

De-identification and Anonymization of Claims Narratives Using Hybrid Neural Architectures

Sujata Rijal¹

¹Gandaki University, Department of Computational Sciences, Bagar Road, Pokhara, Nepal.

Abstract

This paper addresses the challenge of de-identifying sensitive and personal information within insurance claims narratives through the design of hybrid neural architectures that combine multiple layers of advanced language processing. The proposed methodology aims to obfuscate, remove, or replace personally identifying information while preserving the overall coherence and semantic structure of the text for further analysis. By unifying recurrent language models and attention-based transformations with external knowledge resources, this approach ensures robustness to various syntactic and contextual nuances present in real-world claims data. The work builds on the premise that a balance between contextual embedding fidelity and explicit structured constraints can yield high-quality anonymized narratives that retain their factual essence but eliminate the possibility of deducing personal attributes. Both theoretical and practical aspects of the approach are investigated, including novel representations that integrate logic-based constraints for capturing domain-specific rules, and linear algebraic mechanisms to manage large-scale embeddings efficiently. In-depth experiments carried out on de-identified insurance claim datasets confirm that a hybrid model strategy can surpass single-model baselines in terms of precision, recall, and downstream utility. The results also suggest that strategic inclusion of domain-specific lexical constraints reinforces privacy guarantees. The findings highlight a promising new direction for designing anonymization frameworks capable of working with a broad range of domain texts while offering strong theoretical guarantees of de-identification accuracy.

1. Introduction

De-identification of textual data has emerged as a critical task in both natural language processing and privacy-centric data governance domains [1]. In many industries, including finance, healthcare, and insurance, large volumes of unstructured narratives are produced daily. Such records often contain a substantial amount of personal or sensitive information [2]. Within the insurance sector, claims narratives frequently include data such as policyholder names, addresses, contact details, and other personal attributes. The risk of data re-identification has triggered stringent regulations intended to enforce anonymity safeguards [3]. Consequently, there is a growing need for sophisticated approaches that strike an optimal balance between utility and privacy when processing claims narratives. This paper explores how hybrid neural architectures can offer robust mechanisms for identifying and anonymizing sensitive information while maintaining semantic coherence. [4]

Contemporary privacy legislation stipulates that personally identifiable information must be removed or transformed in a way that renders it infeasible to determine an individual's identity. Traditional approaches to textual anonymization have ranged from simple rule-based replacements to more advanced statistical techniques [5]. However, these methods often fail to capture the complexity inherent in unstructured, domain-specific language. Complexities arise, for instance, in the use of context-dependent cues where an entity might be indirectly referenced through occupational details, location identifiers, or specialized insurance terminology. Hence, the development of an effective de-identification system

must incorporate both the lexical diversity of real-world narratives and the underlying syntactic-semantic interplay that can implicitly disclose personal details. [6]

The emergence of deep neural networks has shifted the landscape of natural language processing, enabling models to capture intricate dependencies within sequential data. Recurrent architectures, such as Long Short-Term Memory networks, have demonstrated success in entity recognition tasks [7]. Similarly, transformer-based architectures have gained traction for their ability to handle long-range dependencies more effectively than purely sequential models. Nonetheless, each of these approaches has limitations when dealing with highly specialized domains [8]. A purely recurrent model may not fully capture nuanced domain-specific patterns, while a transformer-based approach might neglect certain local context cues that remain pivotal for accurate identification of personally identifiable elements.

In light of these challenges, hybrid neural architectures hold promise [9]. By combining sequential and attention-based models, it is possible to leverage the complementary advantages of each paradigm. For example, a bidirectional recurrent backbone can capture local dependencies, while self-attention layers can focus on more global relationships [10]. Additionally, external knowledge or logic-based constraints can be integrated into the model to improve interpretability and consistency. Such constraints could help enforce domain rules, for instance ensuring that phone numbers follow particular patterns, or that address formats conform to region-specific standards. The interplay between learned representations and structured symbolic knowledge thereby undergirds a robust approach to anonymization. [11]

The present work also recognizes that de-identification must go beyond surface-level pattern extraction. In the context of insurance claims, seemingly innocuous details may lead to an individual's re-identification if combined with external data sources or background knowledge [12]. Therefore, the architectural design must incorporate a deeper representation of domain semantics. This is where the inclusion of logic statements and symbolic constraints becomes crucial [13]. By defining a set of propositional or first-order logic statements that capture the domain's privacy-critical rules, it is possible to constrain the model's outputs in a manner that a purely data-driven approach might overlook. Moreover, advanced linear algebraic methods provide mechanisms to handle large-scale embeddings and robustly unify them with symbolic constraints [14]. In effect, the hybrid system attempts to maximize recall of sensitive elements while preserving linguistic fluency.

This paper begins by laying out the theoretical foundations of hybrid architectures for text de-identification, highlighting the core principles from both neural network design and symbolic logic constraints. It then details the proposed methodology, including specialized representation layers, constraint-enforced decoding, and training procedures [15]. An extensive experimental section demonstrates the performance of the proposed system on real-world claims narratives, evaluating it against baseline methods. Metrics such as precision, recall, F1-score, and the privacy-utility trade-off are used to provide a comprehensive perspective on the system's capabilities [16]. Finally, discussion on limitations, potential improvements, and broader applicability is presented. Through this work, we aim to establish that hybrid neural architectures, reinforced by structured constraints, can offer a potent and flexible solution for de-identification and anonymization tasks in insurance claims processing. [17]

2. Theoretical Foundations

A thorough exploration of theoretical underpinnings sets the stage for designing de-identification systems that meet rigorous standards of privacy and utility. One guiding principle is the interplay between probability theory, which underlies neural network weight distributions, and logic-based constraints that capture domain-specific knowledge [18]. Both are critical when dealing with high-stakes textual data, where misclassifications can lead to breaches of sensitive information. Several strands of theory—ranging from measure-theoretic probability to formal language theory—inform how we integrate computational models with symbolic formalisms. [19]

Consider the set of all tokens in the domain vocabulary denoted by V . A typical language model may attempt to learn a probability distribution $p(t_i | t_{i-1}, t_{i-2}, \dots)$ over these tokens. For de-identification,

the process extends to identifying sequences $\langle t_j, \dots, t_k \rangle$ such that these sequences represent sensitive entities. Let E be the set of entity labels (for example, PERSON, ADDRESS, PHONE) [20]. The de-identification model can be framed as a mapping function $f : V^n \rightarrow E^n$, where each token in the input sequence is assigned an appropriate label or transformation indicating whether it is sensitive. By augmenting this with an attention mechanism, the model associates each token t_i with a context vector that captures dependencies over the entire sequence [21]. In mathematical terms, let $\mathbf{h}_i$ represent the hidden state for token t_i . In a recurrent architecture, we may have

$$\mathbf{h}_i = \sigma(\mathbf{W}\mathbf{h}_{i-1} + \mathbf{U}\mathbf{x}_i + \mathbf{b}),$$

where σ is a suitable nonlinear activation and $\mathbf{x}_i$ is the embedding of token t_i . In an attention-based layer, the hidden representation for t_i may be augmented by a weighted sum of other context vectors, typically [22]

$$\mathbf{c}_i = \sum_{m=1}^n \alpha_{i,m} \mathbf{h}_m,$$

where the attention weights $\alpha_{i,m}$ depend on similarity measures between $\mathbf{h}_i$ and $\mathbf{h}_m$.

