

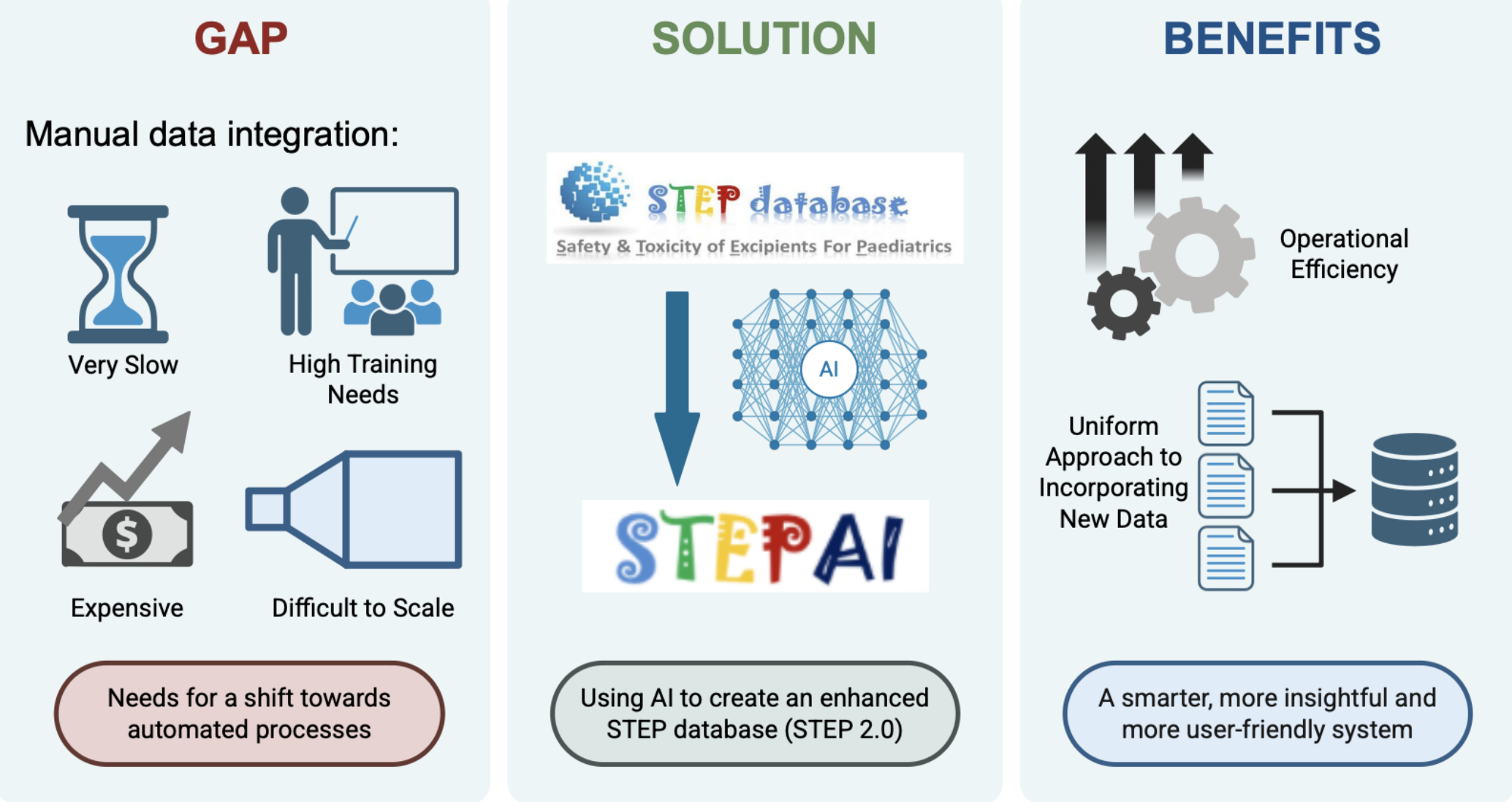
Automated Extraction of Pharmaceutical Knowledge Using Artificial Intelligence for the STEP Database

