

Forschungsprojekt «Open Justice vs. Privacy» und die Weiterentwicklung von Legal NLP

Daniel Kettiger, Kompetenzzentrum für Public Management der Universität Bern

Prof. Dr. Matthias Stürmer, Institut Public Sector Transformation der Berner Fachhochschule

Agenda

1. Ergebnisse des Forschungsprojekts aus juristischer und verwaltungswissenschaftlicher Sicht
2. Ergebnisse des Forschungsprojekts aus technischer Sicht
3. Zukunft von Legal NLP (Natural Language Processing)

Forschungsprojekt «Open Justice vs. Privacy»

Forschungsprojekt im SNF NFP 77 «Digitale Transformation»

- Interdisziplinär: Recht, Computer Science, Verwaltungswissenschaften
- Transdisziplinär: Einbezug von Akteuren rund um die Gerichtsbarkeit

Universität Bern

- Kompetenzzentrum für Public Management (Prof. Dr. Andreas Lienhard)
- Institut für Wirtschaftsinformatik (Prof. Dr. Thomas Myrach)
- Institut für Informatik (Prof. Dr. Matthias Stürmer)
+ Berner Fachhochschule, Institut Public Sector Transformation

Forschungsgegenstand

- Anonymisierung von Gerichtsurteilen



Digitale Transformation
Nationales Forschungsprogramm



1. Ergebnisse des Forschungsprojekts aus juristischer und verwaltungs- wissenschaftlicher Sicht

77
NRP

Digital Transformation
National Research Programme

Arbeiten im Rahmen des Forschungsprojekts

Juristische Forschung

- Dissertation «Publikation von Gerichtsurteilen im Spannungsfeld zwischen Transparenz und Geheimhaltungsinteressen», inkl. Befragung Gerichte CH und Rechtsvergleich Europa
- Masterarbeit: «Staatshaftung für mangelhafte Anonymisierung von publizierten Gerichtsurteilen » (Justice - Justiz - Giustizia 2022/1)
- Masterarbeit: «Datenschutzrechtliche Aspekte bei der Veröffentlichung von Gerichtsurteilen in sekundären Entscheidungsdatenbanken (Justice - Justiz - Giustizia 2024/2)

Verwaltungswissenschaftliche Forschung

- Expertenbefragung zu Aspekten der Anonymisierung von Gerichtsurteilen
- Expertenworkshop

Fazit des NFP77-Projekts

Einige **Schlussfolgerungen und Empfehlungen** aus unserem Forschungsprojekt:

1. Der **Begriff der Anonymisierung** im Gerichts-organisations- und Verfahrensrecht entspricht nicht demjenigen im Datenschutzrecht.
2. Der «interne» Grund für die Anonymisierung von Gerichtsurteilen ist die **Vulnerabilität (Verletzlichkeit) der betroffenen Personen**. Besonders schutzbedürftige Personen (z. B. Kinder, Behinderte) benötigen bei der Veröffentlichung von Urteilen einen besonders sorgfältigen Persönlichkeitsschutz.
3. Das Urteil muss in jedem Fall **nachvollziehbar** bleiben, auch wenn nicht ausgeschlossen ist, dass eine mit den Einzelheiten des Falles bereits vertraute Person den Namen einer Partei erkennen könnte.



Digitale Transformation
Nationales Forschungsprogramm



2. Ergebnisse des Forschungsprojekts aus technischer Sicht

77
NRP

Digital Transformation
National Research Programme

Grundkonzepte moderner KI-Modelle

Pre-Training

Generierung von
«Foundation Models»

→ Viele Daten (grosse Bild- und
Textmengen)
Viel Rechenkapazität (GPUs)
Viel Energie (MWh)



Fine-Tuning

Spezialisierung für
bestimmte
Aufgaben,
«Alignment»
→ kleine
Datenmengen in
hoher Qualität



Inference

Produktive Anwendung
im regulären Betrieb
→ Prompting

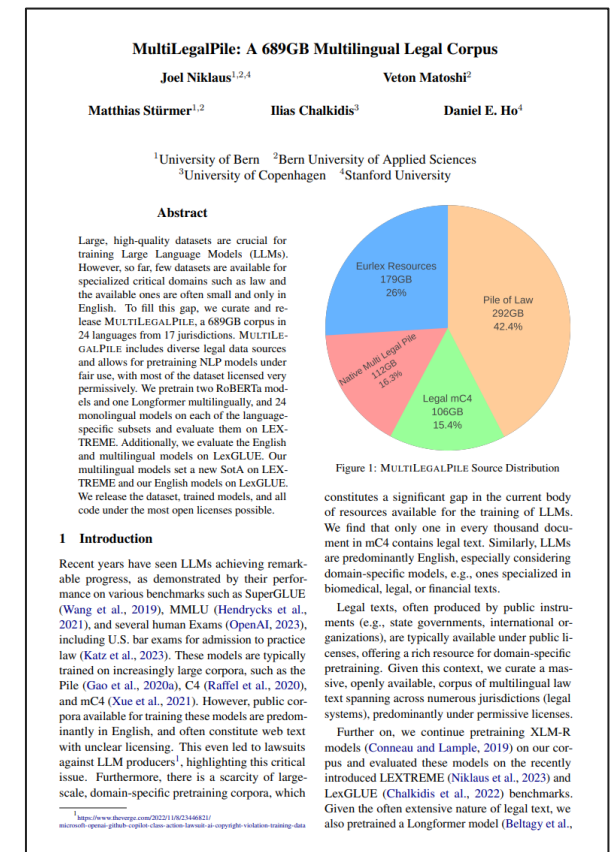
Mehrsprachige Rechtssprachen-Modell fürs Bundesgericht

Neues, juristisches Sprachmodell, das 2022 in einem Rechenzentrum von Google (TPU) auf Rechtstexten trainiert wurde (**Pre-Training**)

Auftrag des **Schweizerischen Bundesgerichts** zur besseren Anonymisierung von Gerichtsentscheiden

Dazu notwendig: Erstellung der **weltweit grössten mehrsprachigen Sammlung juristischer Texte**:

- Native Multi Legal Pile (112 GB)
- Eurlex-Ressourcen (179 GB)
- Legal MC4 (106 GB)
- Pile of Law (292 GB)