However, neural networks on their own may fail to incorporate explicit constraints that are essential in domain-specific de-identification tasks. For instance, certain patterns like date-of-birth or standardized address formats must be replaced systematically [23]. One way to integrate constraints is to treat them as logical assertions about the labels assigned to tokens. Let $P(x)$ denote the predicate that token x represents a personally identifying detail [24]. A simplified logic statement might be $\forall x \in V, [\text{isDigit}(x) \wedge \text{length}(x) = 10] \Rightarrow P(x)$. In a more nuanced scenario, domain-specific rules can be imposed, such as mandatory anonymization for any substring matching patterns of phone numbers. Hence, the model must incorporate additional symbolic variables that represent the presence or absence of these conditions. [25]

In advanced settings, these logic statements can be encoded using matrix representations. Suppose we define a matrix $\mathbf{C}$ such that $\mathbf{C}_{ij} = 1$ if and only if the j -th token in the i -th training example violates a known privacy rule. During training, a penalty term can be introduced to ensure that the system's label prediction aligns with the requirement $\mathbf{C}_{ij} = 0$. This penalty modifies the overall objective function of the neural network, creating a synergy between learned representations and symbolic constraints [26]. Let $\mathbf{y}$ be the predicted labels and $\tilde{\mathbf{y}}$ the ground truth. The standard cross-entropy loss can be denoted as

$$\mathcal{L}_{\text{CE}}(\mathbf{y}, \tilde{\mathbf{y}}) = - \sum_{i=1}^N \sum_{j=1}^n \tilde{y}_{ij} \log(y_{ij}).$$

A constraint-based regularization term $\mathcal{L}_{\text{constraints}}$ might be defined to reflect the degree to which the system satisfies or violates domain rules. The combined objective then becomes [27]

$$\mathcal{L} = \mathcal{L}_{\text{CE}} + \lambda \mathcal{L}_{\text{constraints}},$$

where λ is a hyperparameter. This form ensures that the network not only learns from labeled data but is also guided by logic-based requirements. [28]

From an information-theoretic viewpoint, the objective of de-identification can be interpreted as minimizing the mutual information between the anonymized text and the set of personal identifiers. If $\mathbf{T}$ denotes the random variable representing the text tokens, and $\mathbf{I}$ denotes the random variable representing personal identifiers, we seek a function $g : \mathbf{T} \rightarrow \mathbf{T}'$ that yields an anonymized text $\mathbf{T}'$ with minimal $I(\mathbf{T}'; \mathbf{I})$. In practice, this equates to systematically removing or transforming any aspect of $\mathbf{T}$ that reveals $\mathbf{I}$. At the same time, the transformation must preserve the utility of $\mathbf{T}$ for downstream analytic tasks. This tension between privacy and utility, often studied under the lens of differential privacy or k -anonymity

in database contexts, can be extended to the free-form textual domain by designing models that weigh the significance of each token or phrase.

It is also possible to consider the space of all possible anonymizations of a given text [29]. If Ω represents this space, each element in Ω is a candidate text that meets a certain level of anonymization. We can then define a preference order $\prec$ on Ω based on utility metrics, so that $\mathbf{T}_1 \prec \mathbf{T}_2$ if $\mathbf{T}_2$ yields better performance on a particular downstream task while satisfying the same privacy guarantees. The training process can be framed as identifying an anonymized text that is Pareto-optimal in the trade-off between privacy and utility [30]. Hybrid neural architectures provide a practical mechanism for exploring and approximating this space, especially when large-scale data is involved.

These theoretical considerations lay the groundwork for the practical design described in the subsequent sections [31]. By uniting data-driven embeddings, attention layers, and symbolic logic constraints, we create a model architecture that is both flexible and formally aligned with privacy requirements. This framework sets the stage for detailed descriptions of how the approach is implemented to handle insurance claims narratives, capturing their domain-specific intricacies without sacrificing the rigor demanded by privacy regulations. [32]

3. Proposed Hybrid Neural Architecture

The hybrid neural architecture for de-identification and anonymization combines recurrent and attention-based modules, augmented by external symbolic constraints. The approach is driven by the objective of ensuring that even subtle forms of personal identifiers, such as partial names or location indicators, are consistently detected and transformed [33]. In what follows, we present each of the major components in a continuous narrative without subdividing the section into headings or enumerations.

To begin with, the input to the system is a sequence of tokens that may include words, punctuation, numerical data, and other symbols common in insurance claims. Each token is mapped to an embedding vector via an embedding matrix $\mathbf{E} \in \mathbb{R}^{|V| \times d}$, where $|V|$ is the size of the vocabulary and d is the embedding dimension. This mapping is refined using contextual embeddings derived from a pre-trained model [34]. In the recurrent component, a bidirectional LSTM or GRU processes these embeddings, generating hidden states $\vec{\mathbf{h}}_i$ and $\overleftarrow{\mathbf{h}}_i$, which are then concatenated to yield a contextual representation $\mathbf{h}_i = [\vec{\mathbf{h}}_i; \overleftarrow{\mathbf{h}}_i]$. The concatenation ensures that the representation at each time step accounts for both preceding and succeeding context.

After the recurrent layer, the system applies a self-attention mechanism that assigns varying levels of importance to different tokens when predicting whether a given token is sensitive [35]. Given a query vector $\mathbf{q}_i$, key vectors $\mathbf{k}_j$, and value vectors $\mathbf{v}_j$ for each token j , the attention weights α_{ij} are computed using scaled dot-product attention. In formulaic terms,

$$\alpha_{ij} = \frac{\exp(\mathbf{q}_i \cdot \mathbf{k}_j / \sqrt{d_k})}{\sum_{m=1}^n \exp(\mathbf{q}_i \cdot \mathbf{k}_m / \sqrt{d_k})},$$

where d_k is the dimension of the key vectors [36]. The output of the attention is then $\mathbf{c}_i = \sum_{j=1}^n \alpha_{ij} \mathbf{v}_j$. These outputs are subsequently combined with the recurrent states $\mathbf{h}_i$. One approach is to concatenate the attention output with the recurrent state, giving $\mathbf{o}_i = [\mathbf{h}_i; \mathbf{c}_i]$. The resulting vector $\mathbf{o}_i$ then undergoes further transformation in feed-forward layers to produce label logits for each token, capturing the probability distribution over entity labels, including the sensitive or non-sensitive classification.

