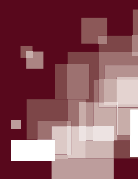


IJDL

International Journal of
DIGITAL LAW



Predictive pattern analysis in administrative and tax litigation through generative AI and regular expressions

Análisis de patrones predictivos en litigios administrativos y tributarios mediante IA generativa y expresiones regulares

Juan Gustavo Corvalán*

Universidad de Buenos Aires (Buenos Aires, Argentina)
corvalanjuang@gmail.com
<https://orcid.org/0009-0004-0241-3236>

Sofía Tammaro**

Universidad de Buenos Aires (Buenos Aires, Argentina)
sofiatammaro97@gmail.com
<https://orcid.org/0009-0005-9592-3362>

Gisel Alvarado***

Universidad de Buenos Aires (Buenos Aires, Argentina)
aAlvaradofgise@gmail.com
<https://orcid.org/0009-0009-2984-5588>

Como citar este artigo/*How to cite this article*: CORVALÁN, Juan Gustavo *et al.* Predictive pattern analysis in administrative and tax litigation through generative AI and regular expressions. *International Journal of Digital Law*, Belo Horizonte, v. 7, e701, 2026. DOI: 10.47975/ijdl.v7.1305.

- * Director of the Innovation and Artificial Intelligence Laboratory at the University of Buenos Aires Law School (Buenos Aires, Argentina). Holds a Doctorate in Legal Sciences from the Universidad del Salvador. Director of the Postgraduate Program in Artificial Intelligence and Law at the University of Buenos Aires (UBA). Director of the Diploma Program in Law 4.0 at Austral University. Co-creator of Prometea, the first predictive artificial intelligence system serving the justice system. Co-creator of PretorAI and Academic Director for the implementation of that system at the Constitutional Court of Colombia. General Project Director under the agreement to implement Prometea at the Inter-American Court of Human Rights. Currently serves as Deputy Prosecutor General for Contentious-Administrative and Tax Matters before the Superior Court of Justice of the Autonomous City of Buenos Aires. *E-mail*: corvalanjuang@gmail.com.
- ** Criminal Justice Agenda Lead at the Innovation and Artificial Intelligence Laboratory (IALAB) at the Faculty of Law, University of Buenos Aires (Buenos Aires, Argentina). Lecturer in the Postgraduate Program on Artificial Intelligence and Law, University of Buenos Aires. *E-mail*: sofiatammaro97@gmail.com.
- *** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Law student at the Catholic University of Cuyo. *E-mail*: alvaradofgise@gmail.com.

Luca Nicolás Forziatti Gangi****

Universidad de Buenos Aires (Buenos Aires, Argentina)
lforziatigangi@uade.edu.ar
<https://orcid.org/0009-0001-6085-6287>

Carina Mariel Papini*****

Universidad de Buenos Aires (Buenos Aires, Argentina)
carinapapini79@gmail.com
<https://orcid.org/0009-0006-4733-1634>

Mariana Sánchez Caparrós*****

Universidad de Buenos Aires (Buenos Aires, Argentina)
mariana.sanchezcaparros@ialab.com.ar
<https://orcid.org/0009-0003-6775-4819>

Agostina Celeste Jara Rey*****

Universidad de Buenos Aires (Buenos Aires, Argentina)
ajarey.contacto@gmail.com
<https://orcid.org/0009-0007-6920-2582>

Florencia Paola Croci*****

Universidad de Buenos Aires (Buenos Aires, Argentina)
florenciacroci@ialab.com.ar
<https://orcid.org/0009-0004-1630-4479>

Lola Ramos Pereyra*****

Universidad de Buenos Aires (Buenos Aires, Argentina)
loliramos99gmail.com
<https://orcid.org/0009-0000-9909-8623>

-
- **** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Bachelor of Laws from the Universidad Argentina de la Empresa. Lawyer. *E-mail:* lforziatigangi@uade.edu.ar.
- ***** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Academic Coordinator of the Continuing Education Program on Artificial Intelligence and Law at the University of Buenos Aires Law School. Currently pursuing a Master's degree in Constitutional Procedural Law at the University of Buenos Aires and a Master's degree in Artificial Intelligence at the Centro Europeo de Posgrado. *E-mail:* carinapapini79@gmail.com.
- ***** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Holds a Doctorate in Legal Sciences from UCA and a Master's degree in Administrative Law from Austral University. *E-mail:* mariana.sanchezcaparros@ialab.com.ar.
- ***** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Lawyer holding a postgraduate degree in Complex Corporate Law Issues from the National University of the Northeast. Currently pursuing a Master's degree in Business Law at the National University of the Northeast. *E-mail:* ajararey.contacto@gmail.com.
- ***** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Law graduate of the University of Buenos Aires. *E-mail:* florenciacroci@ialab.com.ar.
- ***** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Academic Coordination Assistant for the Postgraduate Program in Artificial Intelligence and Law at the University of Buenos Aires. *E-mail:* loliramos99gmail.com.

Melania Gadea*****

Universidad de Buenos Aires (Buenos Aires, Argentina)
 melaniagadea00@gmail.com
<https://orcid.org/0009-0005-3496-6051>

Luciano Della Via*****

Universidad de Buenos Aires (Buenos Aires, Argentina)
 dallavia.ldv@gmail.com
<https://orcid.org/0009-0002-8590-7666>

Recibido/Received: 11.08.2025 / August 11th, 2025**Aprobado/Approved:** 12.02.2026 / February 12th, 2026

Abstract: Text classification is the task of assigning a piece of text to an appropriate category, with the set of possible categories varying by domain. In this study, we evaluate ChatGPT's accuracy in classifying judicial rulings issued by an Argentine court, using prompts based on regular expression matches. We compare its performance to two alternative approaches: (1) classification based on natural language descriptions of the categories, and (2) a traditional regex-based algorithm. The regex patterns were developed by domain experts employed at the partner organization, and all experiments were conducted in Spanish. Our results show that ChatGPT achieves high accuracy when prompted with both natural language descriptions and regex-based category identifications. However, generating outputs with ChatGPT is significantly more time-consuming. Despite this limitation, the model's accessible interface and minimal technical requirements make it a promising tool for automating document classification. Furthermore, the use of regex-based prompts enables human oversight and active involvement in the classification process—an advantage in domains where interpretability and accountability are critical. All code, prompts, and evaluation materials are publicly available at <https://github.com/FAIR-IALAB-UBA/ai-regex>.

Keywords: Large language models. Document classification. Regular expressions. Predictive pattern analysis. Generative AI.