Weltweit grösste, mehrsprachige Rechtstextesammlung

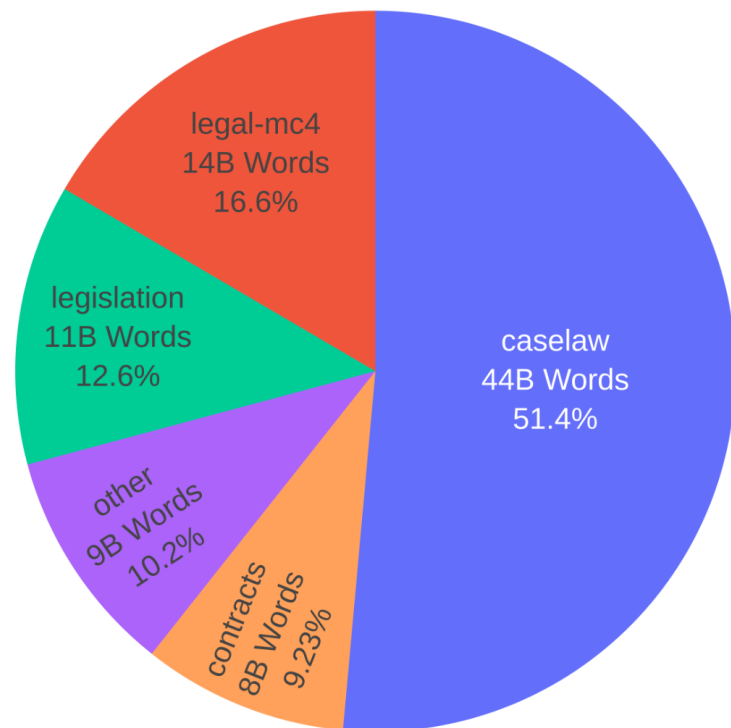


Figure 2. MULTILEGALPILE Text Type Distribution

Language	Text Type	Words	Documents	Words per Document	Jurisdiction	Source	License/Copyright
Native Multi Legal Pile							
bg	legislation	309M	262k	1178	Bulgaria	MARCELL	CC0-1.0
cs	caselaw	571M	342k	1667	Czechia Czechia Czechia	CzCDC Constitutional Court CzCDC Supreme Administrative Court CzCDC Supreme Court	CC BY-NC 4.0 CC BY-NC 4.0 CC BY-NC 4.0
da	caselaw	211M	92k	2275	Denmark	DDSC	CC BY 4.0 and other, depending on the dataset
da	legislation	653M	296k	2201	Denmark	DDSC	CC BY 4.0 and other, depending on the dataset
de	caselaw	1786M	614k	2905	Germany Switzerland	openlegaldata entscheidsuche	ODbL-1.0 similar to CC BY
de	legislation	513M	302k	1698	Germany Switzerland	openlegaldata lexfind	ODbL-1.0 not protected by copyright law
en	legislation	2539M	713k	3557	Switzerland UK	lexfind uk-lex	not protected by copyright law CC BY 4.0
fr	caselaw	1172M	495k	2363	Belgium France Luxembourg Switzerland	jurportal CASS judoc entscheidsuche	not protected by copyright law Open Licence 2.0 not protected by copyright law similar to CC BY
fr	legislation	600M	253k	2365	Switzerland Belgium	lexfind ejustice	not protected by copyright law not protected by copyright law
hu	legislation	265M	259k	1019	Hungary	MARCELL	CC0-1.0
it	caselaw	407M	159k	2554	Switzerland	entscheidsuche	similar to CC BY
it	legislation	543M	238k	2278	Switzerland	lexfind	not protected by copyright law
nl	legislation	551M	243k	2263	Belgium	ejustice	not protected by copyright law
pl	legislation	299M	260k	1148	Poland	MARCELL	CC0-1.0
pt	caselaw	12613M	17M	728	Brazil Brazil Brazil	RulingBR CRETA CJPG	not protected by copyright law CC BY-NC-SA 4.0 not protected by copyright law
ro	legislation	559M	396k	1410	Romania	MARCELL	CC0-1.0
sk	legislation	280M	246k	1137	Slovakia	MARCELL	CC0-1.0
sl	legislation	366M	257k	1418	Slovenia	MARCELL	CC-BY-4.0
total		24236M	23M	1065	Native Multi Legal Pile		
Overall statistics for the remaining subsets							
total		12107M	8M	1457	EU	Eurlex Resources	CC BY 4.0
total		43376M	18M	2454	US (99%), Canada, and EU	Pile of Law	CC BY-NC-SA 4.0; See Henderson* et al. (2022) for details
total		28599M	10M	2454		Legal mC4	ODC-BY

Neues Rechts-Sprachmodell für Bundesgericht trainiert



Startseite / Forschung + Dienstleistungen / Projekte / Training eines Swiss Long Legal BERT Modells

Training eines Swiss Long Legal BERT Modells

Wir werden juristische Texte in deutscher, französischer und italienischer Sprache scrapen, um ein Schweizer Long Legal BERT-Modell zu trainieren, das NLP-Aufgaben in der Schweizer Rechtsdomäne besser erfüllen kann.

Steckbrief

Lead-Departement Wirtschaft	Förderorganisation Andere	Projektleitung Joël Niklaus
Institut Institut Public Sector Transformation (IPST)	Laufzeit (geplant) 15.12.2021 - 31.12.2022	Projektmitarbeitende Alperen Bektas Veton Matoshi
Forschungseinheit Digital Sustainability Lab	Projektverantwortung Prof. Dr. Matthias Stürmer	Partner Schweizerisches Bundesgericht

Vergleich

Domain-spezifisches Modell
verstehet Rechtssprache besser
als generische Modell.

Beispiel: das Wort
«**Verhältnismässigkeit**» ist
maskiert («MASK»)

**Juristisches Rechtssprachen-
Modell** findet Wort sofort
(96% Wahrscheinlichkeit)

Generisches Sprachmodell
vermutet «Unabhängigkeit» (49%),
«Öffentlichkeit» (22%) oder
«Transparenz» (12%)

GPT Legal Question Answering Summarization Simplification

Roberta Legal Bert

[Roberta German](#)
[Bert](#)

Hier zeigen wir ein Modell, welches an unserem Institut darauf vortrainiert wurde, die rechtliche Sprache besser zu verstehen. Dabei wird in einem Text immer ein Wort (mit dem Stichwort) maskiert, und das Modell muss das fehlende Wort voraussagen. Dadurch, dass das Modell auf die rechtliche Sprache spezifiziert wurde, sind die Voraussagen deutlich besser, wie das nachfolgende Beispiel zeigt (BGE 142 II 268 S. 271, Erwägung 4.1): Unser spezialisiertes Modell gibt richtigerweise das Wort "Verhältnismässigkeit" aus, während ein generisches XLM-RoBERTa-Modell deutlich allgemeinere Wörter wie Freiheit, Demokratie oder Öffentlichkeit voraussagt. Beide Modelle haben 354 Millionen Parameter.

Input

Die Beschwerdeführerin rügt sodann eine Verletzung des Verhältnismässigkeitsprinzips. Sie ist der Auffassung, dass die Publikationstätigkeit der WEKO den Grundsatz der <mask> gemäss Art. 5 Abs. 2 und Art. 36 BV wahren müsse.

Löschen Absenden

Classification

verhältnismässigkeit

verhältnismässigkeit	96%
subsidiarität	1%
gleichbehandlung	0%
rechtssicherheit	0%
verhältnismaßigkeit	0%

Input

Die Beschwerdeführerin rügt sodann eine Verletzung des Verhältnismässigkeitsprinzips. Sie ist der Auffassung, dass die Publikationstätigkeit der WEKO den Grundsatz der [MASK] gemäss Art. 5 Abs. 2 und Art. 36 BV wahren müsse.

Löschen Absenden

Classification

Unabhängigkeit

Unabhängigkeit	49%
Öffentlichkeit	22%
Transparenz	12%
Privatsphäre	3%
Information	2%

Vergleich mit ChatGPT

Kleine, Domain-spezifische
Modelle können sogar besser als
ChatGPT sein!

Beispiel: das Wort
«**Arbeitsgericht**» ist maskiert
(«MASK»)

**Juristisches Rechtssprachen-
Modell** findet Wort sofort
(97% Wahrscheinlichkeit)

ChatGPT vermutet «Bank»...

Einführung GPT2 Legal Named Entity Recognition Zero-Shot Interface

Legal German RoBERTa German BERT

Legal German RoBERTa

Statt ein öffentlich verfügbares Foundation-Modell zu verwenden, können wir auch selbst solche Modelle trainieren. Auf dieser Seite zeigen wir ein solches Modell, welches an unserem Institut auf einem Textkorpus mit rechtlicher Sprache vortrainiert wurde. Dadurch soll das Modell Textverarbeitungsaufgaben im Legal- Bereich deutlich besser lösen können als generische Modelle.