A key novelty of this architecture lies in the symbolic constraint integration. During training, the system also receives a set of logical propositions that define domain rules about which tokens or token sequences must be anonymized [37]. These propositions can be stated in first-order logic or via specific pattern-matching rules. For instance, a rule might capture that any token recognized as a partial phone number string must be replaced [38]. In more advanced scenarios, context-based rules might be applied, for instance “any token sequence that follows the phrase ‘policy number’ and matches a digit pattern

must be anonymized.” Such rules are translated into constraint vectors that interact with the network’s output layer. If $\mathbf{z}_i$ denotes the constraint-based label suggestion for token i , then the system incorporates a regularization term that encourages the predicted label $\hat{\mathbf{y}}_i$ to match $\mathbf{z}_i$. The logic statements can be enforced by adding penalty terms into the loss function. In formulaic notation, let us define a function $\delta(i)$ that returns 1 if the token i violates any of the domain rules and 0 otherwise [39]. The penalty term for constraints becomes

$$\mathcal{L}_{\text{constraints}} = \sum_{i=1}^n \delta(i) \Phi(\hat{\mathbf{y}}_i, \mathbf{z}_i),$$

where $\Phi(\hat{\mathbf{y}}_i, \mathbf{z}_i)$ is a measure of divergence between the network’s prediction and the rule-based label for token i . The system thus attempts to reconcile data-driven predictions with symbolic domain knowledge. [40]

Another aspect of the architecture is the use of vector projections to handle a broad range of domain-specific embeddings. Insurance narratives often employ specialized vocabulary, acronyms, and jargon [41]. One can define a mapping from a pretrained embedding space, such as a domain-specific variant of word embeddings, into a lower-dimensional manifold that captures the salient features for de-identification. Suppose $\mathbf{x}_i \in \mathbb{R}^D$ is the original high-dimensional embedding of token i . A projection matrix $\mathbf{P} \in \mathbb{R}^{d \times D}$ maps $\mathbf{x}_i$ to $\mathbf{u}_i \in \mathbb{R}^d$. The transformation [42]

$$\mathbf{u}_i = \mathbf{P}\mathbf{x}_i$$

serves to reduce dimensionality and emphasize features linked to private or sensitive attributes. The parameters of $\mathbf{P}$ can be learned jointly with the rest of the network. This ensures that the final learned representation is aligned with the objective of capturing personal data cues within text.

Upon receiving the final output labels, the system next performs an anonymization step [43]. Each token predicted to belong to a sensitive label category is either replaced with a placeholder (e.g., <PERSON>) or masked in a more sophisticated way. The anonymized output is thus a new token sequence that has minimal references to personal data [44]. To strengthen confidentiality, the system can also incorporate randomization strategies that mitigate the risk of re-identification. One approach is to introduce random synonyms or generic placeholders, thereby reducing the ability of an adversary to reconstruct the original token from partial knowledge [45]. Mathematically, if $\hat{\mathbf{y}}_i$ indicates the predicted label for token t_i , we define a function

$$g(t_i, \hat{\mathbf{y}}_i) = \begin{cases} \text{mask or replace}(t_i) & \text{if } \hat{\mathbf{y}}_i \in \{\text{PERSON, ADDRESS, PHONE, } \dots\}, \\ t_i & \text{otherwise.} \end{cases}$$

The final anonymized text is given by $\langle g(t_1, \hat{\mathbf{y}}_1), g(t_2, \hat{\mathbf{y}}_2), \dots, g(t_n, \hat{\mathbf{y}}_n) \rangle$. This operation is crucial because it not only ensures that sensitive tokens are obfuscated but also maintains lexical and syntactical coherence to the greatest extent possible. The synergy of recurrent, attention-based, and symbolic constraint components helps accurately identify tokens that should be transformed. [46]

Through this integrated design, the model aims to excel in detecting subtle or domain-specific markers of personal identity while adhering to established domain constraints. The final architecture is thus an elegant blend of data-driven deep learning and symbolic knowledge representation, enabling robust, high-quality anonymization of complex textual narratives [47]. The next sections will illuminate the empirical performance of this architecture and its effectiveness in meeting real-world insurance data privacy requirements.

4. Experiments and Evaluation

The system was rigorously evaluated on a corpus of de-identified insurance claims narratives drawn from diverse lines of insurance, such as personal auto, homeowner, and liability claims. Each text

contained varying degrees of personal information, including explicit policyholder names, addresses, occupations, and numeric identifiers [48]. The evaluation was conducted to measure both the predictive accuracy of the model in identifying sensitive tokens and the overall utility of the anonymized text for downstream tasks, like claims analysis or fraud detection.

Each narrative was tokenized and labeled by domain experts, providing ground-truth annotations for sensitive entities [49]. This annotated dataset was partitioned into training, validation, and test sets. The training process involved optimizing the cross-entropy loss with a constraint-based regularization term [50]. Hyperparameters such as embedding dimension, number of attention heads, and recurrent hidden size were tuned based on validation set performance. Additionally, the weight of the constraint penalty term λ was systematically varied to observe its impact on the model's recall and precision for sensitive token detection. [51]

The first step of the evaluation measured entity-level detection metrics. Specifically, the model's predictions were compared to the ground-truth labels at a token level, and standard metrics of precision, recall, and F1-score were computed. The results were compared with two baseline systems: a purely recurrent neural network with no attention layers and a transformer-based model without explicit logic constraints [52]. The purely recurrent system, though able to capture local dependencies, often missed certain contextually implied identifiers. The transformer-based baseline performed better overall but had a tendency to generate false positives in contexts with domain-specific language that closely mimicked personal identifiers [53]. The proposed hybrid approach, by contrast, leveraged both local and global contexts effectively and benefited from symbolic constraints to reduce spurious false positives in domain-specific phrases.

In numerical terms, the hybrid model's performance on detecting sensitive tokens demonstrated a precision of approximately 0.92, a recall of 0.89, and an F1-score of 0.90 [54]. By contrast, the purely recurrent baseline had a precision of 0.88 and recall of 0.83, while the transformer baseline reached a precision of 0.91 but recall of only 0.86. These numeric results highlight that the hybrid model generally outperforms or remains on par with either single-paradigm approach, illustrating the synergy gained from combining recurrent processing, attention, and logic constraints [55]. The gap in recall between the hybrid model and the transformer baseline suggests that symbolic constraints can systematically handle corner cases typically overlooked by purely data-driven approaches.