Resumen: La clasificación de texto consiste en asignar un fragmento de texto a una categoría adecuada. El conjunto de categorías posibles varía según el dominio. En este estudio, evaluamos la precisión de ChatGPT al clasificar sentencias judiciales emitidas por un tribunal argentino mediante indicaciones basadas en coincidencias de expresiones regulares. Comparamos su rendimiento con dos enfoques alternativos: (1) clasificación basada en descripciones de las categorías en lenguaje natural y (2) un algoritmo tradicional basado en expresiones regulares. Los patrones de expresiones regulares fueron desarrollados por expertos del sector de la organización asociada, y todos los experimentos se realizaron en español. Nuestros resultados muestran que ChatGPT alcanza una alta precisión cuando se le solicitan descripciones en lenguaje natural e identificaciones de categorías basadas en expresiones regulares. Sin embargo, generar resultados con ChatGPT requiere mucho más tiempo. A pesar de esta limitación, la interfaz accesible del modelo y sus mínimos requisitos técnicos lo convierten en una herramienta prometedora para automatizar la clasificación de documentos. Además, el uso de indicaciones basadas en expresiones regulares permite la supervisión humana y la participación activa en el proceso de clasificación, una ventaja en dominios donde la interpretabilidad

***** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Law graduate of the University of Buenos Aires. *E-mail:* melaniagadea00@gmail.com.

***** Member of the Innovation and Artificial Intelligence Laboratory (IALAB) at the University of Buenos Aires Law School (Buenos Aires, Argentina). Law graduate of the Pontifical Catholic University of Argentina. Holder of a Master's degree in Law and New Technologies from the School of Legal Practice at the Complutense University of Madrid. *E-mail:* dallavia.ldv@gmail.com.

y la rendición de cuentas son cruciales. Todo el código, las indicaciones y los materiales de evaluación están disponibles públicamente en <<https://github.com/FAIR-IALAB-UBA/ai-regex>>.

Palabras clave: Modelos de lenguaje grandes. Clasificación de documentos. Expresiones regulares. Análisis predictivo de patrones. IA generativa.

Summary: **1** Introduction – **2** Methodology – **3** Experiments and results – **4** Limitations and future work – **5** Conclusions – **6** Acknowledgments – References

1 Introduction

The operational dynamics of today's large organizations are marked by the daily management of an ever-growing volume of information.^{1,2} Every day, public agencies and private companies handle documents that must be classified, stored, and retrieved under time-sensitive conditions. This challenge is particularly pronounced in the public sector, where administrative workloads place significant demands on human teams. The issue becomes especially critical in fields where deadlines are imperative, such as law medicine or public safety.³ For instance, within the Argentine judiciary, criminal courts reported a total of 52,165 convictions.⁴ More recently, in 2024, the Supreme Court of Argentina handled 20,611 cases.⁵

Among the steps involved in every organizational process, document classification is crucial in managing data and may be one of the least time-consuming on an individual basis, yet when accumulated, it can represent a significant portion of the operational capacity of qualified personnel. Moreover, manual classification is prone to human error and inconsistencies,⁶ which can affect the reliability and efficiency of the broader decision-making process.

For this reason, well-trained machine learning (ML) algorithms have the potential to automate this task, enhancing both efficiency and productivity. In general, these models not only reduce the time but also are highly accurate for classification.⁷

¹ KLEIN, L. K; EARL, E.; CUNDICK, D. Reducing information overload in your organization. *Harvard Business Review*, 2023.

² BHAMBRI, Gaurav. Information overload in business organizations and entrepreneurship: An analytical review of the literature. *Business Information Review*, v. 38, issue 4, 2021, p. 193–200.

³ MAHONEY, Christian; GRONVALL, Peter; HUBER-FLIFLET; ZHANG, Jianping. Explainable text classification techniques in legal document review: locating rationales without using human annotated training text snippets. In: *2022 IEEE International Conference on Big Data (Big Data)*. IEEE, p. 2044-2051, 2022.

⁴ NATIONAL JUDICIAL STATISTICS SYSTEM (SNEJ). *Annual Report: Technical Report*. National Judicial Statistics System (SNEJ): Argentina, 2022.

⁵ SUPREME COURT OF ARGENTINA. *Statistics For The Year 2024: Technical report*. Supreme Court of Argentina, 2024.

⁶ PALANIVINAYAGAM, Ashokkumar; ZIAD EL-BAYEH, Claude; DAMAŠEVICIUS, Robertas. Twenty years of machine-learning-based text classification: A systematic review. *Algorithms*, v. 16, issue 5, p. 236, 2023.

⁷ JIANG, Shuo; HU, Jie; MAGEE, Christopher L.; LUO, Jianxi. Deep learning for technical document classification. *IEEE Transactions on Engineering Management*, v. 71, p. 1163-1179, 2022.

However, despite their better accuracy, there is lots of room for improvement,⁸ and their success relies on their capacity to understand complex and non-linear relationships within data. In recent years, there has been an exponential growth in the number of complex documents and texts that require a deeper understanding of ML methods to be able to accurately classify texts in many applications.⁹

In recent years, the emergence of large language models (LLMs) and their rapidly expanding capabilities have led to their adoption across a wide range of societal sectors, including public institutions. In this context, AI has become increasingly “democratized”—a reference to the reduced costs associated with acquiring and deploying AI tools, as well as the development of intuitive interfaces that enable human-AI interaction¹⁰ without requiring extensive training or technical expertise.¹¹

In this context, language models offer the potential to automate a wide range of administrative tasks at low cost and technically accessible. In the public sector, they are being used to support decision-making processes¹² in the national security context,¹³ or to analyse public affairs documents.¹⁴ Additional applications include document classification, information retrieval, and the drafting of routine reports or communications, all of which can significantly enhance institutional efficiency.

However, despite the impressive capabilities and accessibility of LLMs, their performance in information extraction is still not fully satisfactory.¹⁵ While they have demonstrated remarkable proficiency in generating unstructured natural language based on instructions, their output may be inconsistent when required to conform to specific structured formats —¹⁶ a crucial aspect in applications such as entity

⁸ PALANIVINAYAGAM, Ashokkumar; ZIAD EL-BAYEH, Claude; DAMAŠEVICIUS, Robertas. Twenty years of machine-learning-based text classification: A systematic review. *Algorithms*, v. 16, issue 5, p. 236, 2023.

⁹ KOWSARI, Kamram; JAFARI MEIMANDI, Kiana; HEIDARYSAFA, Mojtaba; MENDU, Sanjana; BARNES, Laura; BROWN, Donald. Text classification algorithms: A survey. *Information*, v. 10, issue 4, p. 150, 2019.

¹⁰ ABUKAR AHMED, Abdinor; KHALIF ALI, Mohamed. User-centric adoption of democratized generative AI: Focus on human-machine interaction and overcoming challenges. *International Journal of Engineering Trends and Technology*, v. 72, issue 9, p. 78-95, 2024.

¹¹ SEGER, Elizabeth. *What do we mean when we talk about “AI democratisation”*. 2023. Available at: <<https://www.governance.ai/analysis/what-do-we-mean-when-we-talk-about-ai-democratisation>>.

