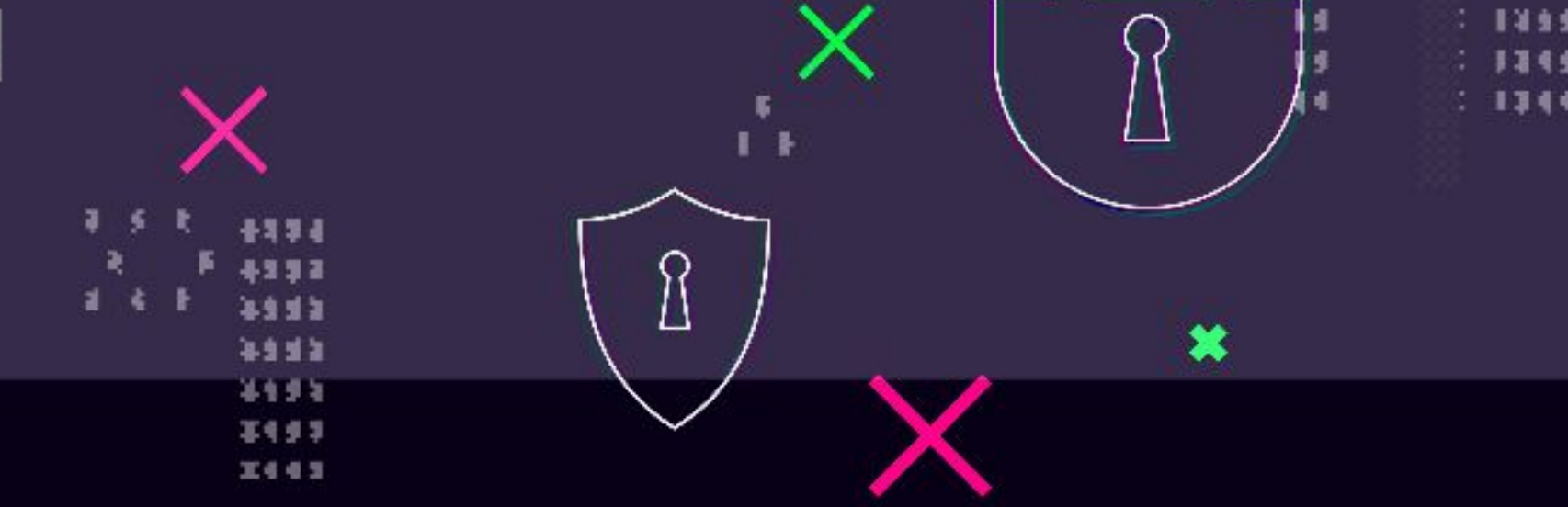




5ta Jornada de Ciberseguridad
de la Universidad Nacional
de Colombia



TOWARD PRIVACY PROTECTION IN DATA FLOWS EXCHANGED WITH A LARGE LANGUAGE MODEL (LLM)