Another layer of the evaluation focused on utility metrics for downstream analytics [56]. De-identified text was used as input to a claims analytics pipeline that extracted structured data points like claim type, severity level, and potential indicators of fraudulent behavior. Comparative performance was measured on anonymized versus raw text. Notably, the pipeline's accuracy on these tasks decreased when anonymized text from the purely recurrent or transformer baseline was used [57]. However, the decrease was less pronounced when the hybrid model's anonymized output was employed, indicating that the latter preserves relevant contextual details even after removing sensitive information. A plausible explanation is that the symbolic constraints ensure only tokens deemed truly sensitive are altered, reducing collateral damage to tokens containing domain-specific cues necessary for subsequent analyses. [58]

Additionally, a set of privacy-focused metrics was computed to quantify the risk of re-identification. The design of these metrics borrowed concepts from k-anonymity and adversarial re-identification scenarios [59]. Specifically, a script performed a risk analysis by attempting to match anonymized tokens to potential personal identifiers in external databases or cross-referencing them against known partial identities. In these experiments, the hybrid model's anonymized outputs resisted re-identification to a greater degree, due to consistent masking or replacement of sensitive tokens [60]. The purely recurrent baseline occasionally omitted partial or indirect references, permitting a skilled adversary to reconstruct certain personal details with a modest success rate. The transformer-only system was generally effective but occasionally replaced legitimate context with placeholders, thus potentially losing important domain signals. [61]

The computational overhead of the hybrid model was also assessed. Combining a recurrent backbone with attention layers and constraint-based regularization introduces additional parameters and computational steps. However, empirical results indicate that the architecture remains feasible for large-scale

batch processing of insurance claims, especially when parallelization strategies are applied [62]. The self-attention components benefit from efficient implementations on modern GPU hardware, while the recurrent elements handle domain-specific contexts effectively without an excessive memory footprint. Constraint-based operations require no more than a linear pass to map tokens to penalty terms, integrated into the backpropagation process. [63]

Overall, the experiments demonstrate that the hybrid neural architecture provides a balanced solution: it surpasses purely recurrent or purely attention-based methods in precise detection of sensitive information, retains high utility in the downstream tasks that rely on textual coherence, and exhibits strong resilience to re-identification attempts. These results affirm the viability of combining neural representation learning with logic-driven symbolic constraints to achieve robust de-identification in complex textual domains like insurance claims [64]. The subsequent discussion further examines these findings, detailing insights into the approach's strengths and possible avenues for refinement.

5. Discussion and Future Work

Several findings in this research highlight the strengths and limitations of the hybrid neural framework, with implications for broader adoption across various sectors dealing with sensitive textual data [65]. First, the synergy between recurrent and attention-based mechanisms has been validated by empirical results. On one hand, recurrent layers capture the local dependencies typically found in sensitive information, such as phone numbers or addresses that adhere to predictable patterns. On the other hand, attention layers identify broader contextual cues indicating the presence of implied personal identifiers [66]. This multifaceted focus results in a more robust detection of sensitive information.

The integration of symbolic constraints plays an equally critical role [67]. In many de-identification tasks, purely data-driven neural models are prone to overfitting to patterns seen in the training data. When new or specialized forms of sensitive information appear, such models may fail to recognize them adequately [68]. By incorporating logic statements that define domain-specific privacy rules, the system can generalize more reliably to previously unseen data. It is also notable that these constraints do not hinder the system's capacity to learn from the raw text; rather, they provide guidance that the model can incorporate to refine its predictions, especially in edge cases where domain knowledge is crucial. [69]

Despite these strengths, several challenges remain. The process of creating and maintaining logic statements or domain-specific constraints can be labor-intensive, as it requires collaboration between domain experts and system developers [70]. Furthermore, symbolic rules may become outdated as the domain evolves. For instance, new policy types or new formats of identifying information might emerge, necessitating continual updates. A potential avenue for future work involves the automated extraction of symbolic constraints from annotated data [71]. Advances in inductive logic programming and semantic pattern discovery could potentially alleviate the manual effort in specifying domain rules.

Another issue revolves around the balance between anonymity and utility [72]. While the results show that the hybrid model preserves considerable downstream utility, certain tasks may require near-verbatim text for accurate predictions. If anonymization is too aggressive, it may degrade the text's semantic content, limiting the system's applicability [73]. Future research could explore adaptive anonymization techniques that tailor transformations based on the needs of subsequent tasks. One might define a function that weighs the significance of each token or phrase for a downstream task, enabling the model to preserve tokens that are crucial for predictive performance while still minimizing the risk of re-identification. [74]

Scalability is also an ongoing concern. Though the hybrid architecture can be parallelized to handle large volumes of text, many organizations process millions of claims or customer interactions daily [75]. The computational cost of advanced attention mechanisms, especially when combined with symbolic constraints, can become nontrivial. Approaches to alleviate this may include sparse attention mechanisms, knowledge distillation, or hierarchical modeling that first filters out obviously non-sensitive segments before applying deeper analysis to smaller subsections. Additionally, the constraints might be

enforced in a two-stage pipeline, where an initial rule-based system tags obvious entities, reducing the burden on the neural components. [76]

The complexity of cross-lingual or multilingual de-identification remains another frontier. Many insurers now operate in multiple countries, requiring them to process claims in diverse languages [77]. The proposed approach could be extended by incorporating multilingual embeddings or by training language-specific sub-modules. However, domain-specific logic constraints may differ across languages due to differences in address formats, naming conventions, or other cultural factors [78]. This raises intriguing questions about how best to fuse cross-lingual embedding spaces with localized symbolic rules, especially if a single global system must handle multiple linguistic contexts.

Privacy-preserving transformations also intersect with legal and ethical considerations [79]. In some jurisdictions, partial redactions might be permissible if the redacted portion still obscures the identity sufficiently. In others, more stringent anonymization standards require complete removal or random replacement of certain token categories. Integrating knowledge of local regulations into the symbolic constraints can help align the system with legal requirements [80]. Indeed, the interplay between computational design and legal frameworks is increasingly shaping how AI systems are built, indicating a need for interdisciplinary approaches.

Finally, interpretability and model auditing remain vital [81]. Although the logic constraints offer some transparency, the neural elements of the system can be opaque. If an entity is mistakenly left unanonymized, it may be challenging for auditors to pinpoint exactly which part of the model is responsible [82]. Future developments may involve advanced interpretability methods, such as layer-wise relevance propagation or attention weight analysis. By examining how much influence constraints have on the model's predictions, system maintainers could better detect failures in real time and reduce potential privacy violations. [83]

In sum, the hybrid neural framework proposed here provides a robust platform for de-identification in insurance claims narratives, combining advanced representation learning with formalized domain knowledge. The success observed in controlled experiments and real-world applications strongly suggests that this model paradigm can be generalized to other high-stakes textual contexts [84]. Continued improvements in constraint learning, multilingual support, and interpretability will likely further enhance its value, enabling safer, more secure handling of sensitive data across industries.

6. Conclusion

This paper has examined the use of hybrid neural architectures to address the dual objectives of de-identification and anonymization in insurance claims narratives. The investigation began with the theoretical considerations that guide the modeling of sensitive token detection, utilizing principles from probability and logic [85]. The proposed system combines the contextual depth of recurrent networks, the global insight of attention mechanisms, and the domain specificity enforced by symbolic constraints. Empirical results across a large and heterogeneous dataset of insurance claims confirm that such an integrated approach surpasses purely recurrent or purely attention-based systems in both effectiveness and robustness. [86]

Beyond predictive metrics, the model's capacity to preserve semantic coherence in de-identified text underscores its practical applicability to real-world claims processing, fraud detection, and analytics. The synergy between advanced representation learning and structured domain knowledge contributes to a comprehensive anonymization pipeline that consistently flags and obscures personal information [87]. By explicitly incorporating logic-based constraints, the system can adapt to complex patterns of sensitive information and ensure adherence to evolving regulatory mandates.