¹² YOON, Sungwook Yoon; et al. Digital innovation in public administration through intelligent public sector automation (IPSA): Strategies and challenges. *Journal of Multimedia Information System*, v. 11, issue 4, p. 249-260, 2024.

¹³ CABELLERO, William N.; JENKINS, Phillip R. On large language models in national security applications. *Stat*, v. 14, issue 2, e70057, 2025.

¹⁴ PEÑA ALMANSA, Alejandro; MORALES MORENO, Aythami; FIÉRRER AGUILAR, Julián; SERNA, Ignacio; ORTEGA GARCÍA, Javier; PUENTE, Íñigo; CÓRDOVA, Jorge; CÓRDOVA, Gonzalo; et. al. Leveraging large language models for topic classification in the domain of public affairs. *Lecture Notes in Computer Science* (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics), 2023.

¹⁵ YE, Junjie; XU, Nuo; WANG, Yikun; ZHOU, Jie; ZHANG, Qi; GUI, Tao; HUANG, Xuanjing. *LLM-DA: Data augmentation via large language models for few-shot named entity recognition*. 2024. Available at: <<https://arxiv.org/abs/2402.14568>>.

¹⁶ LI, Yinghao; RAMPRASAD, Rampi; ZHANG, Chao. *A simple but effective approach to improve structured language model output for information extraction*. 2024. Available at: <<https://arxiv.org/abs/2402.13364>>.

extraction. Also, it has been argued that LLMs still significantly underperform fine-tuned models in the task of text classification,¹⁷ highlighting the need for financial investment to implement tools adapted to users' specific needs.

To address these limitations, and enhance classification accuracy, we integrated LLMs with Regular Expressions (RegEx)—highly effective tools for text matching and manipulation.¹⁸ With this combination, we tried to leverage the accessibility of LLMs alongside the precision of handcrafted patterns.

In this paper, we evaluate ChatGPT's accuracy in classifying judicial rulings issued by an Argentine court, using a prompt based on regular expression (regex) matches. We compare its performance with two alternative approaches: (1) a classification task based on natural language descriptions of the categories, and (2) a regex-based algorithm.

Our experiments show that the language model achieves high accuracy when prompted with both natural language identifications and regex-based category definitions. However, generating outputs with ChatGPT requires significantly more time.¹⁹ Despite this, its user-friendly interface and minimal technical requirements make it accessible to a broad range of users, positioning it as a viable tool for automating document classification tasks. Moreover, the regex-based task with the LLM design enables human oversight and active participation in the classification process, which is particularly valuable in certain contexts.

2 Methodology

Text classification is a problem designed to assign a piece of text a proper class or category.^{20 21} The classes depend on the domain under consideration and range from a variety of topics [20]. In our case, the documents under analysis were judicial rulings, and the categories were defined by legal experts. These experts also designed the regular expressions (regex) to identify relevant features within the rulings and assign appropriate labels. We evaluated ChatGPT's performance in classifying

¹⁷ SUN, Xiaofei; LI, Xiaoya; LI, Jiwei; WU, Fei; GUO, Shangwei; ZHANG, Tianwei; WANG, Guoyin. *Text classification via large language models*. 2023. Available at: <<https://arxiv.org/abs/2305.08377>>.

¹⁸ ZHANG, Shuai; GU, Xiaodong; CHEN, Yuting; SHEN, Beijun. *Infer: step-by-step regex generation via chain of inference*. In: *38th IEEE/ACM International Conference on Automated Software Engineering (ASE)*. IEEE, 2023. p. 1505–1515.

¹⁹ KOSTINA, Arina; DIKAIKOS, Marios D.; STEFANIDIS, Dimosthenis; PALLIS, George Pallis. *Large language models for text classification: case study and comprehensive review*. 2025. Available at: <<https://arxiv.org/abs/2501.08457>>.

²⁰ AGARWAL, Basant; MITTAL, Namita. *Text classification using machine learning methods-a survey*. In: *Proceedings of the Second International Conference on Soft Computing for Problem Solving (SocProS 2012)*. New Dehli: Springer, 2014. p. 701-709.

²¹ WAGH, Vedangi; KHANDVE, Snehal; JOSHI, Isha; WANI, Apurva; KALE, Geetanjali; JOSHI, Raviraj. *Comparative study of long document classification*. In: *TENCON 2021-2021 IEEE Region 10 Conference (TENCON)*. 2021. p. 732–737.

these documents using both the expert-defined regex-based categories and natural language descriptions of the same categories. We compared its performance against that of a regex-based algorithm built using the same expert-designed expressions.

We selected a real-world organization, involving actual domain experts and a genuine need for automating the document classification task, in order to ensure the external validity of the evaluation and enhance the applicability of the results to real-world settings.²² Additionally, since the selected organization is an Argentinian public body, all experiments were conducted in Spanish, including the documents and the prompts provided to the language model.

2.1 Dataset construction

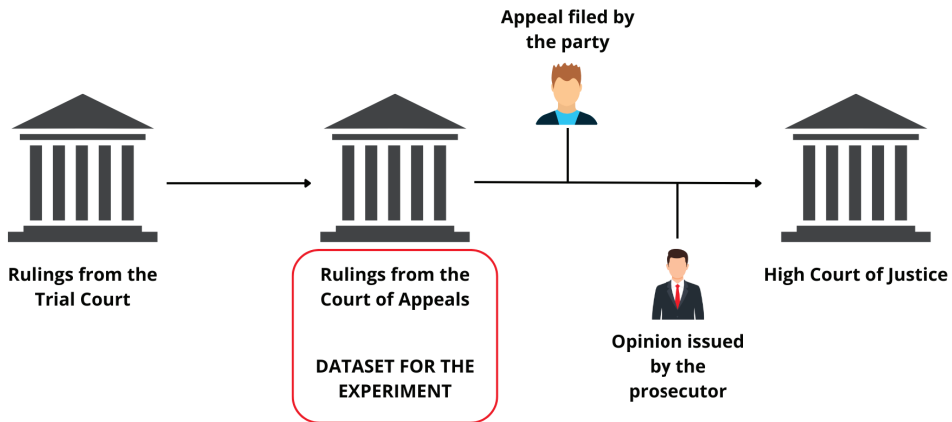
To construct the dataset, we collected 223 judicial rulings related to public employment issued by the Court of Appeals in Administrative, Tax and Regulatory Matters of the City of Buenos Aires, Argentina (Cámara de Apelaciones en lo Contencioso, Administrativo y Tributario de la Ciudad Autónoma de Buenos Aires (CABA)).

Appeals may be filed by the parties against these rulings, and the Prosecutor's Office (Fiscalía General Adjunta en lo Contencioso, Administrativo y Tributario de la Ciudad Autónoma de Buenos Aires (CABA)) must issue an opinion assessing whether the appeal should be admitted or rejected before the case is forwarded to the High Court of Justice (Tribunal Superior de Justicia), as shown in Figure 1. These rulings may fall under one of four thematic categories, making it essential for the Prosecutor to identify the relevant thematic axis upon receiving the appeal in order to ensure proper legal analysis and processing.