Sifan Hu¹; Sagar Uprety¹; Sridhar Radhakrishnan²; Chandra Sekharan³; Peter Silvera⁴; Ranjana Srivastava⁴; James Cummins⁵; Kristen Porter⁵; Santosh Reddy Nella⁶; Daniel Schaufelberger⁷; Bruno Hancock⁸; Savani Deshpande¹; Smita Salunke¹

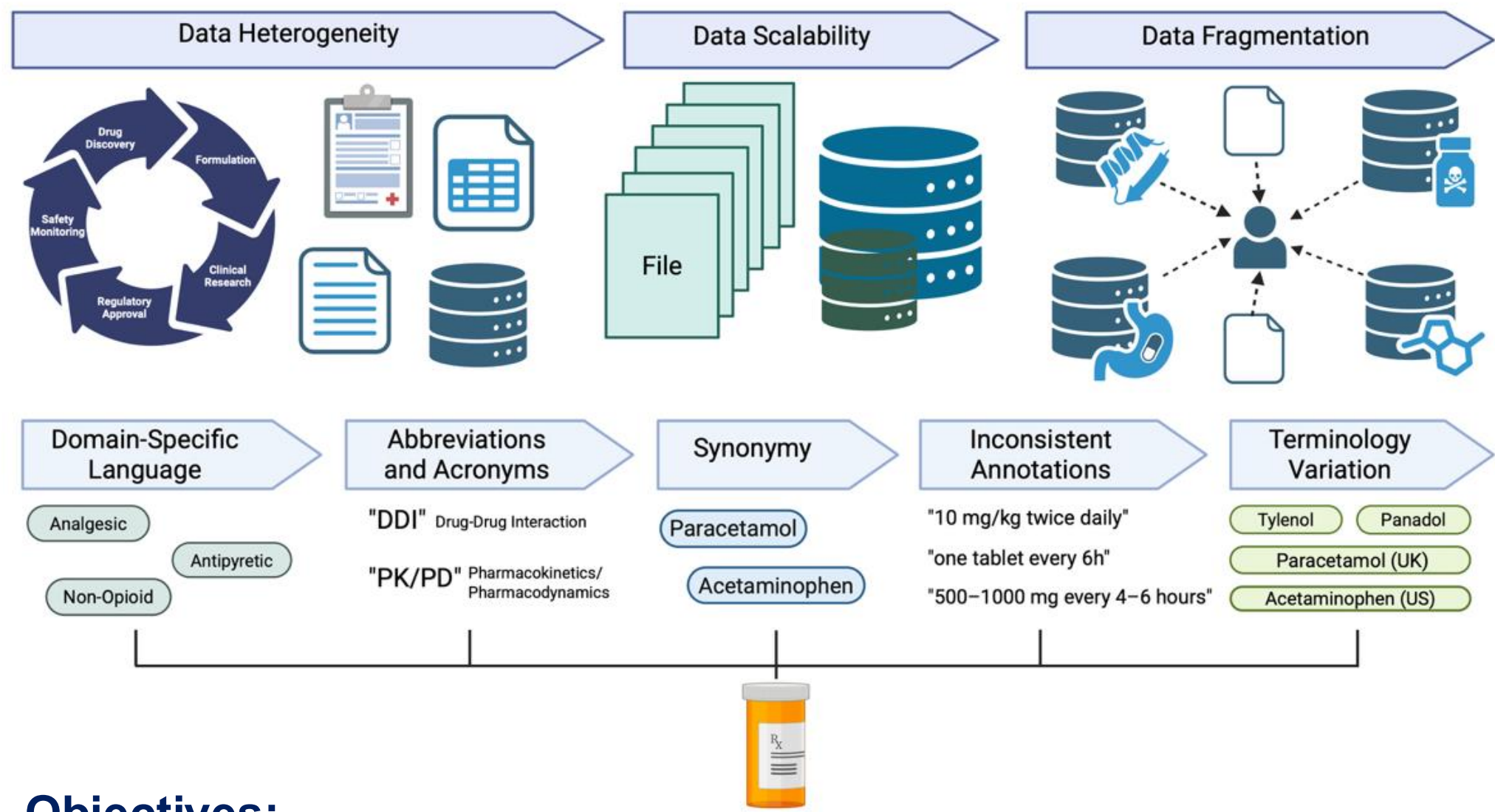
¹ UCL School of Pharmacy, London, UK; ² School of Computer Science, University of Oklahoma, USA; ³ Department of Computer Science, Texas A&M University Corpus Christi, USA; ⁴ Advanced Bioscience Laboratories, Inc, USA; ⁵ National Institute of Allergy and Infectious Diseases (NIAID), USA; ⁶ Vincatis Technologies Private Limited, India; ⁷ Schaufelberger Consulting LLC, USA; ⁸ Aleurites LLC, USA