The study also outlined several directions for future research, including automated rule extraction, adaptive anonymization tailored to downstream tasks, and the integration of privacy-preserving mechanisms in multi-lingual contexts [88]. In addition, it highlighted challenges related to balancing privacy and utility, maintaining scalability, and ensuring interpretability. As privacy concerns continue

to intensify, the criticality of robust de-identification methodologies will only grow. Hybrid neural architectures, enriched by domain-specific symbolic rules, represent a promising pathway for building secure and context-aware anonymization systems, aligning with both the business imperatives of data-driven decision-making and the ethical imperative to safeguard personal information. [89]

References

- [1] S. Paheding, A. Saleem, M. F. H. Siddiqui, N. Rawashdeh, A. Essa, and A. A. Reyes, "Advancing horizons in remote sensing: a comprehensive survey of deep learning models and applications in image classification and beyond," *Neural Computing and Applications*, vol. 36, pp. 16727–16767, 8 2024.
- [2] R. Ramon-Gonen, A. Dori, and S. Shelly, "Towards a practical use of text mining approaches in electrodiagnostic data.," *Scientific reports*, vol. 13, pp. 19483–, 11 2023.
- [3] J. R. Machireddy, "Automation in healthcare claims processing: Enhancing efficiency and accuracy," *International Journal of Science and Research Archive*, vol. 09, no. 01, pp. 825–834, 2023.
- [4] Y. Zheng, W. Gan, Z. Chen, Z. Qi, Q. Liang, and P. S. Yu, "Large language models for medicine: a survey," *International Journal of Machine Learning and Cybernetics*, vol. 16, pp. 1015–1040, 8 2024.
- [5] Y. Lin and N. Nixon, "Transitioning to online instructions and covid-19 response: A view from mining emergent college students discourse in online discussion forum," *International Journal of Artificial Intelligence in Education*, vol. 34, pp. 706–731, 6 2024.
- [6] O. Miranda, S. M. Kiehl, X. Qi, M. D. Brannock, T. Kosten, N. D. Ryan, L. Kirisci, Y. Wang, and L. Wang, "Enhancing post-traumatic stress disorder patient assessment: leveraging natural language processing for research of domain criteria identification using electronic medical records.," *BMC medical informatics and decision making*, vol. 24, pp. 154–, 6 2024.
- [7] M. Guo, Y. Ma, E. Eworuke, M. Khashei, J. Song, Y. Zhao, and F. Jin, "Identifying covid-19 cases and extracting patient reported symptoms from reddit using natural language processing.," *Scientific reports*, vol. 13, pp. 13721–, 8 2023.
- [8] S. Ahmed, I. E. Nielsen, A. Tripathi, S. Siddiqui, R. P. Ramachandran, and G. Rasool, "Transformers in time-series analysis: A tutorial," *Circuits, Systems, and Signal Processing*, vol. 42, pp. 7433–7466, 7 2023.
- [9] J. I. Park, J. W. Park, K. Zhang, and D. Kim, "Advancing equity in breast cancer care: natural language processing for analysing treatment outcomes in under-represented populations.," *BMJ health & care informatics*, vol. 31, pp. e100966–e100966, 7 2024.
- [10] C. Carini and A. A. Seyhan, "Tribulations and future opportunities for artificial intelligence in precision medicine.," *Journal of translational medicine*, vol. 22, pp. 411–, 4 2024.
- [11] J. M. Harrison, A. Yala, P. Mikhael, J. Roldan, D. Ciprani, T. Michelakos, L. Bolm, M. Qadan, C. Ferrone, C. F.-D. Castillo, K. D. Lillemoe, E. Santus, and K. Hughes, "Successful development of a natural language processing algorithm for pancreatic neoplasms and associated histologic features.," *Pancreas*, vol. 52, pp. e219–e223, 4 2023.
- [12] M. Chaudhary, K. Kosyluk, S. Thomas, and T. Neal, "On the use of aspect-based sentiment analysis of twitter data to explore the experiences of african americans during covid-19.," *Scientific reports*, vol. 13, pp. 10694–, 7 2023.
- [13] G. Arnold, M. Klasic, C. Wu, M. Schomburg, and A. York, "Finding, distinguishing, and understanding overlooked policy entrepreneurs," *Policy Sciences*, vol. 56, pp. 657–687, 11 2023.
- [14] J. K. Scroggins, I. I. Hulchaf, S. Harkins, D. Scharp, H. Moen, A. Davoudi, K. Cato, M. Tadiello, M. Topaz, and V. Barcelona, "Identifying stigmatizing and positive/preferred language in obstetric clinical notes using natural language processing.," *Journal of the American Medical Informatics Association : JAMIA*, vol. 32, pp. 308–317, 11 2024.
- [15] D. Clunie, A. Taylor, T. Bisson, D. Gutman, Y. Xiao, C. G. Schwarz, D. Greve, J. Gichoya, G. Shih, A. Kline, B. Kopchick, and K. Farahani, "Summary of the national cancer institute 2023 virtual workshop on medical image de-identification-part 2: Pathology whole slide image de-identification, de-facing, the role of ai in image de-identification, and the nci midi datasets and pipeline.," *Journal of imaging informatics in medicine*, vol. 38, pp. 16–30, 7 2024.
- [16] K. D. Forbus, C. Riesbeck, L. Birnbaum, K. Livingston, A. B. Sharma, and L. C. Ureel, "A prototype system that learns by reading simplified texts.," in *AAAI Spring Symposium: Machine Reading*, pp. 49–54, 2007.