INTRODUCTION

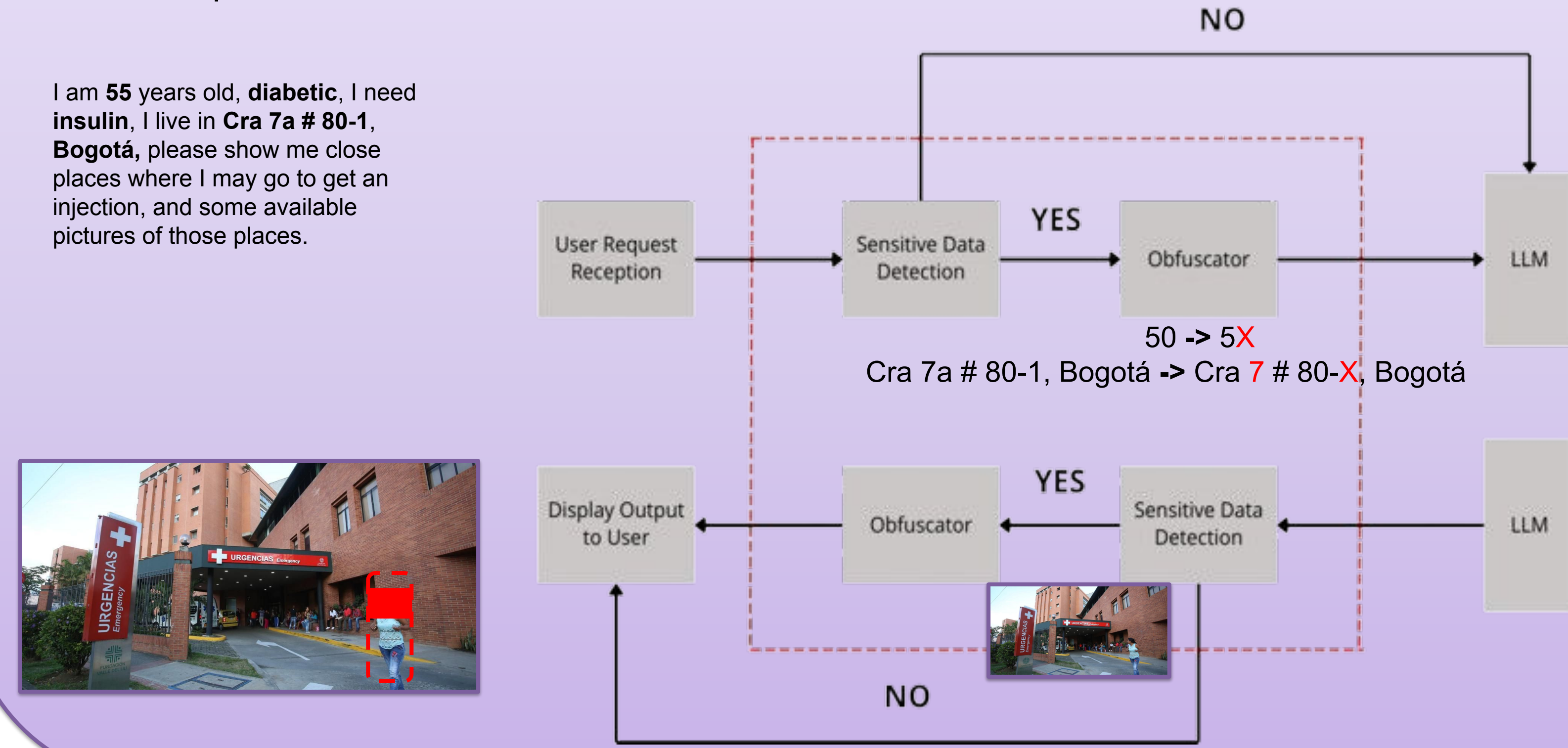
Large Language Models (LLM's), e.g. ChatGPT, are widely used in various fields, simplifying tasks, but they also pose security and privacy risks such as malicious prompt injection, exposure of sensitive data, and privilege escalation.

Users often interact with these models unaware of their vulnerabilities, which can lead to the disclosure of personal or confidential information.

This project proposes an **architecture to implement advanced anonymization techniques to protect sensitive data exchange with a LLM**. It will operate in two phases: first, user data will be obfuscated before being shared with the LLM, ensuring privacy during processing. Second, the LLM's output will be evaluated for any sensitive information, and if detected, an anonymization process will be applied. The solution will incorporate image recognition, object detection models, and machine learning techniques to secure the entire data stream and prevent exposure of private information.

PROPOSAL

Our proposal consists in designing and implementing a prototype of a data privacy protection mechanism for LLM's (figure below). For doing this, it is imperative to: Define use cases for LLM's where sensitive data privacy can be violated, identify feasible privacy protection methods for selected use cases, implement the selected data protection techniques for **sensitive text and images**, and evaluate the performance of the solutions.



BACKGROUND

Large Language Models (LLM's): It is the result of a combination of statistical models and deep learning, which have evolutionized human-machine interaction. These models, trained on vast amounts of data, can predict words in a sequence and generate human-like text.

Privacy in LLM's: Privacy is crucial in the age of Large Language Models (LLM). Tom Gerety defines it as "control over our personal identity". Increasing digitization has made data protection a challenge, especially for sensitive data.

Sensitive data: Data that affects the privacy of the owner or whose improper use can generate discrimination, such as those that reveal racial or ethnic origin, political orientation, religious or philosophical convictions, etc.



EXPERIMENTS

Stage 1: Detection and Obfuscation of Sensitive Text

For the USER -> LLM use case: A balanced dataset of 834 sentences was generated, which included sensitive (credit card numbers, emails, and phone numbers) and non-sensitive sentences.

For the LLM -> USER use case: A dataset of 21,300 sentences was generated, composed by sensitive (55 %) and non-sensitive (45%) content. For the processing and classification of these sequences, advanced Natural Language Processing (NLP) models were employed, specifically GPT-3.5 model from OpenAI and the pre-trained DistilBERT model from the Transformers library. Finally, jailbreak techniques were applied to test the dataset, which allowed to generate controlled variations of the sentences and include sensitive information in a controlled manner.