Die zwei hier gezeigten Foundation-Modelle wurden etwas anders trainiert als das vorher gezeigte GPT2-Modell: Wir nehmen einen Text und maskieren ein Wort davon (mit dem Stichwort

Input

In der Folge leitete E. beim <mask> Zürich zwei Prozesse gegen seine frühere Arbeitgeberin ein. Im einen verlangte er Lohnfortzahlung, Genugtuung und die Ausstellung eines Arbeitszeugnisses; im andern Ersatz für Schaden, der ihm aus dem Verhalten der Bank entstanden sei.

Clear Submit

Classification

arbeitsgericht 97%

bezirksgericht 3%

gericht 0%

kanton 0%

verwaltungsgericht 0%

ChatGPT 4o mini ▾

was ist das maskierte wort?

In der Folge leitete E. beim <mask> Zürich zwei Prozesse gegen seine frühere Arbeitgeberin ein. Im einen verlangte er Lohnfortzahlung, Genugtuung und die Ausstellung eines Arbeitszeugnisses; im andern Ersatz für Schaden, der ihm aus dem Verhalten der Bank entstanden sei.

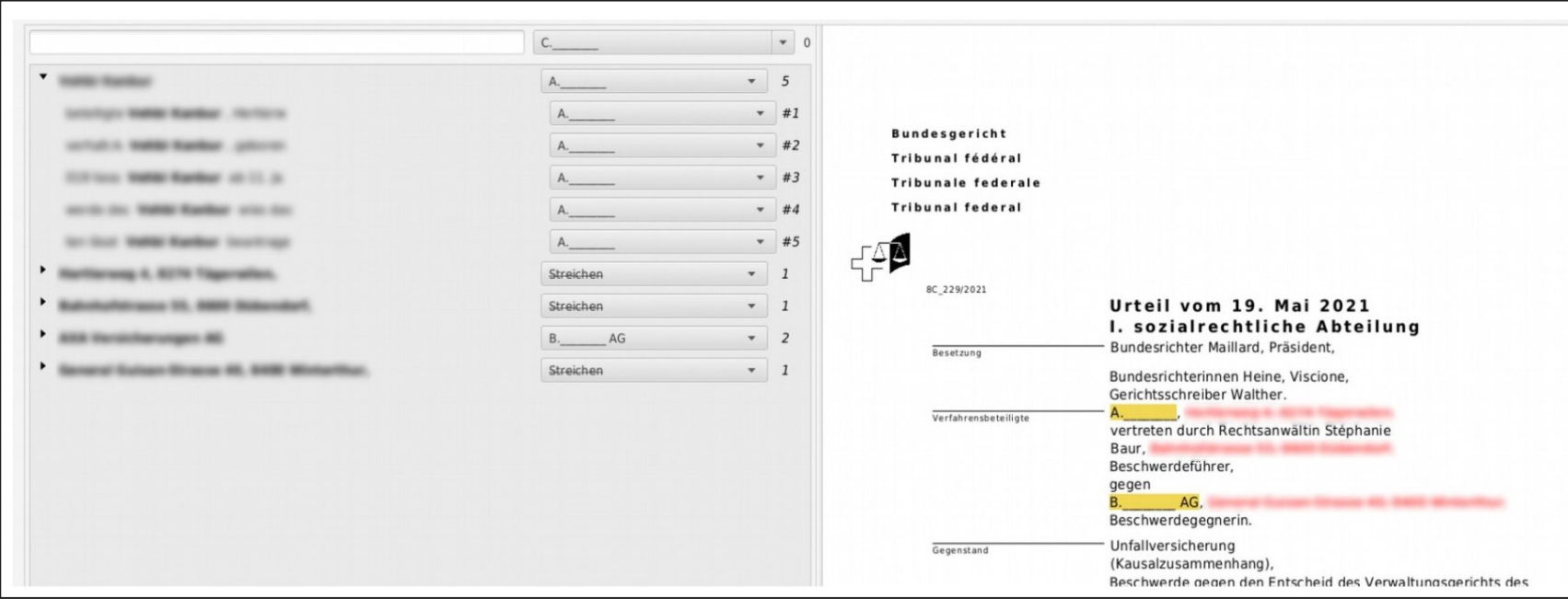
Das maskierte Wort in dem Satz ist wahrscheinlich **"Bank"**. Der Satz lautet dann:

"In der Folge leitete E. beim **Bank** Zürich zwei Prozesse gegen seine frühere Arbeitgeberin ein. Im einen verlangte er Lohnfortzahlung, Genugtuung und die Ausstellung eines Arbeitszeugnisses; im andern Ersatz für Schaden, der ihm aus dem Verhalten der Bank entstanden sei."

📄 🌟 🗣️ 🔊 ↺

Verbesserung der Anonymisierung mittels NLP

Betrieb von KI-Modell durch interne Informatikabteilung des Schweizer Bundesgerichts für Anonymisierungs-Vorschläge



The screenshot shows a web-based interface for anonymizing court rulings. On the left, there is a sidebar with a list of cases. The main area is divided into two columns. The left column contains a list of cases with dropdown menus for selecting entities to anonymize (e.g., 'A.', 'B.', 'AG'). The right column displays the anonymized text of a ruling from May 19, 2021. The ruling text is partially redacted with yellow boxes, indicating the areas where the model has suggested anonymization.

Automatic Anonymization of Swiss Federal Supreme Court Rulings

Joel Niklaus^{1,2,3 *}

Robin Mamié^{4 *}

Matthias Stürmer^{1,2}

Daniel Brunner⁴

Marcel Gygli²

¹University of Bern ²Bern University of Applied Sciences
³Stanford University ⁴Swiss Federal Supreme Court

Abstract

Releasing court decisions to the public relies on proper anonymization to protect all involved parties, where necessary. The Swiss Federal Supreme Court relies on an existing system that combines different traditional computational methods with human experts. In this work, we enhance the existing anonymization software using a large dataset annotated with entities to be anonymized. We compared BERT-based models with models pre-trained on in-domain data. Our results show that using in-domain data to pre-train the models further improves the F1-score by more than 5% compared to existing models. Our work demonstrates that combining existing anonymization methods, such as regular expressions, with machine learning can further reduce manual labor and enhance automatic suggestions.

already supported in their work through an application called *Anom2* (see Figure 1). *Anom2* provides access to various methods and algorithms for finding and replacing text entities (e.g., with regular expressions). The aim of this work is to enhance the capabilities of *Anom2* with Machine Learning capabilities that provide the user with more suggestions that need to be anonymized. Our results show that this approach allows users to find more elements that require anonymization.

2 Related Work

For identifying elements that might require anonymization, a process called Named Entity Recognition (NER) is employed. Traditionally, NER recognizes and categorizes text parts according to a set of semantic categories like *Location* (LOC), *Organization* (ORG), or *Person* (PER) (Benikova et al., 2014). As these classes are not enough for the anonymization of court cases (Leitner et al., 2020) suggested expanding this list to seven coarse and 19 fine-grained classes, including entities such as *Judge* (RR), or *Lawyer* (AN). Using this dataset (Dajti et al., 2023) fine-tuned GermanBERT (Chen et al., 2020), clearly outperforming a BiLSTM-CRF+ model. Similar approaches have been applied and tested in other languages, such as Romanian (Pais et al., 2021), Greek (Angelidis et al., 2018), Portuguese (Luz de Araujo et al., 2018), and multilingually (de Gibert et al., 2022; Niklaus et al., 2023a).

Domain-specific pretraining has flourished in the legal domain recently. Chalkidis et al. (2020) pretrained LegalBERT on EU and UK legislation, ECHR and US cases, and US contracts. Zheng et al. (2021) pretrained CaseHoldBERT on US case law, while Henderson et al. (2022) trained PoL-BERT on the 256 GB Pile of Law corpus. Niklaus and Gjofer (2022) pretrained Longformer (Beltagy et al., 2020) models using the Replaced

* Equal contribution.