- [17] R. Giorgino, M. Alessandri-Bonetti, M. D. Re, F. Verdoni, G. M. Peretti, and L. Mangiavini, "Google bard and chatgpt in orthopedics: Which is the better doctor in sports medicine and pediatric orthopedics? the role of ai in patient education.," *Diagnostics (Basel, Switzerland)*, vol. 14, pp. 1253–1253, 6 2024.
- [18] A. Li, S. Ma, R. Bandyo, M. Ranjan, E. Chiang, J. Tiong, J. Y. Jiang, J. Ryu, O. Jafari, C. I. Amos, N. R. Fillmore, J. La, K. A. Martin, C. J. Shah, and A. Jotwani, "Implementation of automated thrombosis and bleeding risk assessment in cancer patients undergoing systemic therapy," *Blood*, vol. 144, pp. 3641–3641, 11 2024.
- [19] A. Rafiei, R. Moore, S. Jahromi, F. Hajati, and R. Kamaleswaran, "Meta-learning in healthcare: A survey," *SN Computer Science*, vol. 5, 8 2024.
- [20] K. Rangan and Y. Yin, "A fine-tuning enhanced rag system with quantized influence measure as ai judge.," *Scientific reports*, vol. 14, pp. 27446–, 11 2024.
- [21] A. K. Saxena, R. R. Dixit, and A. Aman-Ullah, "An lstm neural network approach to resource allocation in hospital management systems," *International Journal of Applied Health Care Analytics*, vol. 7, no. 2, pp. 1–12, 2022.
- [22] Y. Zhang, J. R. Golbus, E. Wittrup, K. D. Aaronson, and K. Najarian, "Enhancing heart failure treatment decisions: interpretable machine learning models for advanced therapy eligibility prediction using ehr data.," *BMC medical informatics and decision making*, vol. 24, pp. 53–, 2 2024.
- [23] S. Nie, S. Zhang, Y. Zhao, X. Li, H. Xu, Y. Wang, X. Wang, and M. Zhu, "Machine learning applications in acute coronary syndrome: Diagnosis, outcomes and management.," *Advances in therapy*, vol. 42, pp. 636–665, 12 2024.
- [24] Y. Li, W. Tao, Z. Li, Z. Sun, F. Li, S. Fenton, H. Xu, and C. Tao, "Artificial intelligence-powered pharmacovigilance: A review of machine and deep learning in clinical text-based adverse drug event detection for benchmark datasets.," *Journal of biomedical informatics*, vol. 152, pp. 104621–104621, 3 2024.
- [25] A. Sharma, M. Witbrock, and K. Goolsbey, "Controlling search in very large commonsense knowledge bases: a machine learning approach," *arXiv preprint arXiv:1603.04402*, 2016.
- [26] B. Roth, R. Kampalath, K. Nakashima, S. Shieh, T.-L. Bui, and R. Houshyar, "Revenue and cost analysis of a system utilizing natural language processing and a nurse coordinator for radiology follow-up recommendations.," *Current problems in diagnostic radiology*, vol. 52, pp. 367–371, 5 2023.
- [27] M. Delsoz, H. Raja, Y. Madadi, A. A. Tang, B. M. Wirosko, M. Y. Kahook, and S. Yousefi, "The use of chatgpt to assist in diagnosing glaucoma based on clinical case reports.," *Ophthalmology and therapy*, vol. 12, pp. 3121–3132, 9 2023.
- [28] A. N. Berman, C. Ginder, Z. A. Sporn, V. Tanguturi, M. K. Hidrue, L. B. Shirkey, Y. Zhao, R. Blankstein, A. Turchin, and J. H. Wasfy, "Natural language processing for the ascertainment and phenotyping of left ventricular hypertrophy and hypertrophic cardiomyopathy on echocardiogram reports.," *The American journal of cardiology*, vol. 206, pp. 247–253, 9 2023.
- [29] W. T. Kerr and K. N. McFarlane, "Machine learning and artificial intelligence applications to epilepsy: a review for the practicing epileptologist.," *Current neurology and neuroscience reports*, vol. 23, pp. 869–879, 12 2023.
- [30] E. Warner, J. Lee, W. Hsu, T. Syeda-Mahmood, C. E. Kahn, O. Gevaert, and A. Rao, "Multimodal machine learning in image-based and clinical biomedicine: Survey and prospects.," *International journal of computer vision*, vol. 132, pp. 3753–3769, 4 2024.
- [31] N. A. Stadnick, G. A. Aarons, H. N. Edwards, A. W. Bryl, C. L. Kuelbs, J. L. Helm, and L. Brookman-Frazee, "Cluster randomized trial of a team communication training implementation strategy for depression screening in a pediatric healthcare system: a study protocol.," *Implementation science communications*, vol. 5, pp. 117–, 10 2024.
- [32] L. C. Backer, "The soulful machine, the virtual person, and the "human" condition: an encounter with jan m. broekman, knowledge in change: The semiotics of cognition and conversion (cham, switzerland: Springer nature, 2023).," *International Journal for the Semiotics of Law - Revue internationale de Sémiotique juridique*, vol. 37, pp. 969–1083, 2 2024.
- [33] A. Abhishek and A. Basu, "A framework for disambiguation in ambiguous iconic environments," in *AI 2004: Advances in Artificial Intelligence: 17th Australian Joint Conference on Artificial Intelligence, Cairns, Australia, December 4-6, 2004. Proceedings 17*, pp. 1135–1140, Springer, 2005.
- [34] L. Jacaruso, "Insights into the nutritional prevention of macular degeneration based on a comparative topic modeling approach.," *PeerJ. Computer science*, vol. 10, pp. e1940–e1940, 3 2024.

- [35] A. Sharma and K. Forbus, "Modeling the evolution of knowledge in learning systems," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 26, pp. 669–675, 2012.
- [36] Y. Shen, Y. Xu, J. Ma, W. Rui, C. Zhao, L. Heacock, and C. Huang, "Multi-modal large language models in radiology: principles, applications, and potential.," *Abdominal radiology (New York)*, 12 2024.
- [37] M. Abouelyazid and C. Xiang, "Machine learning-assisted approach for fetal health status prediction using cardiotocogram data," *International Journal of Applied Health Care Analytics*, vol. 6, no. 4, pp. 1–22, 2021.
- [38] J. Lu, Y. Yan, K. Huang, M. Yin, and F. Zhang, "Do we learn from each other: Understanding the human-ai co-learning process embedded in human-ai collaboration," *Group Decision and Negotiation*, 12 2024.
- [39] L. He, M. Moldenhauer, K. Zheng, and H. Ma, "Analyzing free-text clinical narratives for veterans with lymphoid malignancies using natural language processing (nlp).," *Journal of Clinical Oncology*, vol. 41, pp. e13576–e13576, 6 2023.
- [40] M. Muniswamaiah, T. Agerwala, and C. C. Tappert, "Automatic visual recommendation for data science and analytics," in *Advances in Information and Communication: Proceedings of the 2020 Future of Information and Communication Conference (FICC), Volume 2*, pp. 125–132, Springer, 2020.
- [41] P. K. Vutukuri, "Claim extraction process automation," *International Journal of Science and Research (IJSR)*, vol. 13, pp. 1338–1342, 6 2024.
- [42] A. F. Syafiandini, J. Yoon, S. Lee, C. Song, E. Yan, and M. Song, "Examining between-sectors knowledge transfer in the pharmacology field," *Scientometrics*, vol. 129, pp. 3115–3147, 5 2024.
- [43] A. Dalal, S. Ranjan, Y. Bopaiah, D. Chembachere, N. Steiger, C. Burns, and V. Daswani, "Text summarization for pharmaceutical sciences using hierarchical clustering with a weighted evaluation methodology.," *Scientific reports*, vol. 14, pp. 20149–, 8 2024.
- [44] S. Parsa, S. Somani, R. Dudum, S. S. Jain, and F. Rodriguez, "Artificial intelligence in cardiovascular disease prevention: Is it ready for prime time?," *Current atherosclerosis reports*, vol. 26, pp. 263–272, 5 2024.
- [45] J. L. Sorensen, M. M. West, A. M. Racila, O. A. Amao, B. J. Matt, S. Bentler, A. R. Kahl, M. E. Charlton, A. T. Seaman, and S. H. Nash, "Challenges in collecting information on sexual orientation and gender identity for cancer patients: perspectives of hospital and central cancer registry abstractors.," *Cancer causes & control : CCC*, vol. 35, pp. 1433–1445, 7 2024.
- [46] E. Okeh, "Transforming healthcare: A comprehensive approach to mitigating bias and fostering empathy through ai-driven augmented reality," *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, pp. 23753–23754, 3 2024.
- [47] M. Muniswamaiah, T. Agerwala, and C. Tappert, "Big data in cloud computing review and opportunities," *arXiv preprint arXiv:1912.10821*, 2019.
- [48] H. Li, R. C. Gerkin, A. Bakke, R. Norel, G. Cecchi, C. Laudamiel, M. Y. Niv, K. Ohla, J. E. Hayes, V. Parma, and P. Meyer, "Text-based predictions of covid-19 diagnosis from self-reported chemosensory descriptions.," *Communications medicine*, vol. 3, pp. 104–, 7 2023.
- [49] Y. Artsi, V. Sorin, E. Konen, B. S. Glicksberg, G. Nadkarni, and E. Klang, "Large language models for generating medical examinations: systematic review.," *BMC medical education*, vol. 24, pp. 354–, 3 2024.
- [50] S. E. Davis, L. Zabotka, R. J. Desai, S. V. Wang, J. C. Maro, K. Coughlin, J. J. Hernández-Muñoz, D. Stojanovic, N. H. Shah, and J. C. Smith, "Use of electronic health record data for drug safety signal identification: A scoping review.," *Drug safety*, vol. 46, pp. 725–742, 6 2023.
- [51] R. Yang, Q. Zeng, K. You, Y. Qiao, L. Huang, C.-C. Hsieh, B. Rosand, J. Goldwasser, A. Dave, T. Keenan, Y. Ke, C. Hong, N. Liu, E. Chew, D. Radev, Z. Lu, H. Xu, Q. Chen, and I. Li, "Ascle-a python natural language processing toolkit for medical text generation: Development and evaluation study.," *Journal of medical Internet research*, vol. 26, pp. e60601–e60601, 10 2024.
- [52] T. A. Meier, M. S. Refahi, G. Hearne, D. S. Restifo, R. Munoz-Acuna, G. L. Rosen, and S. Woloszynek, "The role and applications of artificial intelligence in the treatment of chronic pain.," *Current pain and headache reports*, vol. 28, pp. 769–784, 6 2024.
- [53] K. S. Allen, D. R. Hood, J. Cummins, S. Kasturi, E. A. Mendonca, and J. R. Vest, "Natural language processing-driven state machines to extract social factors from unstructured clinical documentation.," *JAMIA open*, vol. 6, pp. ooad024–, 4 2023.

