving beyond traditional relevance-based evaluation and incorporating dimensions such as credibility, harmfulness, and user risk in the design and assessment of information access systems.

Furthermore, the project provides evidence that large-scale analysis of user-generated content can support the early detection of health-related risks and the identification of misleading or harmful information. The integration of scalable data-collection infrastructure with advanced modeling techniques has enabled the study of both information quality and user behavior at scale, offering new insights into how people interact with health information online. Overall, the project’s outcomes lay the groundwork for future research on safer, more trustworthy information access systems, with potential applications in public health monitoring, decision support, and the development of responsible AI technologies in the health domain.

Declaration on Generative AI

During the preparation of this work, the author(s) used Claude Opus and Grammarly to Perform Grammar and spelling checks. Furthermore, the author(s) used Google Gemini to generate the images for Figure 1. After using these tool(s)/service(s), the author(s) reviewed and edited the content as needed and take(s) full responsibility for the publication's content.

Acknowledgments

This work was supported by MICIU/AEI/10.13039/ 501100011033 through project PID2022-137061OB-C22 (co-funded by ERDF) and by the Xunta de Galicia – Consellería de Cultura, Educación, Formación Profesional e Universidades through grants ED431G 2023/04 and ED431C 2022/19 (co-funded by ERDF). The authors also acknowledge the project “Cátedra de IA aplicada a la Medicina Personalizada de Precisión” (Cátedras ENIA, TSI-100932-2023-3), funded by the Ministerio de Transformación Digital y Función Pública (Secretaría de Estado de Digitalización e Inteligencia Artificial) and NextGenerationEU. The second author was supported by a Juan de la Cierva Grant (JDC2023-052296-I), funded by MCIN/AEI/10.13039/501100011033 and FSE+.

References

- [1] X. Carrera, M. Fernández-Pichel, D. E. Losada, Generating effective health-related queries for promoting reliable search results, in: Proceedings of the 48th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '25, Association for Computing Machinery, New York, NY, USA, 2025, p. 2627–2631. URL: <https://doi.org/10.1145/3726302.3730202>. doi:10.1145/3726302.3730202.
- [2] M. Fernández-Pichel, J. C. Pichel, D. E. Losada, Evaluating search engines and large language models for answering health questions, npj Digital Medicine 8 (2025) 153. URL: <https://doi.org/10.1038/s41746-025-01546-w>. doi:10.1038/s41746-025-01546-w.
- [3] M. Fernández-Pichel, M. Bink, D. E. Losada, D. Elswailer, A user study on people's perception to the credibility of online health information., in: Proceedings of the The 4th Workshop on Reducing Online Misinformation through Credible Information Retrieval (ROMCIR) hold in conjunction with ECIR 2024, 2024.
- [4] M. Fernández-Pichel, S. Meyer, M. Bink, A. Frummet, D. E. Losada, D. Elswailer, Improving the reliability of health information credibility assessments., in: Proceedings of the The 3rd Workshop on Reducing Online Misinformation through Credible Information Retrieval (ROMCIR) hold in conjunction with ECIR 2023, 2023.
- [5] X. Carrera, M. Fernández-Pichel, D. E. Losada, Query harmfulness prediction (qhp): a new challenge for safer retrieval system, in: Proceedings of the 48th European Conference in Information Retrieval, ECIR '26, 2026.
- [6] P. Landrove, M. Fernández-Pichel, D. E. Losada, On the viability of exploiting large language models for misinformation annotation, in: Proceedings of the 48th European Conference in Information Retrieval, ECIR '26, 2026.
- [7] D. Elswailer, S. Ateia, M. Bink, G. Donabauer, M. Fernández Pichel, A. Frummet, U. Kruschwitz, D. E. Losada, B. Ludwig, S. Meyer, N. Pascual Presa, Query smarter, trust better? exploring search behaviours for verifying news accuracy, in: Proceedings of the 48th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '25, Association for Computing Machinery, New York, NY, USA, 2025, p. 515–526. URL: <https://doi.org/10.1145/3726302.3730067>. doi:10.1145/3726302.3730067.
- [8] M. E. Aragón, A. P. López-Monroy, M. M. y Gómez, D. E. Losada, Adapting language models for mental health analysis on social media, Artificial Intelligence in Medicine 168 (2025) 103217. URL:

<https://www.sciencedirect.com/science/article/pii/S0933365725001526>. doi:<https://doi.org/10.1016/j.artmed.2025.103217>.