These four thematic categories account for approximately 80% of the cases requiring an opinion from the Prosecutor. Therefore, automating the classification process can significantly reduce the time and effort required from the public official and their supporting staff.

²² For further discussion on external validity, see: WEIDINGER, Laura; et al. Sociotechnical safety evaluation of generative AI systems. 2023. Available at: <<https://arxiv.org/abs/2310.11986>>; PASKOV, Patricia; BERGLUND, Lukas; SMITH, Everett; SODER, Lisa. *GPAI evaluations standards taskforce: towards effective AI governance*. 2024. Available at: <<https://arxiv.org/abs/2411.13808>>; BIDERMAN, Stella Biderman, et al. *Lessons from the trenches on reproducible evaluation of language models*. 2024. Available at: <<https://arxiv.org/abs/2405.14782>>; REUEL, Anka; et al. Open problems in technical AI governance. 2024. Available at: <<https://arxiv.org/abs/2407.14981>>; BURDEN, John. *Evaluating AI evaluation: perils and prospects*. 2024. Available at: <<https://arxiv.org/abs/2407.09221>>.

Figure 1 – The figure illustrates the stage of the judicial process from which the rulings included in the dataset are extracted



The 223 rulings were manually annotated by expert staff from the Prosecutor's Office. This manual annotation by legal experts ensured high-quality labeling. Furthermore, since the rulings are not publicly available online, the risk of data contamination during the LLM training is avoided.

The four categories into which the rulings were classified were: (1) Didactic Materials (Material Didactico)–39 rulings; (2) “Franqueros”–7 rulings, (3) “Fonaindo”–67 rulings, and (4) Critical Activities (Actividad Critica)–110 rulings.

2.2 Task configuration and instructions

Basically, the task of the LLM is to classify a document into a category based on whether the majority of patterns associated with that category are present in the text. To this end, we applied prompt engineering techniques–chain-of-thoughts (CoT)–and incorporated widely recommended components, such as role specification, output formatting, and contextual information.²³ The prompt used in this evaluation was structured as follows (the full prompt is contained in Appendix 1 A):

- (1) Role-playing:** Also referred to as the “Persona” theme. We provide the role's description in the prompt to make the LLMs adopt a specific perspective to solve the given task. In this case, we ask the LLM to act as an expert in document classification based on pattern recognition within the text.

²³ SAHOO, Pranab; SINGH, Ayush Kumar; SAHA, Sriparna; JAIN, Vinija; MONDAL, Samrat; CHADHA, Aman. *A systematic survey of prompt engineering in large language models: techniques and applications*. 2024. Available at: <<https://arxiv.org/abs/2402.07927>>.

- (2) **Input provided:** From the outset, we informed the LLM of the type of input it would receive along with the specific task instructions, which included: (1) a document to be classified—in this case, a judicial ruling; (2) the four possible categories to which the ruling could belong; and (3) a list of patterns associated with each category.
- (3) **Context:** We clarified that the document to be classified is a second-instance judicial decision concerning public sector employment, issued by an Argentine court.
- (4) **Steps:** To perform the classification task, we employed the Chain-of-Thought (CoT) prompting technique, guiding the language model to follow a four-step process:
- Pattern Identification: Analyze the document and identify which patterns from each category are present
 - Relative Coverage Calculation: Since the number of patterns varied across categories, we asked the model to compute the proportion of patterns detected relative to the total number of patterns in each category, rather than relying on absolute counts.
 - Document Classification: Assign the document to the category with the highest percentage of pattern coverage, regardless of which category had the highest absolute number of pattern matches.
 - Classification Justification: Present the results as a checklist, specifying which patterns were identified, the category to which they belong, and the percentage of total patterns detected for each category.
- (5) **Categories:** The four categories mentioned above: (1) Didactic Materials (Material Didactico); (2) “Franqueros”, (3) “Fonaindo”, and (4) Critical Activities (Actividad Cr.tica). In the prompt, we explicitly instructed that the provided document had to be classified into one of these four categories, with no option to leave it unclassified.
- (6) **Output format:** We instructed the LLM to generate its output in a checklist format, listing all the patterns associated with each category and indicating, for each case, which patterns were identified and which were not. This design aims to enhance the explainability of the classification process, allowing users to easily verify the actual presence or absence of the patterns in the source text.
- (7) **Warnings:** We instructed the language model to avoid repeating the full prompt in its output and to explicitly identify and analyze all relevant patterns.
- (8) **Patterns for each category:** The model was provided with a list of patterns corresponding to each category. These patterns were of two types: natural language expressions and regular expressions (regex).

The regular expressions were manually identified, created, and validated by legal experts from the Prosecutor's Office as the most representative expressions for each category. The number of patterns per category was as follows: Didactic Material—6 regular expressions and 3 natural language expressions; Fonaindo—13 regular expressions and 2 natural language expressions; Franqueros—7 regular expressions and 1 natural language expression; and Critical Activity—13 regular expressions and 2 natural language expressions.

(9) The judicial ruling.

2.3 Baselines

We compared ChatGPT's performance in classifying documents based on pattern matching with two baselines: (1) a classification task using natural language descriptions of the categories, and (2) a regex-based algorithm.

In the first baseline, the prompt design followed a structure similar to that described in Section 2.2, but instead of categories defined by patterns, we used natural language descriptions. The warnings and contextual information provided in the prompt were identical to those used in the pattern-based task. We also incorporated role-playing, instructing the model to act as an expert in document classification. The categories were then listed, each accompanied by a brief (approximately five-line) natural language description, written by the same domain experts who had designed the regex patterns for the other task. The full prompt used for this baseline is included in Appendix 1 B.

3 Experiments and results

Figure 2 presents the results of the comparison between three classification approaches applied to legal rulings: (1) ChatGPT using expert-defined regex-based category prompts (ChatGPT Regex-based), (2) ChatGPT using natural language descriptions of the same categories (ChatGPT NL), and (3) the standalone regex-based algorithm. Performance was evaluated across four metrics: Precision, Recall, F1 Score, and Accuracy. The regex algorithm consistently outperforms ChatGPT across all metrics, although the differences are relatively modest. Notably, ChatGPT—despite being slightly less accurate—exhibits strong and well-balanced performance, suggesting its viability as an alternative classification tool, particularly due to its accessibility and ease of use.

Figure 3 illustrates the time required for document classification by the regular expression-based algorithm and ChatGPT Regex-based. While the former performs the task almost instantaneously—within a fraction of a second—ChatGPT requires

significantly more time. This substantial difference reflects the computational complexity of generating outputs in LLMs, especially when processing long documents and performing multi-step reasoning.