- [54] M. Muniswamaiah, T. Agerwala, and C. C. Tappert, "Approximate query processing for big data in heterogeneous databases," in *2020 IEEE international conference on big data (big data)*, pp. 5765–5767, IEEE, 2020.
- [55] M. V. Mazzolenis, G. N. Mourra, S. Moreau, M. E. Mazzolenis, I. H. Cerda, J. Vega, J. S. Khan, and A. Thérond, "The role of virtual reality and artificial intelligence in cognitive pain therapy: A narrative review.," *Current pain and headache reports*, vol. 28, pp. 881–892, 6 2024.
- [56] A. A. Mäkitie, R. O. Alabi, S. P. Ng, R. P. Takes, K. T. Robbins, O. Ronen, A. R. Shaha, P. J. Bradley, N. F. Saba, S. Nuyts, A. Triantafyllou, C. Piazza, A. Rinaldo, and A. Ferlito, "Artificial intelligence in head and neck cancer: A systematic review of systematic reviews.," *Advances in therapy*, vol. 40, pp. 3360–3380, 6 2023.
- [57] A. Mudgal, U. Kush, A. Kumar, and A. Jafari, "Multimodal fusion: advancing medical visual question-answering," *Neural Computing and Applications*, vol. 36, pp. 20949–20962, 8 2024.
- [58] H. Papneja and N. Yadav, "Self-disclosure to conversational ai: a literature review, emergent framework, and directions for future research," *Personal and Ubiquitous Computing*, 8 2024.
- [59] J. R. Machireddy, "Machine learning and automation in healthcare claims processing," *Journal of Artificial Intelligence General science (JAIGS) ISSN: 3006-4023*, vol. 6, no. 1, pp. 686–701, 2024.
- [60] L. Tang, Z. Sun, B. Idnay, J. G. Nestor, A. Soroush, P. A. Elias, Z. Xu, Y. Ding, G. Durrett, J. F. Rousseau, C. Weng, and Y. Peng, "Evaluating large language models on medical evidence summarization.," *NPJ digital medicine*, vol. 6, pp. 158–, 8 2023.
- [61] C. Y. Chen, A. Christoffels, R. Dube, K. Enos, J. E. Gilbert, S. Koyejo, J. Leigh, C. Liquido, A. McKee, K. Noe, T.-Q. Peng, and K. Taiuru, "Increasing the presence of bipoc researchers in computational science.," *Nature computational science*, vol. 4, pp. 646–653, 9 2024.
- [62] J. Chen and S. Tajdini, "A moderated model of artificial intelligence adoption in firms and its effects on their performance," *Information Technology and Management*, 4 2024.
- [63] S. Niu, J. Ma, Q. Yin, Z. Wang, L. Bai, and X. Yang, "Modelling patient longitudinal data for clinical decision support: A case study on emerging ai healthcare technologies," *Information Systems Frontiers*, 7 2024.
- [64] G. M. Silverman, G. Rajamani, N. E. Ingraham, J. K. Glover, H. S. Sahoo, M. Usher, R. Zhang, F. Ikramuddin, T. E. Melnik, G. B. Melton, and C. J. Tignanelli, "A symptom-based natural language processing surveillance pipeline for post-covid-19 patients.," *Studies in health technology and informatics*, vol. 310, pp. 860–, 1 2024.
- [65] N. Brugnone, N. Benkler, P. Revay, R. Myhre, S. Friedman, S. Schmer-Galunder, S. Gray, and J. Gentile, "Is from ought? a comparison of unsupervised methods for structuring values-based wisdom-of-crowds estimates," *Journal of Computational Social Science*, vol. 7, pp. 1327–1377, 4 2024.
- [66] M. M. Dagli, Y. Ghenbot, H. S. Ahmad, D. Chauhan, R. Turlip, P. Wang, W. C. Welch, A. K. Ozturk, and J. W. Yoon, "Development and validation of a novel ai framework using nlp with llm integration for relevant clinical data extraction through automated chart review.," *Scientific reports*, vol. 14, pp. 26783–, 11 2024.
- [67] B. Zaidat, Y. S. Lahoti, A. Yu, K. S. Mohamed, S. K. Cho, and J. S. Kim, "Artificially intelligent billing in spine surgery: An analysis of a large language model.," *Global spine journal*, vol. 15, pp. 1113–1120, 12 2023.
- [68] null Muhammad Bilal Ahmad Jamil and null Duryab Shahzadi, "A systematic review a conversational interface agent for the export business acceleration," *Lahore Garrison University Research Journal of Computer Science and Information Technology*, vol. 7, pp. 37–49, 8 2023.
- [69] L. E. Young, Y. Nan, E. Jang, and R. Stevens, "Digital epidemiological approaches in hiv research: a scoping methodological review.," *Current HIV/AIDS reports*, vol. 20, pp. 470–480, 11 2023.
- [70] A. K. Saxena, "Evaluating the regulatory and policy recommendations for promoting information diversity in the digital age," *International Journal of Responsible Artificial Intelligence*, vol. 11, no. 8, pp. 33–42, 2021.
- [71] Z. Xiao, J. Zhu, Y. Wang, P. Zhou, W. H. Lam, M. A. Porter, and Y. Sun, "Detecting political biases of named entities and hashtags on twitter," *EPJ Data Science*, vol. 12, 6 2023.
- [72] K. Zhang, R. Zhou, E. Adhikarla, Z. Yan, Y. Liu, J. Yu, Z. Liu, X. Chen, B. D. Davison, H. Ren, J. Huang, C. Chen, Y. Zhou, S. Fu, W. Liu, T. Liu, X. Li, Y. Chen, L. He, J. Zou, Q. Li, H. Liu, and L. Sun, "A generalist vision-language foundation model for diverse biomedical tasks.," *Nature medicine*, vol. 30, pp. 3129–3141, 8 2024.