159

Proceedings of the Natural Legal Language Processing Workshop 2023, pages 159–165
December 7, 2023 ©2023 Association for Computational Linguistics

Joel Niklaus, Robin Mamié, Matthias Stürmer, Daniel Brunner, und Marcel Gygli 2023 „Automatic Anonymization of Swiss Federal Supreme Court Rulings“.

In Proceedings of the Natural Legal Language Processing Workshop 2023, 159–65. Singapore: Association for Computational Linguistics, 2023.

<https://doi.org/10.18653/v1/2023.nllp-1.16>

Prüfen der De-Anonymisierungs-Risiken mit LLMs

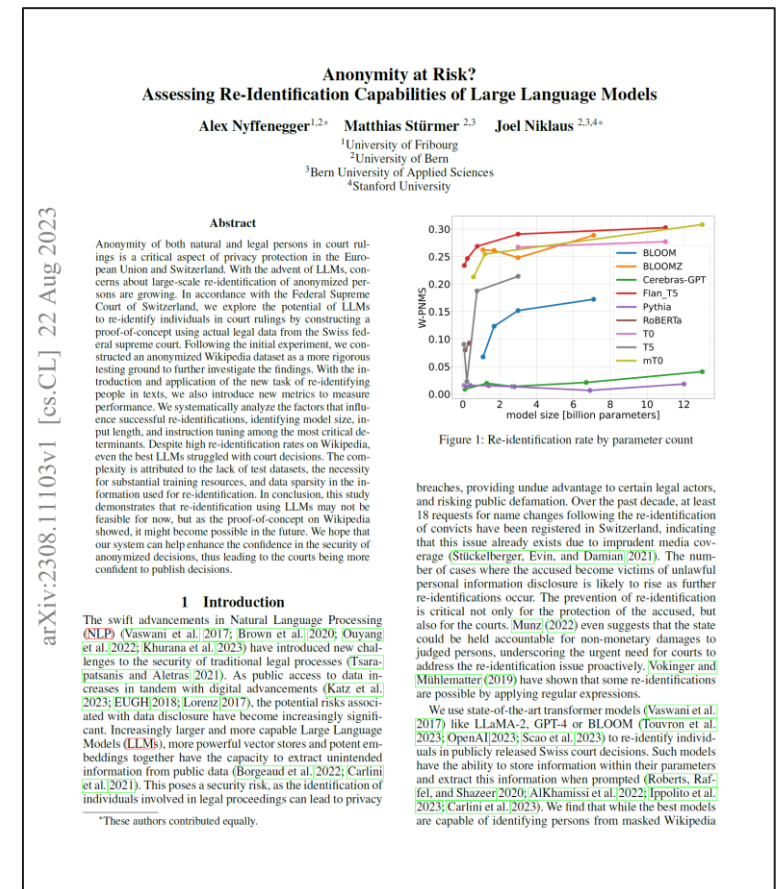
Experimente mit der Re-Identifizierung von Gerichtsurteilen und Wikipedia-Seiten

Testen verschiedener De-Anonymisierungs-Tasks:

- „Maske ausfüllen“
- „Fragenbeantwortung“
- „Text-Erzeugung“

Erstellung von Metriken zur Messung der De-Anonymisierung:

- Partial Name Match Score (PNMS)
- Normalisierte Levenshtein-Distanz (NLD)
- Ergebnis der Nachnamenübereinstimmung (LNMS)
- Gewichtete Teilnamensübereinstimmung (W-PNMS)



Fazit des NFP77-Projekts

Einige **Schlussfolgerungen und Empfehlungen** aus unserem Forschungsprojekt:

1. Zurzeit besteht **kein Risiko** der vollautomatisierten De-Anonymisierung von Gerichtsurteilen durch künstliche Intelligenz (KI)
2. Um künftige Re-Identifikation zu verhindern, wird der Einsatz von **KI-unterstützten Anonymisierungs-Tools** empfohlen
3. Das **technische Potenzial** und damit das Risiko zur Re-Identifizierung ist systematisch zu **überwachen** und periodisch zu evaluieren



Digitale Transformation
Nationales Forschungsprogramm



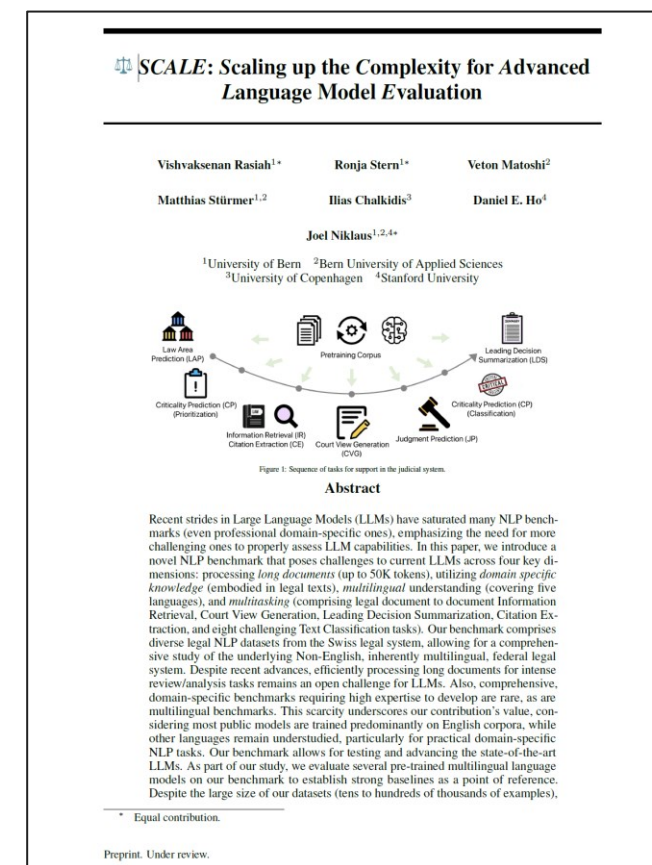
3. Zukunft von Legal NLP (Natural Language Processing)



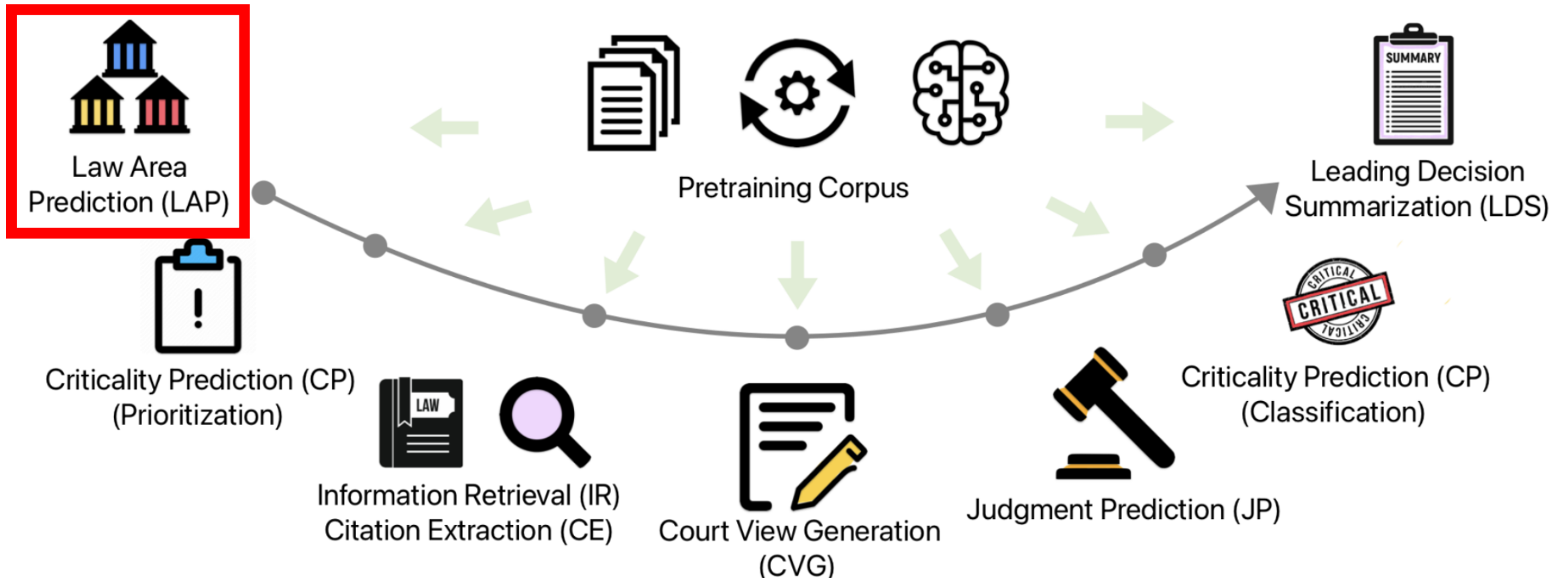
Benchmarks (Qualitätsmessung) typischer Gerichtsaufgaben