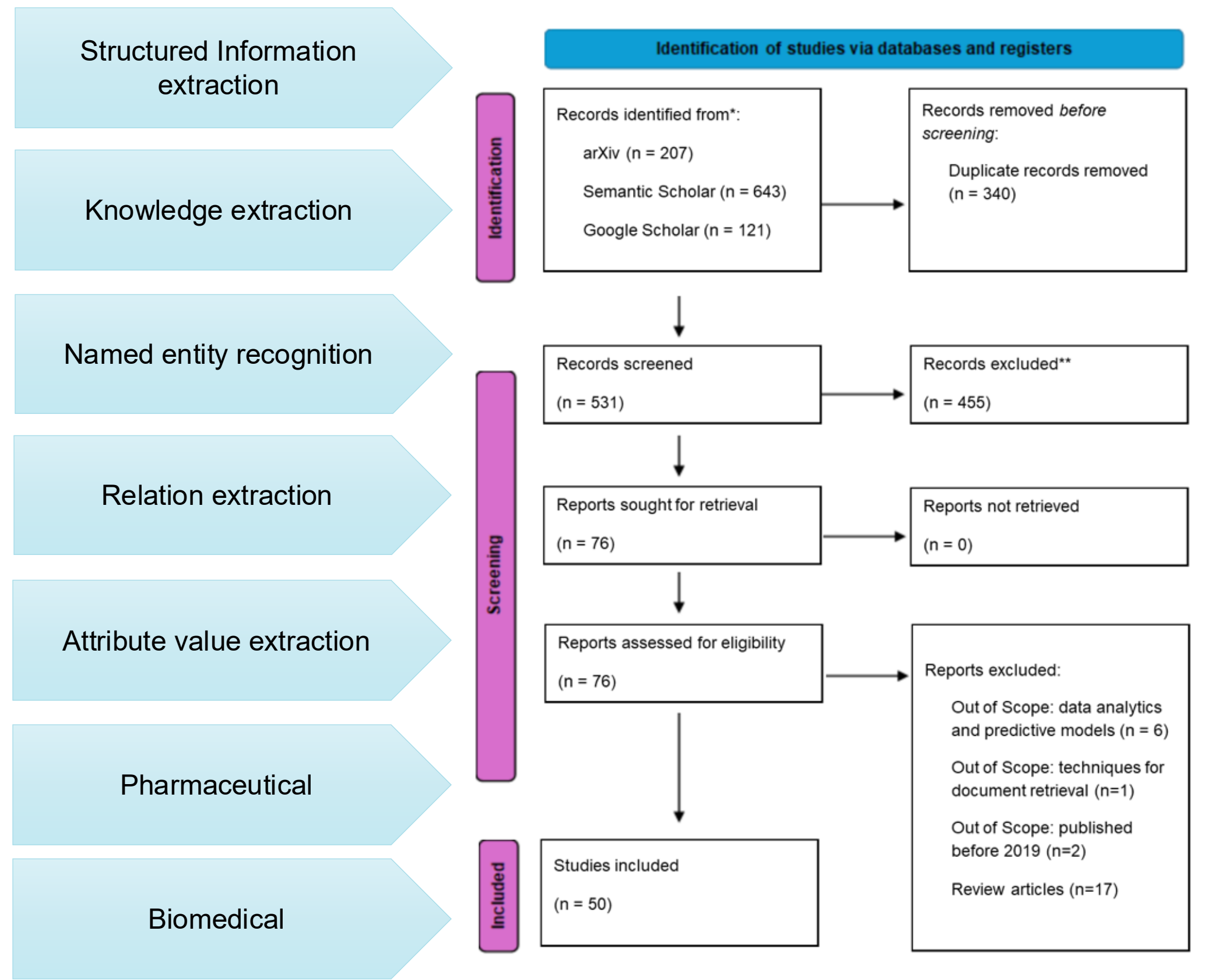
INTRODUCTION



Challenges in pharmaceutical knowledge extraction:



SURVEY STRATEGY



CONCLUSION

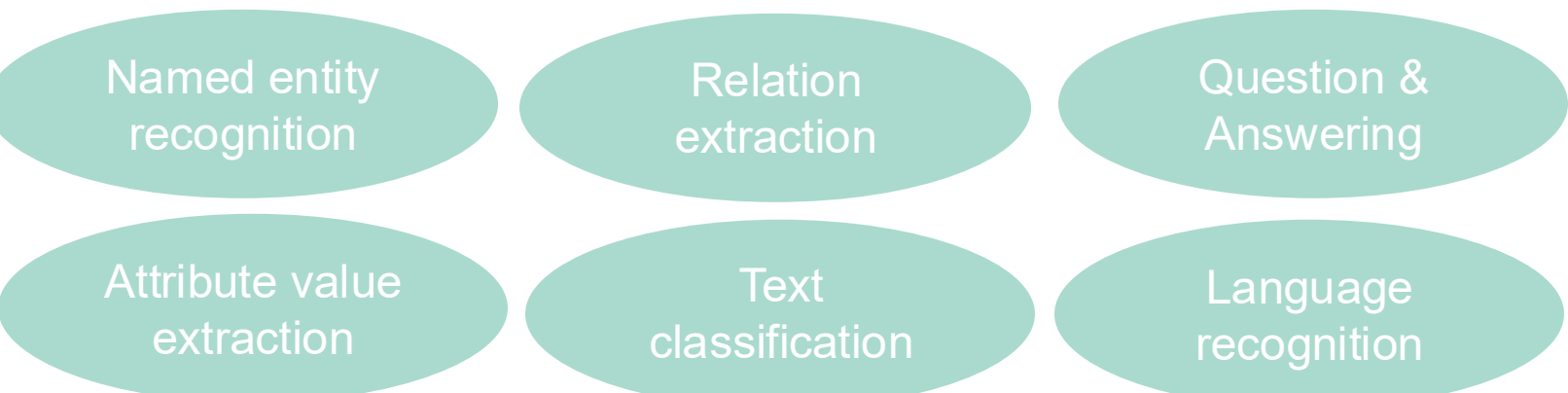
The results of this review informed model selection for applying AI-enabled information extraction to the STEP database, although direct evaluation on STEP data is ongoing.