- [73] J. A. R. Ivorra, E. Trallero-Araguas, M. L. Lasanta, L. Cebrián, L. Lojo, B. López-Muñiz, J. Fernández-Melon, B. Núñez, L. Silva-Fernández, R. V. Cabello, P. Ahijado, I. D. la Morena Barrio, N. C. Torrijo, B. Safont, E. Ornilla, J. Restrepo, A. Campo, J. L. Andreu, E. Díez, A. L. Robles, E. Bollo, D. Benavent, D. Vilanova, S. L. Valdés, and R. Castellanos-Moreira, "Prevalence and clinical characteristics of patients with rheumatoid arthritis with interstitial lung disease using unstructured healthcare data and machine learning," *RMD open*, vol. 10, pp. e003353–e003353, 1 2024.
- [74] N. Venishetty and O. A. Raheem, "Commentary on: Frequently asked questions on erectile dysfunction: evaluating artificial intelligence answers with expert mentorship.," *International journal of impotence research*, 5 2024.
- [75] Y. Liu, Y. Luo, and A. M. Naidech, "Big data in stroke: How to use big data to make the next management decision.," *Neurotherapeutics : the journal of the American Society for Experimental NeuroTherapeutics*, vol. 20, pp. 744–757, 3 2023.
- [76] C. J. Lobel, D. A. Laney, J. Yang, D. Jacob, A. Rickheim, C. Z. Ogg, D. Clynes, and J. Dronen, "Fdrisk: development of a validated risk assessment tool for fabry disease utilizing electronic health record data.," *Journal of rare diseases (Berlin, Germany)*, vol. 3, pp. 2–, 1 2024.
- [77] M. Muniswamaiah, T. Agerwala, and C. C. Tappert, "Federated query processing for big data in data science.," in *2019 IEEE International Conference on Big Data (Big Data)*, pp. 6145–6147, IEEE, 2019.
- [78] S. N. Goldman, A. T. Hui, S. Choi, E. K. Mbamalu, P. Tirabady, A. S. Eleswarapu, J. A. Gomez, L. M. Alvandi, and E. D. Fornari, "Applications of artificial intelligence for adolescent idiopathic scoliosis: mapping the evidence.," *Spine deformity*, vol. 12, pp. 1545–1570, 8 2024.
- [79] Z. He, A. Yan, A. Gentili, J. McAuley, and C.-N. Hsu, "'nothing abnormal': Disambiguating medical reports via contrastive knowledge infusion," *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37, pp. 14232–14240, 6 2023.
- [80] H. Friedman, M. Li, K. L. Harvey, I. Griesemer, D. Mohr, A. M. Linsky, and D. Gurewich, "Identifying veterans with a higher risk of social needs using cluster analysis.," *Journal of general internal medicine*, vol. 40, pp. 385–392, 10 2024.
- [81] Y. M. Hwang, S. N. Piekos, A. G. Paquette, Q. Wei, N. D. Price, L. Hood, and J. J. Hadlock, "Accelerating adverse pregnancy outcomes research amidst rising medication use: parallel retrospective cohort analyses for signal prioritization.," *BMC medicine*, vol. 22, pp. 495–, 10 2024.
- [82] M. Alsaiddi, M. T. Jan, A. Altaher, H. Zhuang, and X. Zhu, "Tackling the class imbalanced dermoscopic image classification using data augmentation and gan," *Multimedia Tools and Applications*, vol. 83, pp. 49121–49147, 10 2023.
- [83] D. David, O. Zloto, G. Katz, R. Huna-Baron, V. Vishnevskia-Dai, S. Armarnik, N. A. Zauberman, E. M. Barnir, R. Singer, A. Hostovsky, and E. Klang, "The use of artificial intelligence based chat bots in ophthalmology triage.," *Eye (London, England)*, vol. 39, pp. 785–789, 11 2024.
- [84] Z. B. Millman, D. J. Holt, M. S. Keshavan, L. Seidman, A. Breier, M. E. Shenton, D. Öngür, S. V. Abram, J. P. Hua, S. Nicholas, B. J. Roach, S. K. Keedy, J. A. Sweeney, D. H. Mathalon, J. M. Ford, S. Erhardt, Y. Gravenfors, C. Savage, G. Prettyman, P. Didier, J. W. Kable, T. D. Satterthwaite, C. Davatzikos, R. T. Shinohara, M. E. Calkins, M. Elliott, R. C. Gur, R. Gur, and D. H. Wolf, "Acnp 62nd annual meeting: Poster abstracts p501 – p753," *Neuropsychopharmacology*, vol. 48, pp. 355–495, 12 2023.
- [85] H. Kotzab, I. Özge Yumurtacı Hüseyinoğlu, and J. Fischer, "Foundations and areas of covid-19 pandemic logistics and supply chain management research: the early and mid-pandemic view," *Review of Managerial Science*, 11 2024.
- [86] R. Golan, A. Swayze, Z. M. Connelly, J. Loloi, K. Campbell, A. C. Lentz, K. Watts, A. Small, M. Babar, R. D. Patel, R. Ramasamy, and S. Loeb, "Can artificial intelligence evaluate the quality of youtube videos on erectile dysfunction?," *BJU international*, vol. 135, pp. 431–433, 11 2024.
- [87] A. Sharma and K. D. Forbus, "Graph-based reasoning and reinforcement learning for improving q/a performance in large knowledge-based systems," in *2010 AAAI Fall Symposium Series*, 2010.
- [88] Y. Wang, E. Willis, V. K. Yeruva, D. Ho, and Y. Lee, "A case study of using natural language processing to extract consumer insights from tweets in american cities for public health crises.," *BMC public health*, vol. 23, pp. 935–, 5 2023.
- [89] C. Reich and B. Meder, "The heart and artificial intelligence-how can we improve medicine without causing harm.," *Current heart failure reports*, vol. 20, pp. 271–279, 6 2023.