Neue Sammlung juristischer NLP-Benchmarks mit einzigartigen Merkmalen:

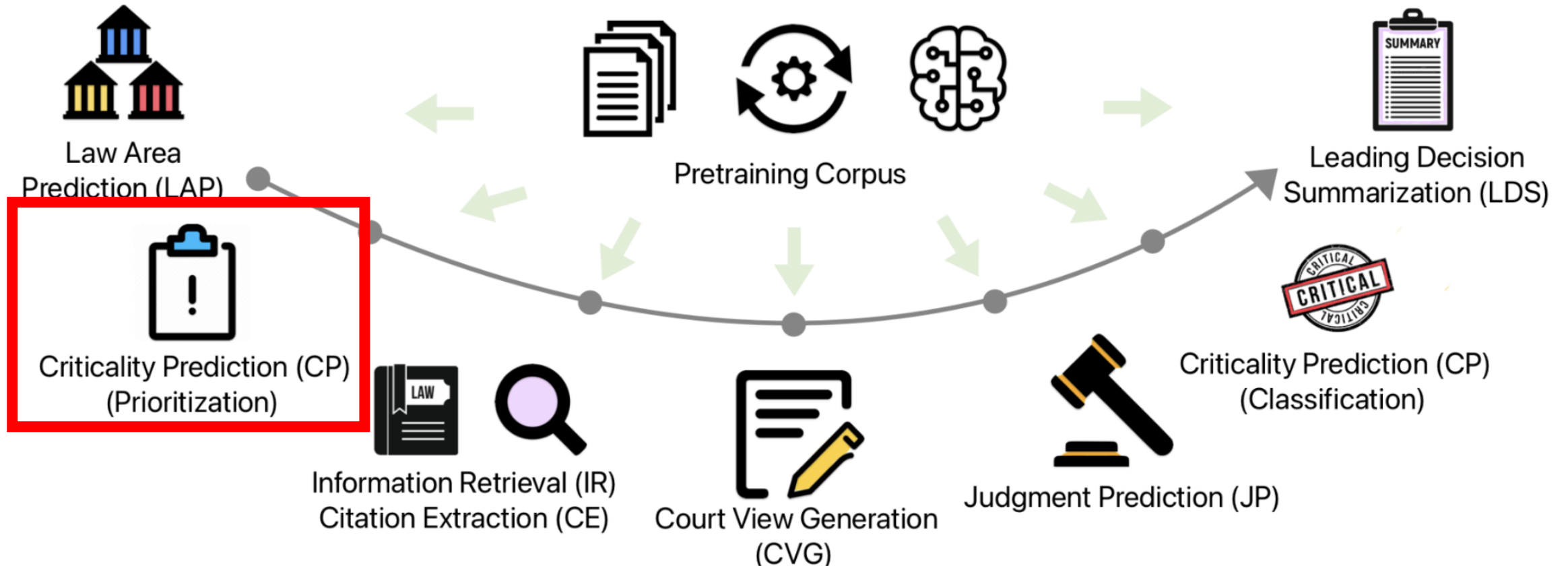
1. **Verarbeitung langer Dokumente**
→ bis zu 50'000 Tokens (Wortteile)
2. **Domänenspezifisches Wissen**
→ in juristischen Texten verankert
3. **Mehrsprachiges Textverständnis**
→ in fünf Sprachen (DE, FR, IT, RM, EN)
4. **Multitasking (unterschiedliche Aufgaben)**
→ Information Retrieval, Court View Generation, Leading Decision Summarization, Citation Extraction, Text Classification



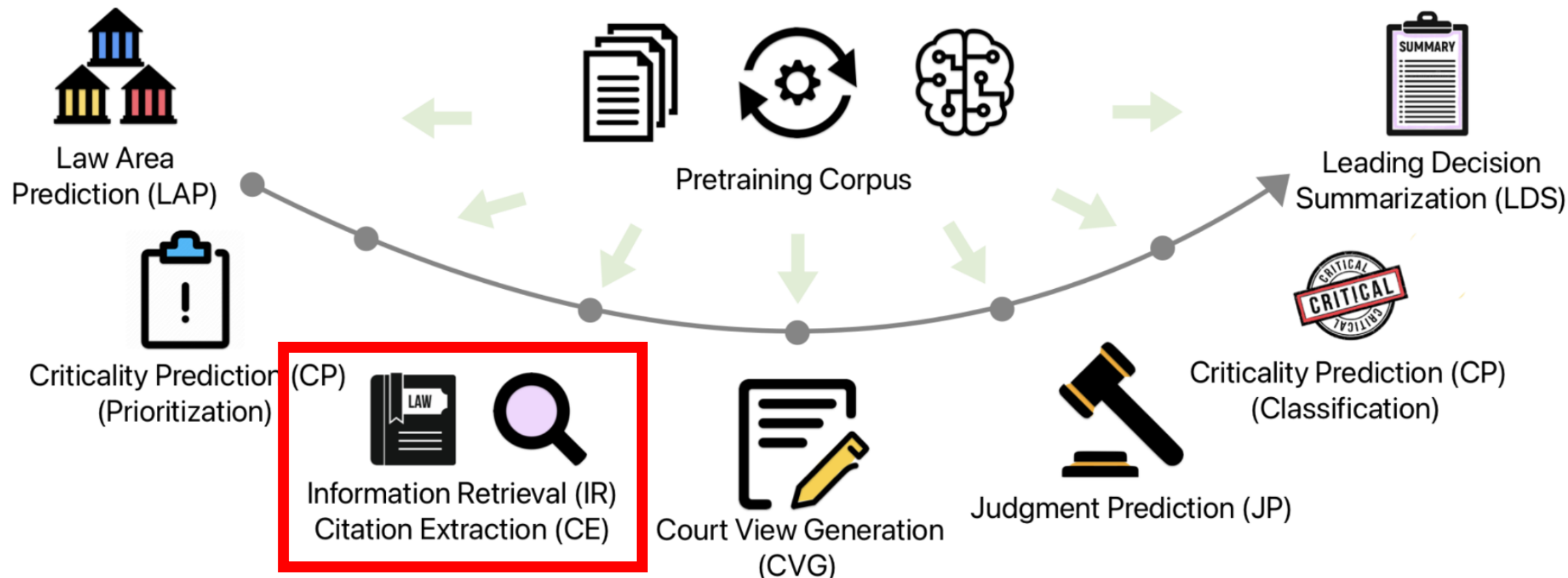
Legal NLP für Gerichte: 1) Vorhersage des Rechtsgebiets