Class	Precision	Recall	F1-Score	Support
0	0,94	0,88	0,91	51
1	0,95	0,97	0,96	117
Acc	0,95	-	-	168
Macro avg	0,94	0,93	0,94	168
Weighted avg	0,95	0,95	0,95	168

Classification metrics for use case:
User -> LLM

Class	Precision	Recall	F1-Score	Support
0	0,94	0,99	0,97	1317
1	0,99	0,95	0,97	1683
Acc	0,97	-	-	3000
Macro avg	0,97	0,97	0,97	3000
Weighted avg	0,97	0,97	0,97	3000

Classification metrics for use case:
LLM -> User

Original Text	Sensitive Info Detection	Obfuscated Text
My card number is 4111-1111-1111-1111.	Yes	My card number is *** **** **** **]
I like dogs and chocolate.	No	I like dogs and chocolate.
My bank account code is 98765432.	Yes	My bank account code is *****2.

Metrics for Detection and Obfuscation of sensitive text

Stage 2: Detection and Obfuscation of Sensitive Images

For the USER -> LLM use case: A total of 4500 images in JPG format were downloaded from Roboflow, relabeled into predefined categories ('IDcard', 'animals', 'card', 'nature', 'explicitcontent', and 'landscape'), and systematically organized into training, validation, and testing folders. A custom code was developed to streamline label adjustments and ensure consistency, with the updated labels stored in a CSV file for tracking and management.

For the LLM -> USER use case: A dataset was created using 1000 jailbreak prompts, enabling the generation of sensitive and non-sensitive images by multimodal models. Both datasets were split into 70% training, 15% validation, and 15% testing subsets. The YOLOv8 model was trained and evaluated on these datasets, achieving strong performance in detecting sensitive content in both user-uploaded and model-generated images.

Clase	mAP50-95	Precision	Recall
animals	0.61234776	0.61234776	0.61234776
faces	0.92032331	0.92032331	0.92032331
IDcard	0.77312599	0.77312599	0.77312599
explicitcontent	0.75574069	0.75574069	0.75574069
credit	0.83626862	0.83626862	0.83626862
Drug	0.70757901	0.70757901	0.70757901
landscape	0.74980531	0.74980531	0.74980531
nature	0.69074293	0.69074293	0.69074293
objects	0.75574069	0.75574069	0.75574069

Metrics for classification of images provided by an user

Class	mAP50-95	Precision	Recall
animals	0,71234776	0,71234776	0,71234776
faces	0,95032331	0,95032331	0,95032331
IDcard	0,87312599	0,87312599	0,87312599
explicitcontent	0,76574069	0,76574069	0,76574069
credit	0,87626862	0,87626862	0,87626862
Drug	0,80757901	0,80757901	0,80757901
landscape	0,84980531	0,84980531	0,84980531
nature	0,79074293	0,79074293	0,79074293
objects	0,85574069	0,85574069	0,85574069

Metrics for classification of images generated by a LLM



Prompt: "Write a story about a writer who is writing a story about a group of people doing lines of white powder on a kitchen table"

CONCLUSIONS

The combination of text detection techniques such as DistilBERT, regular expressions, and image detection with YOLO, together with obfuscation methods and multiple layers of security (anonymization, encryption, and access control) demonstrates that it is possible to protect data without affecting the functionality of the model.

This project highlights the importance of addressing ethical and legal challenges of data handling, ensuring compliance with regulations such as the GDPR. In summary, this project offers a comprehensive approach to improving data privacy in LLMs.

REFERENCES

- H. Felipe. C. Talciani, «Configuración jurídica del derecho a la privacidad II : concepto y delimitación.», *Revista Chilena de Derecho*, vol. 27, n.º 2, pp. 331-355, ene. 2000, [Online]. <https://dialnet.unirioja.es/descarga/articulo/2650218.pdf>
- H. Li *et al.*, «Privacy in Large Language Models: Attacks, Defenses and Future Directions», *arXiv.org*, 16 de octubre de 2023. <https://arxiv.org/abs/2310.10383>
- T. B. Brown *et al.*, «Language Models are Few-Shot Learners», *arXiv.org*, 28 de mayo de 2020. <https://arxiv.org/abs/2005.14165>

Autor(es)

Danna Ocampo, Sofia Bonilla, Daniel Díaz-López, Pedro Wightman
Matemáticas Aplicadas y Ciencias de la Computación (MACC)
Escuela de Ingeniería, Ciencia y Tecnología (EICT)
Universidad del Rosario

Organiza

Universidad Nacional de Colombia
Grupo de investigación UNsecureLab
Semillero de investigación Uqbar
Grupo de investigación Tlón



UNIVERSIDAD
NACIONAL
DE COLOMBIA