References

- Neumann, M., et al., *ScispaCy: Fast and Robust Models for Biomedical Natural Language Processing*. *ArXiv*, 2019. [arXiv:1902.07669](https://arxiv.org/abs/1902.07669).
- Kocaman, V. and D. Talay, *Accurate Clinical and Biomedical Named Entity Recognition at Scale*. *Software Impacts*, 2022. 13: p. 100373. DOI: 10.1016/j.simpa.2022.100373.
- Sung, M., et al., *BERN2: an advanced neural biomedical named entity recognition and normalization tool*. *Bioinformatics*, 2022. 38(20): p. 4837-4839. DOI: 10.1093/bioinformatics/btad285.
- Lee, J., et al., *BioBERT: a pre-trained biomedical language representation model for biomedical text mining*. *Bioinformatics*, 2020. 36(4): p. 1234-1240. DOI: 10.1093/bioinformatics/btad282.
- Jahan, I., et al., *A comprehensive evaluation of large Language models on benchmark biomedical text processing tasks*. *Computers in Biology and Medicine*, 2024. 171: p. 108189.
- Li, M., et al., *BiomedRAG: A retrieval augmented large language model for biomedicine*. *Journal of Biomedical Informatics*, 2023. 162: p. 104769. DOI: 10.1016/j.jbi.2024.104769.