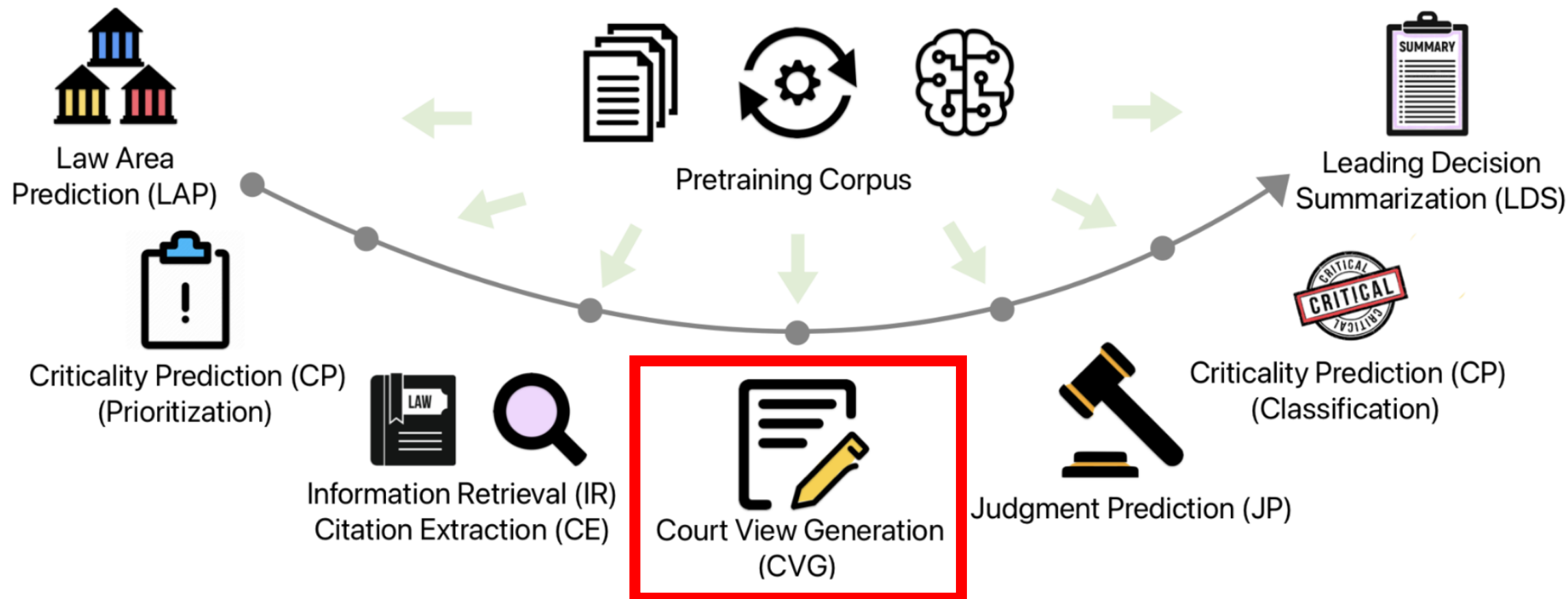
Legal NLP für Gerichte: 2) Vorhersage der Relevanz (Priorität)



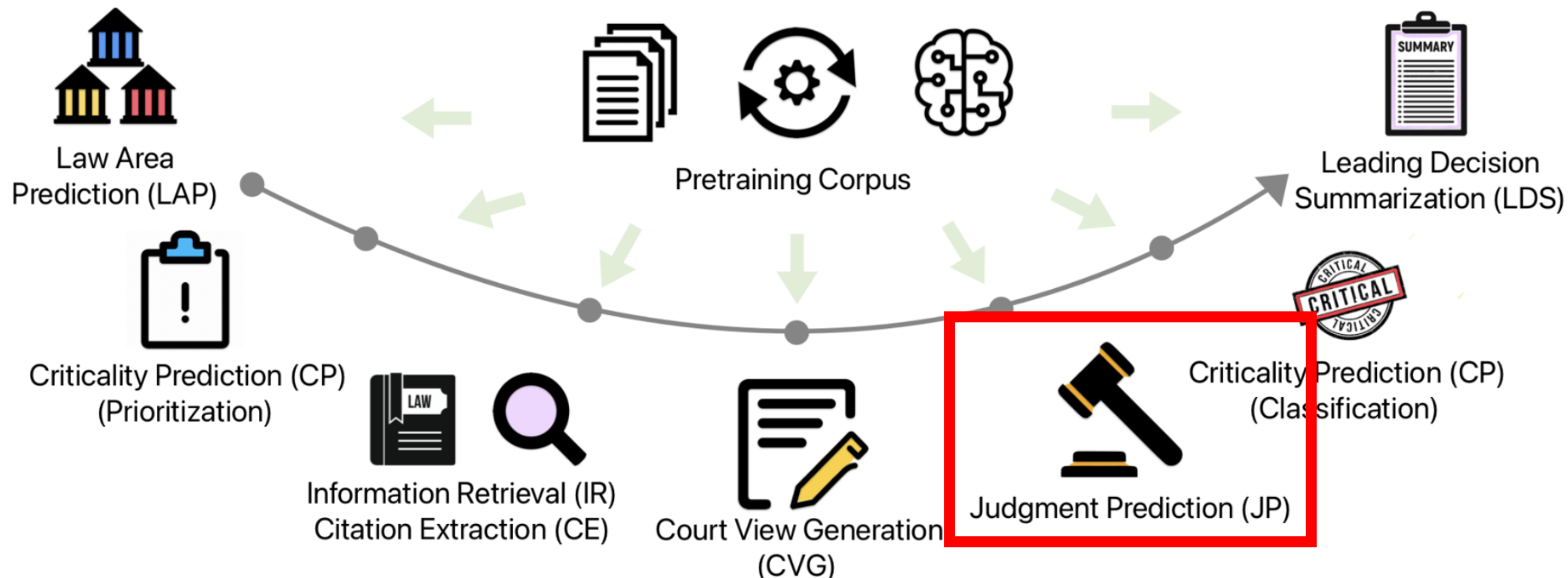
Legal NLP für Gerichte: 3) Informations-Extraktion (Suche)