- [9] M. E. Aragón, A. P. López-Monroy, M. M. y Gómez, D. E. Losada, Online expressions, offline struggles: Using social media to identify depression-related symptoms, *Online Social Networks and Media* 50 (2025) 100338. URL: <https://www.sciencedirect.com/science/article/pii/S2468696425000394>. doi:<https://doi.org/10.1016/j.osnem.2025.100338>.
- [10] M. Couto, M. Fernández-Pichel, M. E. Aragón, D. E. Losada, Analyzing gambling addictions: A Spanish corpus for understanding pathological behavior, in: C. Christodouloupoulos, T. Chakraborty, C. Rose, V. Peng (Eds.), *Findings of the Association for Computational Linguistics: EMNLP 2025*, Association for Computational Linguistics, Suzhou, China, 2025, pp. 17610–17619. URL: <https://aclanthology.org/2025.findings-emnlp.955/>. doi:10.18653/v1/2025.findings-emnlp.955.
- [11] R. Fernández-Iglesias, M. Fernández-Pichel, M. E. Aragón, D. E. Losada, DepressMind: A depression surveillance system for social media analysis, in: N. Aletras, O. De Clercq (Eds.), *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics: System Demonstrations*, Association for Computational Linguistics, St. Julians, Malta, 2024, pp. 35–43. URL: <https://aclanthology.org/2024.eacl-demo.5/>. doi:10.18653/v1/2024.eacl-demo.5.
- [12] M. E. Aragón, A. P. López-Monroy, L. C. González, D. E. Losada, M. Montes-y Gómez, DisorBERT: A double domain adaptation model for detecting signs of mental disorders in social media, in: A. Rogers, J. Boyd-Graber, N. Okazaki (Eds.), *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Association for Computational Linguistics, Toronto, Canada, 2023, pp. 15305–15318. URL: <https://aclanthology.org/2023.acl-long.853/>. doi:10.18653/v1/2023.acl-long.853.
- [13] M. Couto, J. Parapar, D. E. Losada, Exploiting topic analysis models to explore psychological dimensions in social media data, *Scientific Reports* 16 (2026) 6047. URL: <https://doi.org/10.1038/s41598-026-36339-y>. doi:10.1038/s41598-026-36339-y.
- [14] E. A. Ríssola, M. E. Aragón, D. E. Losada, F. Crestani, On the incidence of depression symptoms on social media, *Journal of Computational Social Science* 8 (2025) 48. URL: <https://doi.org/10.1007/s42001-025-00377-9>. doi:10.1007/s42001-025-00377-9.
- [15] M. Couto, A. Perez, J. Parapar, D. E. Losada, Temporal word embeddings for early detection of psychological disorders on social media, *Journal of Healthcare Informatics Research* (2025). URL: <https://doi.org/10.1007/s41666-025-00186-9>. doi:10.1007/s41666-025-00186-9.
- [16] M. E. Aragón, J. Parapar, D. E. Losada, Delving into the depths: Evaluating depression severity through BDI-biased summaries, in: A. Yates, B. Desmet, E. Prud'hommeaux, A. Zirikly, S. Bedrick, S. MacAvaney, K. Bar, M. Ireland, Y. Ophir (Eds.), *Proceedings of the 9th Workshop on Computational Linguistics and Clinical Psychology (CLPsych 2024)*, Association for Computational Linguistics, St. Julians, Malta, 2024, pp. 12–22. URL: <https://aclanthology.org/2024.clpsych-1.2/>. doi:10.18653/v1/2024.clpsych-1.2.
- [17] A. Pérez, M. Fernández-Pichel, J. Parapar, D. E. Losada, Depresym: A depression symptom annotated corpus and the role of large language models as assessors of psychological markers, *Language Resources and Evaluation* 59 (2025) 2737–2762. URL: <https://doi.org/10.1007/s10579-025-09831-6>. doi:10.1007/s10579-025-09831-6.
- [18] M. Fernández-Pichel, M. Prada-Corral, D. E. Losada, J. C. Pichel, P. Gamallo, An unsupervised perplexity-based method for boilerplate removal, *Natural Language Engineering* 30 (2024) 132–149. doi:10.1017/S1351324923000049.