- Chen, Q., et al., *An extensive benchmark study on biomedical text generation and mining with ChatGPT*. *Bioinformatics*, 2023. 39(5): DOI: 10.1093/bioinformatics/btad557.
- Chellagurki, P., et al., *Biomedical Relation Extraction Using LLMs and Knowledge Graphs*. In *2024 IEEE 10th International Conference on Big Data Computing Service and Machine Learning Applications (BigDataService)*. 2024.
- Zhang, J., et al., *A Study of Biomedical Relation Extraction Using GPT Models*. *AMIA Jt Summits Transl Sci Proc*, 2024. 2024: p. 391-400.

RESULTS

Information Extraction Tasks:



Datasets and Evaluation Metrics:

Chemical & Drug Dataset	Disease & Gene/Protein Datasets	Clinical Datasets
BC4CHEMD, BC5CDR, DDI, CHEMPROT, KD-DTI	NCBI-Disease, Linnaeus, BC2GM, JNLPBA, GAD	THYME, EU-ADR

There are also uses of custom datasets, which involves manual or AI-assisted annotation, with information retrieved from clinical records, public databases, and scientific literature.

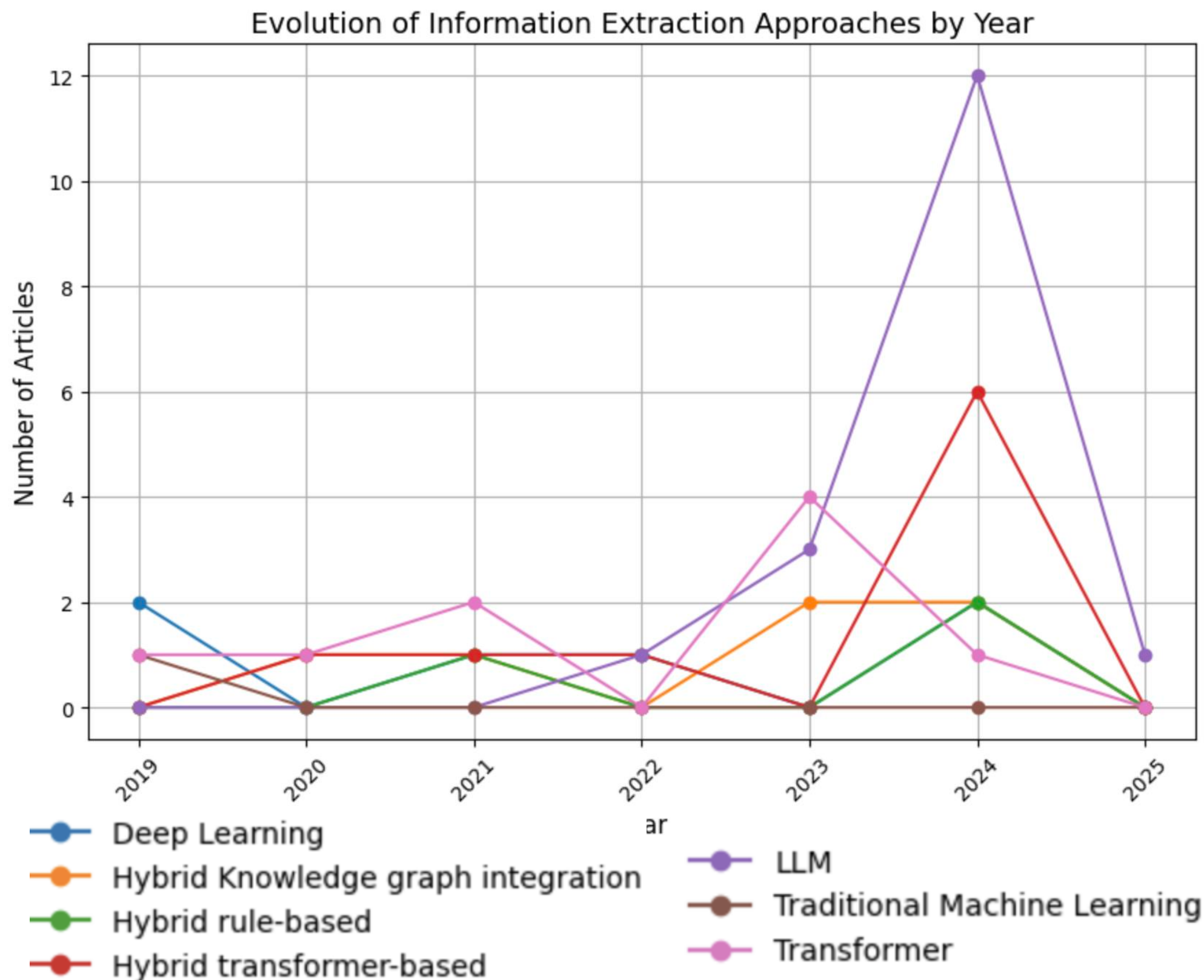
Models are assessed using standard NLP metrics:

Precision: The proportion of true positive results among all predicted positives.

Recall: The proportion of true positives identified out of all actual positives.

F1-score: The harmonic mean of precision and recall, widely used as the main performance indicator in biomedical NER and RE tasks.

Accuracy: Occasionally reported, though less common for imbalanced classification tasks like NER.



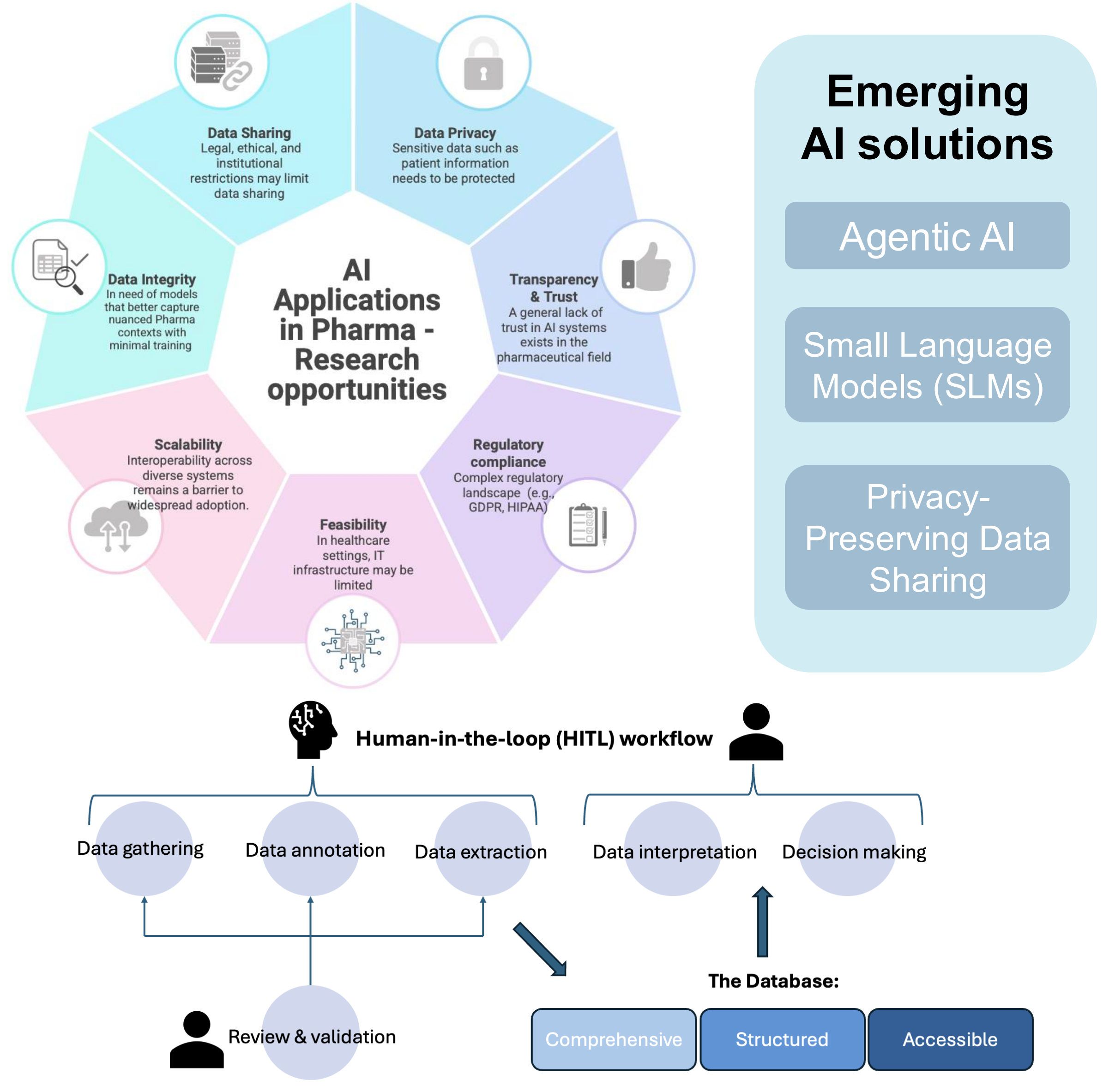
COMPARATIVE ANALYSIS & LIMITATIONS

	Strength	Limitations	Use Cases
Traditional ML	Fast, interpretable	Poor contextual modeling	Pharmaceutical production records processing
Deep Learning	Handling Structural Complexity	Compute-intensive, high data dependency	Chemical NER
Transformers	High accuracy	Expensive training	Medical/clinical/pharmaceutical NER
Hybrid	Balanced performance	Complex implementation	KG-based clinical/pharmaceutical RE
Generative LLMs	Few shots adaptability (Small samples)	Hallucinations, inconsistency	Attribute extraction from unstructured texts

Representative Metrics for Different Models on NER and RE Benchmarks

Architecture	Benchmark dataset	F1 score	References