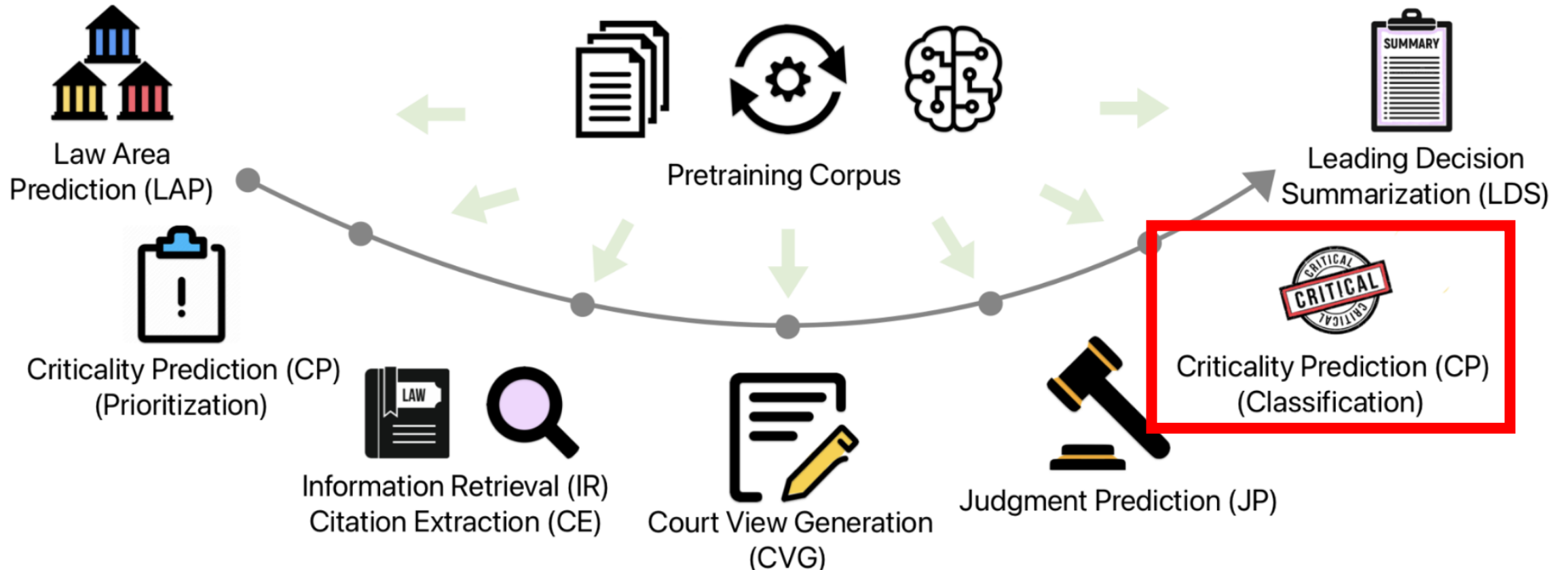
Legal NLP für Gerichte: 4) Generierung der Gerichtsperspektive



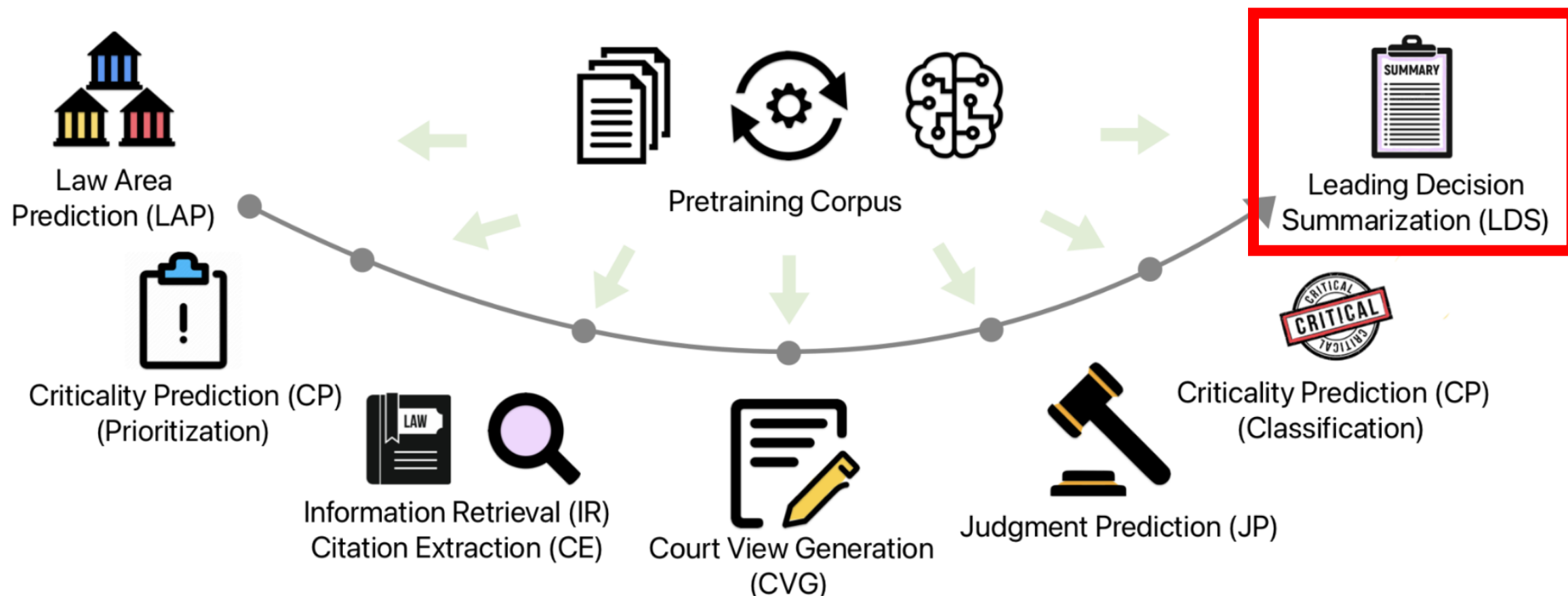
Legal NLP für Gerichte: 5) Urteilsvorhersage



Legal NLP für Gerichte: 6) Relevanz (Klassifizierung von Urteilen)



Legal NLP für Gerichte: 7) Entscheidungszusammenfassung



Legal Reasoning

- **Benchmark** (=Testing) von rechtlichen Schlussfolgerungen
- 162 Aufgaben in **6 Fachgebieten**, erstellt von Jurist:innen:
 1. Problemerkennung
 2. Regelaufwurf
 3. Regelanwendung
 4. Regelschluss
 5. Interpretation
 6. rhetorisches Verständnis
- Prüfung von **20 Open Source und proprietären KI-Modellen**

LEGALBENCH: A Collaboratively Built Benchmark for Measuring Legal Reasoning in Large Language Models

Neel Guha¹, Julian Nyarko^{1,1}, Daniel E. Ho^{1,1}, Christopher Ré^{1,1}, Adam Chilton², Aditya Narayana³, Alex Chohlas-Wood⁴, Austin Peters⁵, Brandon Waldon¹, Daniel N. Rockmore⁶, Diego Zambrano⁷, Dmitry Tatishnan⁸, Enam Hoque⁹, Faiz Surani¹, Frank Fagan¹⁰, Galit Sarfaty⁷, Gregory M. Dickinson¹¹, Haggai Porat¹², Jason Hegland¹, Jessica Wu¹, Joe Nudell¹, Joel Niklaus¹, John Nay¹⁰, Jonathan H. Choi¹¹, Kevin Tobia¹², Margaret Hagan¹³, Megan Ma¹⁰, Michael Livermore¹⁴, Nikon Rasumov-Rahe¹, Nils Holzenberger¹⁵, Noam Kolt¹, Peter Henderson¹, Sean Rehaag¹⁶, Sharad Goel¹⁷, Shang Gao¹⁸, Spencer Williams¹⁸, Sunny Gandhi¹⁹, Tom Zur⁸, Varun Iyer¹, and Zehua Li¹

¹Stanford University, ²University of Chicago, ³Maxime Tools, ⁴Dartmouth College, ⁵LawBeta, ⁶South Texas College of Law Houston, ⁷University of Toronto, ⁸St. Thomas University Benjamin L. Crump College of Law, ⁹Harvard Law School, ¹⁰Stanford Center for Legal Informatics - CodeX, ¹¹University of Southern California, ¹²Georgetown University Law Center, ¹³Stanford Law School, ¹⁴University of Virginia, ¹⁵Télécom Paris, Institut Polytechnique de Paris, ¹⁶Osgoode Hall Law School, York University, ¹⁷Harvard Kennedy School, ¹⁸Golden Gate University School of Law, ¹⁹Luddy School of Informatics - Indiana University Bloomington, ²⁰Casetext

Abstract

The advent of large language models (LLMs) and their adoption by the legal community has given rise to the question: what types of legal reasoning can LLMs perform? To enable greater study of this question, we present LEGALBENCH: a collaboratively constructed legal reasoning benchmark consisting of 162 tasks covering six different types of legal reasoning. LEGALBENCH was built through an interdisciplinary process, in which we collected tasks designed and hand-crafted by legal professionals. Because these subject matter experts took a leading role in construction, tasks either measure legal reasoning capabilities that are practically useful, or measure reasoning skills that lawyers find interesting. To enable cross-disciplinary conversations about LLMs in the law, we additionally show how popular legal frameworks for describing legal reasoning—which distinguish between its many forms—correspond to LEGALBENCH tasks, thus giving lawyers and LLM developers a common vocabulary. This paper describes LEGALBENCH, presents an empirical evaluation of 20 open-source and commercial LLMs, and illustrates the types of research explorations LEGALBENCH enables.