	RegEx Algorithm	ChatGPT Regex-based	ChatGPT NL
Precision	1,00	0,94	0,97
Recall	1,00	0,93	0,97
F1 Score	1,00	0,93	0,97
Accuracy	1,00	0,93	0,97

Figure 2 – Results comparing ChatGPT performance of document classification and the RegEx-based algorithm

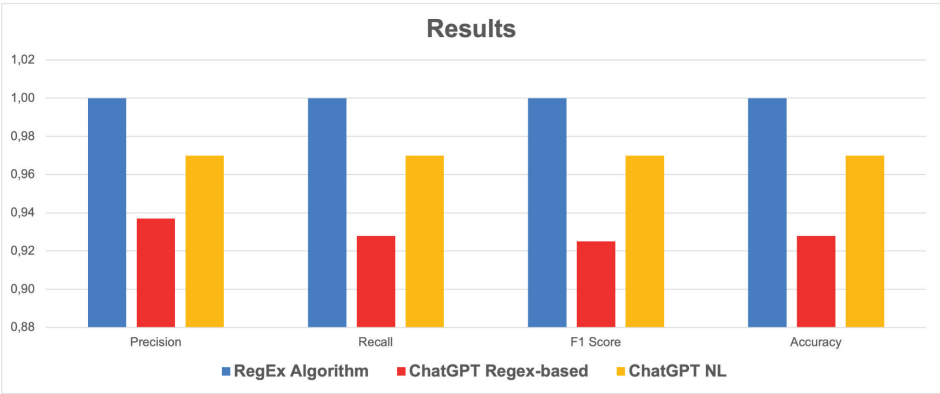
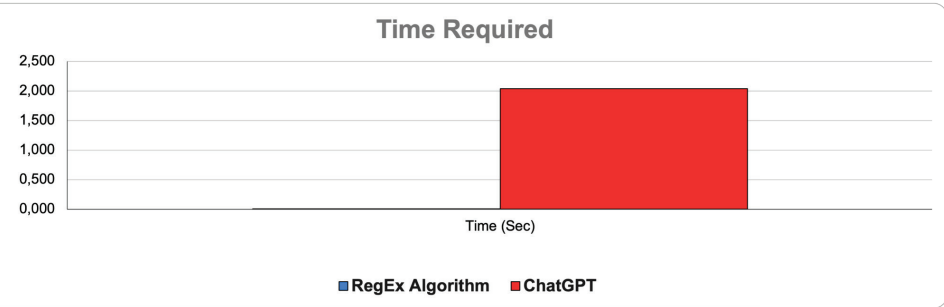


Figure 3 – The figure compares the inference time between the regex-based classification method and ChatGPT



We also conducted a qualitative analysis of the model’s outputs in the regex-matching task to assess whether it followed all the instructions provided in the prompt. We found that, in most cases, the model failed to fully comply with key aspects of the task. Specifically, it often omitted the analysis of all regular

expressions across all categories in the output. It also misreported the number of patterns per category—for instance, stating that the “Fonaindo” category contained 9 regex patterns when it actually included 13. Furthermore, the model sometimes miscalculated pattern coverage: in some cases, it claimed 100% pattern match in one category but ultimately assigned the document to a different category with only 80% match. Although the final classification was correct, the reasoning presented was inconsistent and rendered the output non-explainable to the user. Additionally, in no instance did the model analyze or report on the presence of natural language patterns, despite having been explicitly instructed to do so.

4 Limitations and future work

First, since the entire evaluation was conducted in Spanish, future work could explore whether model performance and output explainability improve when the task is performed in English, given that LLMs generally exhibit stronger performance in that language.²⁴ ²⁵ Additionally, our dataset is limited in scope; extending it to encompass a broader range of domains—such as healthcare—as well as tasks involving multi-label classification would enhance the generalizability and robustness of our findings. Further research could also include the evaluation of additional LLMs, such as Claude or Gemini, to assess the consistency of results across models. Another important limitation relates to prompt sensitivity: outcomes may be influenced by the specific prompt design.²⁶

5 Conclusion

In this paper, we evaluated the performance of ChatGPT in document classification using regular expressions, and compared it to two alternatives: ChatGPT prompted with natural language descriptions of the categories, and a traditional regex-based algorithm. While the language model still faces important limitations in document

²⁴ ROHERA, Pritika; GINIMAV, Chaitrali; SAWANT, Gayatri; JOSHI, Raviraj. *Better to ask in english? Evaluating factual accuracy of multilingual LLMs in english and low-resource languages*. Available at: <<https://arxiv.org/abs/2504.20022>>.

²⁵ SCHUT, Lisa; GAL, Yarin; FARQUHAR, Sebastian. *Do multilingual LLMs think in english? 2025*. Available at: <<https://arxiv.org/abs/2502.15603>>.

²⁶ See also: PEREZ, Ethan; et al. Discovering language model behaviors with model-written evaluations. In: Findings of the Association for Computational Linguistics. Toronto: ACL, 2023. p. 13387–13434; SCLAR, Melanie; CHOI, Yejin; TSVETKOV, Yulia; SUHR, Alane. *Quantifying language models' sensitivity to spurious features in prompt design or: how I learned to start worrying about prompt formatting*. 2023. Available at: <<https://arxiv.org/abs/2310.11324>>; FRONTIER MODEL FORUM. *Issue brief: Early best practices for frontier ai safety evaluations*. 2024. Available at: <<https://www.frontiermodelforum.org/updates/early-best-practices-for-frontier-ai-safety-evaluations/>>; MCINTOSH, Timothy R., et al. Inadequacies of large language model benchmarks in the era of generative artificial intelligence. *IEEE Transactions on Artificial Intelligence*, 2025. Available at: <<https://arxiv.org/abs/2402.09880>>.

classification— particularly in terms of consistency, explainability, and alignment with detailed instructions— it offers a significant advantage in terms of accessibility. Unlike traditional methods that require technical expertise to implement and maintain, LLMs enable non-technical users to automate complex classification tasks with minimal setup. This shift has the potential to democratize access to automation tools within public administration and other domains, expanding the reach and impact of document processing technologies even as technical performance continues to evolve.

6 Acknowledgments

We would like to thank Facundo Nieto, and María Victoria Carro, Mario Alejandro Leiva, and Denise Alejandra Mester from the FAIR team for their assistance with the technical aspects of this paper and their advice on designing a reliable evaluation.