ScispaCy sci md	BC4CHEMD	84.55	Neumann et al. (2019) ¹
BiLSTM-CNN-Char	BC4CHEMD	93.6	Kocaman et al. (2022) ²
BERN2	BC4CHEMD	92.8	Sung et al. (2022) ³
BioBERT	BC4CHEMD	92.36	Lee et al. (2020) ⁴
GPT-3.5 (zero-shot)	BC4CHEMD	26.01	Jahan et al. (2024) ⁵
LlaMA2-13b-BiomedRAG	DDI	80.00	Li et al. (2025) ⁶
ChatGPT (zero-shot)	DDI	51.62	Chen et al. (2023) ⁷
LlaMA2-7b (fine-tuned)	EU-ADR	93.18	Chellagurki et al. (2024) ⁸
GPT-3.5-turbo (zero-shot)	EU-ADR	80.9	Zhang et al. (2024) ⁹

EMERGING TRENDS & RESEARCH OPPORTUNITIES



Author's recommendations

- Use domain-specific pretraining to enhance relevance and accuracy.
- Incorporate normalization techniques and ontology alignment to ensure consistency across varied data sources
- Validate model outputs using a human-in-the-loop process or RAG to reduce errors and improve trust
- Fine-tune models to better capture the language and structure unique to pharmaceutical data.
- Combine symbolic approaches with neural models for better coverage and interpretability
- Evaluate models using multiple, domain appropriate benchmarks or create benchmark on custom data to assess generalizability and task performance.