1 Introduction

Advances in large language models (LLMs) are leading American legal professionals to reexamine the practice of law [52, 61, 158, 53, 12] [3]. Proponents have argued that LLMs could alter how lawyers approach tasks ranging from brief writing to corporate compliance [158]. By making legal services more accessible, they could eventually help alleviate the United States' long-standing access-to-justice crisis [33, 132]. This perspective is informed by the observation that LLMs possess special properties

¹In using "LLMs", we are referring to language models which evince in-context learning capabilities (also referred to as "foundation models" [43]). This behavior has traditionally been observed in models with at least a billion parameters.



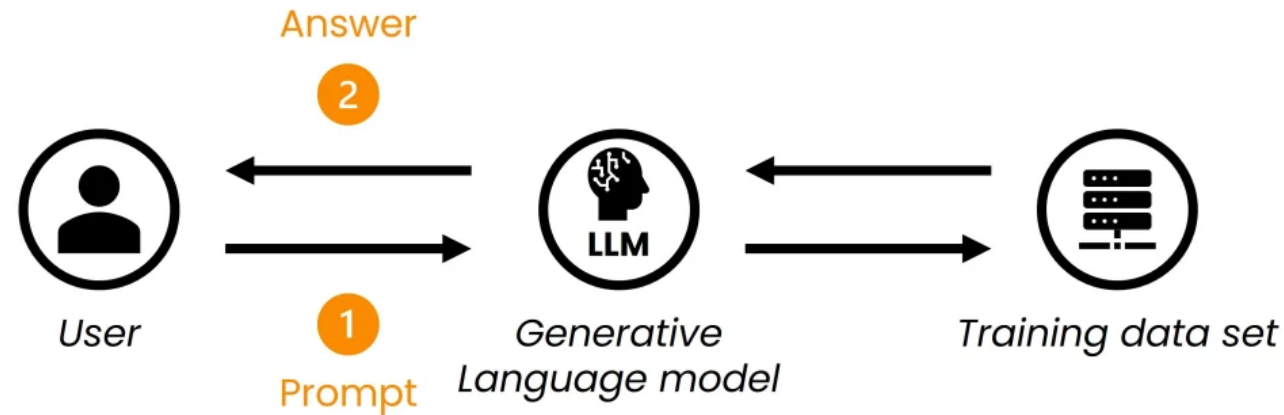
Vielen Dank für Ihre
Aufmerksamkeit

eJustice.ch · Postfach · 3001 Bern · +41 58 462 47 23 · info@ejustice.ch



Backup Folien

Reguläres Prompting von einem LLM (bspw. ChatGPT)



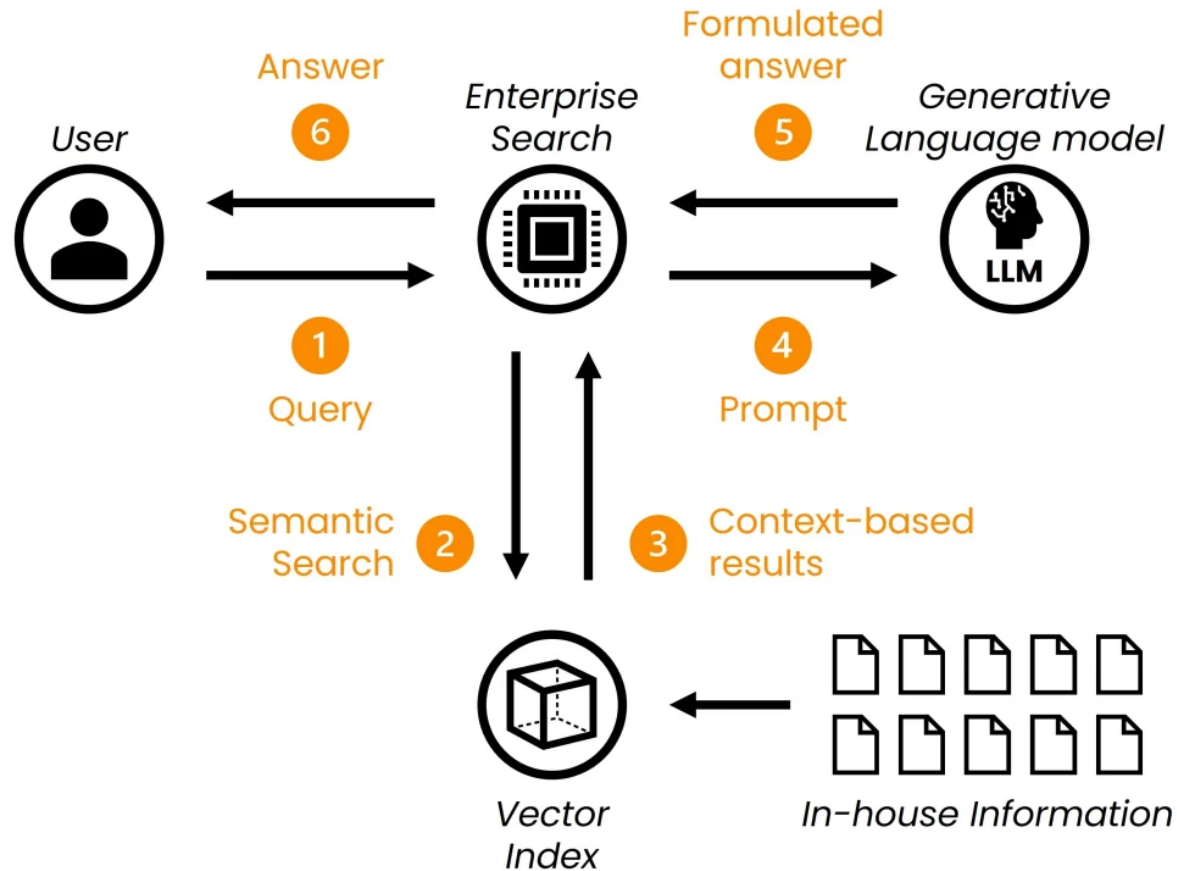
Probleme:

Kombiniert nach Wahrscheinlichkeit → halluziniert, nicht zuverlässig

Kann keine Fragen mit Wissen ausserhalb der Trainingsdaten beantworten, nur Fragen mit Informationen, die im Trainingsdatensatz enthalten sind

Kann keine Quellen zu den Antworten hinzufügen → unklar woher Informationen

Retrieval Augmented Generation (RAG) Verfahren hat Potenzial



Vorteile von RAG:

- **Wenig Halluzinationen**
→ LLM dient nur zur Formulierung des Textes und erstellt keine neuen Inhalte
- **Quellenangaben** möglich
→ Herkunft der Informationen ist verfügbar
- **Einfache Aktualisierung** der Informationen → aktuelle Daten schnell verfügbar
- **Zugriffsrechte und Relevanz** → nur erlaubte Informationen zugänglich