References

- ABUKAR AHMED, Abdinor; KHALIF ALI, Mohamed. User-centric adoption of democratized generative AI: Focus on human-machine interaction and overcoming challenges. *International Journal of Engineering Trends and Technology*, v. 72, issue 9, p. 78-95, 2024.
- AGARWAL, Basant; MITTAL, Namita. Text classification using machine learning methods-a survey. In: *Proceedings of the Second International Conference on Soft Computing for Problem Solving* (SocProS 2012). New Dehli: Springer, 2014. p. 701-709.
- BHAMBRI, Gaurav. Information overload in business organizations and entrepreneurship: An analytical review of the literature. *Business Information Review*, v. 38, issue 4, 2021, p. 193–200.
- BIDERMAN, Stella Biderman, et al. *Lessons from the trenches on reproducible evaluation of language models*. 2024. Available at: <<https://arxiv.org/abs/2405.14782>>; REUEL, Anka; et al. Open problems in technical AI governance. 2024. Available at: <<https://arxiv.org/abs/2407.14981>>.
- BURDEN, John. *Evaluating AI evaluation: perils and prospects*. 2024. Available at: <<https://arxiv.org/abs/2407.09221>>.
- CABELLERO, William N.; JENKINS, Phillip R. On large language models in national security applications. *Stat*, v. 14, issue 2, e70057, 2025.
- FRONTIER MODEL FORUM. *Issue brief: Early best practices for frontier ai safety evaluations*. 2024. Available at: <<https://www.frontiermodelforum.org/updates/early-best-practices-for-frontier-ai-safety-evaluations/>>.
- JIANG, Shuo; HU, Jie; MAGEE, Christopher L.; LUO, Jianxi. Deep learning for technical document classification. *IEEE Transactions on Engineering Management*, v. 71, p. 1163-1179, 2022.
- KLEIN, L. K; EARL, E.; CUNDICK, D. Reducing information overload in your organization. *Harvard Business Review*, 2023.
- KOSTINA, Arina; DIKAIKOS, Marios D.; STEFANIDIS, Dimosthenis; PALLIS, George Pallis. *Large language models for text classification: case study and comprehensive review*. 2025. Available at: <<https://arxiv.org/abs/2501.08457>>.
- KOWSARI, Kamram; JAFARI MEIMANDI, Kiana; HEIDARYSAFA, Mojtaba; MENDU, Sanjana; BARNES, Laura; BROWN, Donald. Text classification algorithms: A survey. *Information*, v. 10, issue 4, p. 150, 2019.

- LI, Yinghao; RAMPRASAD, Rampi; ZHANG, Chao. *A simple but effective approach to improve structured language model output for information extraction*. 2024. Available at: <<https://arxiv.org/abs/2402.13364>>.
- MAHONEY, Christian; GRONVALL, Peter; HUBER-FLIFLET; ZHANG, Jianping. Explainable text classification techniques in legal document review: locating rationales without using human annotated training text snippets. In: *2022 IEEE International Conference on Big Data (Big Data)*. IEEE, p. 2044-2051, 2022.
- MCINTOSH, Timothy R., et al. Inadequacies of large language model benchmarks in the era of generative artificial intelligence. *IEEE Transactions on Artificial Intelligence*, 2025. Available at: <<https://arxiv.org/abs/2402.09880>>.
- NATIONAL JUDICIAL STATISTICS SYSTEM (SNEJ). *Annual Report*: Technical Report. National Judicial Statistics System (SNEJ): Argentina, 2022.
- PALANIVINAYAGAM, Ashokkumar; ZIAD EL-BAYEH, Claude; DAMAŠEVICIUS, Robertas. Twenty years of machine-learning-based text classification: A systematic review. *Algorithms*, v. 16, issue 5, p. 236, 2023.
- PASKOV, Patricia; BERGLUND, Lukas; SMITH, Everett; SODER, Lisa. *GPAI evaluations standards taskforce: towards effective AI governance*. 2024. Available at: <<https://arxiv.org/abs/2411.13808>>.
- PEÑA ALMANSA, Alejandro; et al. Leveraging large language models for topic classification in the domain of public affairs. *Lecture Notes in Computer Science* (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics), 2023.
- PEREZ, Ethan; et al. Discovering language model behaviors with model-written evaluations. In: *Findings of the Association for Computational Linguistics*. Toronto: ACL, 2023. p. 13387–13434.
- ROHERA, Pritika; GINIMAV, Chaitrali; SAWANT, Gayatri; JOSHI, Raviraj. *Better to ask in english?* Evaluating factual accuracy of multilingual LLMs in english and low-resource languages. Available at: <<https://arxiv.org/abs/2504.20022>>.
- SAHOO, Pranab; SINGH, Ayush Kumar; SAHA, Sriparna; JAIN, Vinija; MONDAL, Samrat; CHADHA, Aman. *A systematic survey of prompt engineering in large language models: techniques and applications*. 2024. Available at: <<https://arxiv.org/abs/2402.07927>>.
- SCHUT, Lisa; GAL, Yarin; FARQUHAR, Sebastian. *Do multilingual LLMs think in english?* 2025. Available at: <<https://arxiv.org/abs/2502.15603>>.
- SCLAR, Melanie; CHOI, Yejin; TSVETKOV, Yulia; SUHR, Alane. *Quantifying language models' sensitivity to spurious features in prompt design or: how I learned to start worrying about prompt formatting*. 2023. Available at: <<https://arxiv.org/abs/2310.11324>>.
- SEGER, Elizabeth. *What do we mean when we talk about "AI democratisation"*. 2023. Available at: <<https://www.governance.ai/analysis/what-do-we-mean-when-we-talk-about-ai-democratisation>>.
- SUN, Xiaofei; LI, Xiaoya; LI, Jiwei; WU, Fei; GUO, Shangwei; ZHANG, Tianwei; WANG, Guoyin. *Text classification via large language models*. 2023. Available at: <<https://arxiv.org/abs/2305.08377>>.
- SUPREME COURT OF ARGENTINA. *Statistics For The Year 2024*: Technical report. Supreme Court of Argentina, 2024.
- WAGH, Vedangi; KHANDVE, Snehal; JOSHI, Isha; WANI, Apurva; KALE, Geetanjali; JOSHI, Raviraj. Comparative study of long document classification. In: *TENCON 2021-2021 IEEE Region 10 Conference (TENCON)*. 2021. p. 732–737.
- WEIDINGER, Laura; et al. Sociotechnical safety evaluation of generative AI systems. 2023. Available at: <<https://arxiv.org/abs/2310.11986>>.

YE, Junjie; XU, Nuo; WANG, Yikun; ZHOU, Jie; ZHANG, Qi; GUI, Tao; HUANG, Xuanjing. *LLM-DA: Data augmentation via large language models for few-shot named entity recognition*. 2024. Available at: <<https://arxiv.org/abs/2402.14568>>.

YOON, Sungwook Yoon; et. al. Digital innovation in public administration through intelligent public sector automation (IPSA): Strategies and challenges. *Journal of Multimedia Information System*, v. 11, issue 4, p. 249-260, 2024.

ZHANG, Shuai; GU, Xiaodong; CHEN, Yuting; SHEN, Beijun. Infere: step-by-step regex generation via chain of inference. In: *38th IEEE/ACM International Conference on Automated Software Engineering (ASE)*. IEEE, 2023. p. 1505–1515.

Informação bibliográfica deste texto, conforme a NBR 6023:2025 da Associação Brasileira de Normas Técnicas (ABNT):

CORVALÁN, Juan Gustavo *et al.* Predictive pattern analysis in administrative and tax litigation through generative AI and regular expressions. *International Journal of Digital Law*, Belo Horizonte, v. 7, e701, 2026. DOI: 10.47975/ijdl.v7.1305.

APPENDIX A

Regex-Based Prompt Provided to ChatGPT

Objetivo: Actúa como un experto en clasificar documentos mediante la detección de patrones en el texto. Entrada proporcionada:

Predictive Pattern Analysis in Administrative and Tax Litigation Through Generative AI and Regular Expressions

- Un documento que debes clasificar.
- Una lista de 4 categorías posibles: 1. Material didáctico; 2. Franqueros; 3. Fonaindo y 4. Actividad crítica. El documento necesariamente debe clasificarse en una de ellas.
- Los patrones que caracterizan a cada categoría, los cuales se dividen en: Expresiones en lenguaje natural Expresiones regulares.

Contexto: El documento a clasificar es una sentencia de segunda instancia sobre empleo público emitida por la Cámara de Apelaciones en lo Contencioso, Administrativo y Tributario de CABA, Argentina. Tu tarea es analizarlo para identificar qué patrones de cada categoría aparecen en el texto y asignarlo a la categoría con mayor coincidencia.

Pasos que deben ejecutarse para realizar la tarea:

1. Identificación de patrones: Analiza el documento e identifica qué patrones de cada categoría están presentes, considerando tanto expresiones en lenguaje natural como expresiones regulares.
2. Cálculo de cobertura relativa: Para cada categoría, determina el porcentaje de patrones detectados sobre el total de patrones posibles en esa categoría, en lugar de contar únicamente la cantidad absoluta. Incluye tanto las expresiones regulares como las expresiones en lenguaje natural.
3. Clasificación del documento: El documento será asignado a la categoría en la que se detecte el mayor porcentaje de patrones, no necesariamente la que tenga más patrones en términos absolutos.
4. Explicación de la clasificación: Presenta los resultados en formato de checklist, especificando qué patrones fueron encontrados, en qué categoría, y qué porcentaje del total de patrones de cada categoría fueron detectados.

Categorías:

- Actividad Crítica
- Fonaindo
- Franqueros
- Material Didáctico

Formato de Salida: [“Checklist de patrones detectados en cada categoría: – Diferencias salariales: [V Patrón 1, X Patrón 2, V Patrón 3] – Fonaindo: [X Patrón 1, X Patrón 2, V Patrón 3] – Franqueros: [V Patrón 1, V Patrón 2, X Patrón 3] – Material Didáctico: [X Patrón 1, X Patrón 2, X Patrón 3] Conclusión: De (X) patrones proporcionados para la categoría (nombre de la categoría), se detectaron (Y) patrones en el documento, que son los siguientes: (lista de patrones encontrados). Por lo tanto, se concluye que el documento pertenece a la categoría (categoría asignada).”]

Advertencias:

No repitas esta solicitud en la respuesta, ejecútala directamente. Deberás mencionar y analizar la totalidad de las expresiones regulares y expresiones en lenguaje natural detectadas, independientemente de que pertenezcan o no a la categoría final clasificada.

Patrones que caracterizan a cada categoría:

Material didáctico: Expresiones regulares (6 en total)

- 1) (?i)\bMaterial\s+Didáctico\s+mensual\b
- 2) (?i)\bMaterial\s+Didáctico\b(?:\s+(mensual|Bicentenario))
- 3) (?i)\bMaterial\s+Didáctico\s+Bicentenario\b 4) (?i)\bCarácter\s+remunerativo\s+de\s+los\s+suplementos\s+denominados\s*\(\s*Material\s+Didáctico,\s*s*Material\s+Didáctico\s+mensual,\s*s*Material\s+Didáctico\s+Bicentenario\s*\)\s+liquidados\s+en\s+sus\s+salarios\s+con\s+los\s+c[ó]digos\s+(093,\s*397,\s*s*493)\b
- 5) (?i)\bInconstitucionalidad\s+de\s+los\s+decretos\s+(N[oº]?s*)?(751\04,\s*s*547\2005,\s*s*1294\07)\b
- 6) (?i)\bdeclarar\s+el\s+carácter\s+remunerativo\s+de\s+los\s+suplementos\s+denominados\s*[\s“”]\s*Material\s+Didáctico\s*[\s“”]\s+o\s*[\s“”]\s*s*Material\s+Didáctico\s+Mensual\s*[\s“”]\s+y\s+ordenar\s+al\s+GCBA\s+liquidarlos\s+con\s+dicho\s+carácter\s+incluirllos\s+en\s+la\s+base\s+de\s+cálculo\s+de\s+SAC\s+y\s+abonar\s+a\s+la\s+actora\s+las\s+diferencias\s+salariales\s+devengadas\s+por\s+estos\s+conceptos\s+por\s+los\s+períodos\s+no\s+prescriptos\s+con\s+intereses\b

Expresiones en lenguaje natural (3 en total)

Predictive Pattern Analysis in Administrative and Tax Litigation Through Generative AI and Regular Expressions

- 1) “Material Didáctico Mensual”
- 2) “Material Didáctico”
- 3) “Material Didáctico Bicentenario”

Fonaindo: Expresiones regulares (13 en total)

- 1)^(?i)Fondo\s+Nacional\s+de\s+Incentivo\s+Docente\s+FONAINDO\s+incorrecta\s+liquidaci[oó]n\s+del\s+c[oó]digo\s+399\$ 2)^(?i)docentes\s+dependientes\s+de\s+la\s+Secretar[ií]a\s+de\s+Educaci[oó]n\$
- 3)^(?i)Gobierno\s+de\s+la\s+Ciudad\s+de\s+Buenos\s+Aires\s+diferencias\s+salariales\s+provenientes\s+de\s+la\s+incorrecta\s+liquidaci[oó]n\$
- 4)^(?i)Fondo\s+Nacional\s+de\s+Incentivo\s+Docente\s+Fonaindo\$
- 5)^(?i)c[oó]digo\s+399\$
- 6)^(?i)personal\s+docente\s+en\s+actividad\s+y\s+retirado\$ 7)^(?i)diferencias\s+salariales\s+resultantes\s+del\s+modo\s+(err[oó]neolerroneo)\s+del\s+que\s+se\s+devengaron\s+sus\s+haber\$
- 8)^(?i)car[aá]cter\s+remunerativo\s+y\s+bonificable\s+del\s+suplemento\$ 9)^(?i)car[aá]cter\s+remunerativo\s+del\s+suplemento\$
- 10)^(?i)car[aá]cter\s+bonificable\s+del\s+suplemento\$
- 11)^(?i)Fondo\s+Nacional\s+de\s+Incentivo\s+Docente\s+\(FONAINDO\)\$
- 12)^(?i)Diferencias\s+salariales\$ 13)^(?i)Sueldo\s+Anual\s+Complementario\s+\(SAC\)\$

Expresiones en lenguaje natural (2 en total)

- 1) “Bruno, Marcelo José y otros”
- 2) “Falta de legitimación pasiva”

Franqueros: Expresiones regulares (7 en total)

- 1) \b[eE]nfermer[oa]s?\s+franquer[oa]s?\b
- 2) \b[jJ]ornada\s+laboral(es)?\b
- 3) \b[sS]eis\s+\(?6\)?\s+horas\s+diarias\b
- 4) \b[tT]reinta\s+\(?30\)?\s+horas\s+semanales\b
- 5) \b[áa]rea\s+insalubre\b
- 6) \b[IL]ey\.\?s*n?[ºº]?\s*24\.\?004\b
- 7) \b[oO]rdenanza\.\?s*n?[ºº]?\s*40\.\?403\b

Expresiones en lenguaje natural (1 en total)

- 1) “Resolución conjunta 499/MHFGC/20”

Actividad Crítica: Expresiones regulares (13 en total)

- 1) \bSuplemento\s+especial\s+por\s+área\s+crítica\b
- 2) \bActividades\s+en\s+áreas\s+críticas\b
- 3) (\bDecreto\s+(2851|2154)\89\b
- 4) \bLey\s+6035\471\b
- 5) \bderecho\s+a\s+cobro\s+del\s+suplemento\s+por\s+actividad\s+crítica\b
- 6) \bincremento\s+sostenido\s+del\s+número\s+de\s+pacientes\b
- 7) \bcarencia\s+de\s+personal\s+idóneo\b
- 8) \bderecho\s+a\s+la\s+retribución\s+justa\b
- 9) \bsuplemento\s+por\s+área\s+crítica\b
- 10) \bterapia\s+intensiva\b
- 11) \bfunción\s+crítica\b
- 12) \bordenanza\s+41455\b 13) \bcarrera\s+municipal\s+de\s+profesionales\s+de\s+la\s+salud\b

Expresiones en lenguaje natural (2 en total)

- 1) “Paz, Héctor Damián c/GCBA s/ empleo público (excepto cesantía o exoneraciones)”
- 2) “Idalgo, Sergio Fernando y otros”

APPENDIX B

Prompt with Natural Language Category Descriptions Provided to ChatGPT

Rol y contexto: Actúa como un experto en clasificar documentos. El documento a clasificar es una sentencia de segunda instancia sobre empleo público emitida por la Cámara de Apelaciones en lo Contencioso, Administrativo y Tributario de CABA, Argentina. Tu tarea es analizarlo para identificar a qué categoría pertenece el texto a partir de la descripción de las categorías que se va a proporcionar.

Categorías: Las categorías a las que puede pertenecer el documento son las siguientes. El documento necesariamente debe clasificarse en una de ellas.

- Actividad Crítica
- Fonaindo
- Franqueros
- Material Didáctico

Descripción de las categorías:

Actividad crítica: Las sentencias de esta categoría responden a apelaciones interpuestas por empleados del Gobierno de la Ciudad de Buenos Aires (GCBA) que ejercen funciones en hospitales (médicos o enfermeros, generalmente) y reclaman que se les integre al sueldo el “suplemento especial por Área Crítica”, así como también, las diferencias salariales derivadas de la falta de pago de dicho rubro. El mismo, suele encontrarse en las sentencias como “Suplemento por Actividad Crítica” y lo perciben todos aquellos empleados que trabajen en áreas críticas de los hospitales (quirófanos, unidades de cuidados intensivos (UCI), etc.).

Fonaindo: Las sentencias de esta categoría responden a apelaciones interpuestas por docentes que trabajan para el GCBA y reclaman, en este caso, que se declare como remunerativo y bonificable un suplemento de la Ley Nacional 25.053 — Fondo Nacional de Incentivo Docente (FO.NA.IN.DO, código 399)—.

Franqueros: Las sentencias de esta categoría responden a apelaciones interpuestas por empleados del GCBA (que suelen trabajar en establecimientos de salud, como enfermeros o médicos) y requieren que se les reduzcan o adecúen las jornadas de trabajo a determinada cantidad de horas, ya que la extensión excesiva es considerada insalubre. Estas peticiones suelen justificarse en las condiciones particulares de

su labor, que incluye la cobertura de francos, fines de semana, feriados, asuetos y otros días no laborables.

Material didáctico: Las sentencias de esta categoría responden a apelaciones interpuestas por docentes que trabajan para el GCBA y reclaman que se declaren como remunerativos ciertos rubros que se les pagan como no remunerativos. En el pedido que hacen, además, solicitan que se les abonen las diferencias salariales generadas por los períodos no abonados. Esos rubros que reclaman se suelen encontrar en las sentencias como “Material Didáctico”, “Material Didáctico Mensual” y “Material Didáctico del Bicentenario”.

Advertencias: No repitas esta solicitud en la respuesta, ejecútala directamente.

Informaciones adicionales

Additional information

Editores responsables	
Editor en jefe	Emerson Gabardo
Editor adjunto	Lucas Bossoni Saikali

Declaración de autoría y especificación de la contribución al artículo. <i>Statement of Authorship and Individual Contribution</i>	
Juan Gustavo Corvalán	Contribuciones: 1. Conceptualización; 4. Adquisición de fondos; 6. Metodología; 10. Supervisión; 13. Redacción – borrador original; 14. Redacción – revisión y edicin.
Sofía Tammaro	Contribuciones: 3. Análisis formal; 5. Investigación; 6. Metodología; 7. Software; 11. Validacin; 13. Redacción – borrador original.
Gisel Alvarado	Contribuciones: 1. Conceptualización; 5. Investigación; 6. Metodología; 10. Supervisión; 14. Redacción – revisión y edicin.
Luca Nicolás Forziatti Gangi	Contribuciones: 3. Análisis formal; 5. Investigación; 7. Software; 11. Validacin; 12. Visualizacin.
Carina Mariel Papini	Contribuciones: 2. Curación de datos; 5. Investigación; 6. Metodología; 11. Validacin.

Declaración de autoría y especificación de la contribución al artículo. <i>Statement of Authorship and Individual Contribution</i>	
Mariana Sánchez Caparrós	Contribuciones: 2. Curación de datos; 5. Investigación; 11. Validación; 14. Redacción – revisión y edición.
Agostina Celeste Jara Rey	Contribuciones: 2. Curación de datos; 5. Investigación; 11. Validación.
Florencia Paola Croci	Contribuciones: 2. Curación de datos; 5. Investigación; 11. Validación.
Lola Ramos Pereyra	Contribuciones: 2. Curación de datos; 5. Investigación; 11. Validación; 12. Visualización.
Melania Gadea	Contribuciones: 2. Curación de datos; 5. Investigación; 11. Validación.
Luciano Della Via	Contribuciones: 2. Curación de datos; 5. Investigación; 11. Validación.