



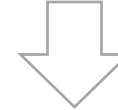
Making Large Information Resources Machine-readable with Machine Learning

Carlos Soto
Brookhaven National Laboratory

Nov 16, 2023

Outline

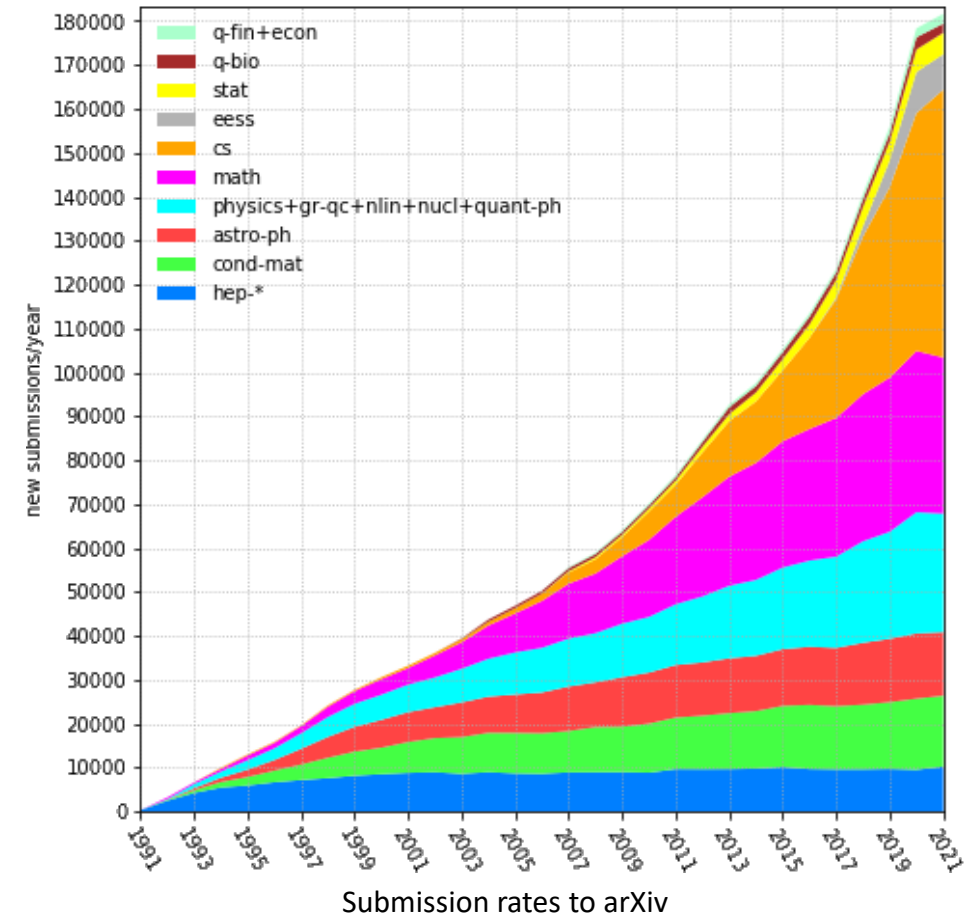
- Information at Scale
- Successes with Text and Language
- Documents
- Tables
- Figures & Charts
- A Vision for ML-driven Knowledge Curation



```
{
  "orders": [
    {
      "orderno": "748745375",
      "date": "June 30, 2088 1:54:23 AM",
      "trackingno": "TN0039291",
      "custid": "11045",
      "customer": [
        {
          "custid": "11045",
          "fname": "Sue",
          "lname": "Hatfield",
          "address": "1409 Silver Street",
          "city": "Ashland",
          "state": "NE",
          "zip": "68003"
        }
      ]
    }
  ]
}
```

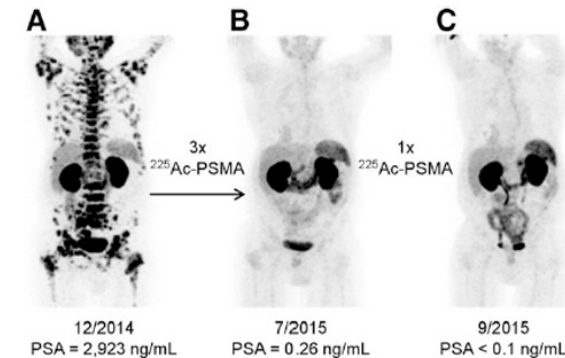
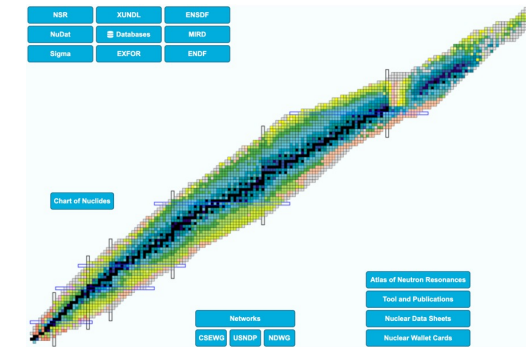
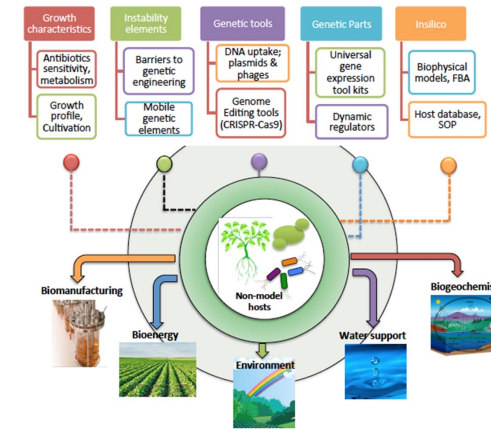

Information at Scale

- Many missions and efforts require analyzing vast information resources, much of whose contents are **not machine-readable**. These resources:
 - Inform the design, targets, and parameters of **new experiments**
 - Serve as **context** for interpreting and comparing results
 - Offer opportunity to perform **large-scale analyses**
- **Massive scale** limits effective use, inhibits full potential
 - e.g., millions of existing articles, publication rates accelerating
 - **Impossible to manually survey** all relevant works



Some of my Motivating Applications

- Biology
 - Genetic toolkit and protein-protein interaction (PPI) **discovery**
- Chemistry
 - Parameterization of isotopes separations **experiments**
- Nuclear Physics
 - **Curation** of experimental nuclide structure quantities into authoritative resource
- Commonalities:
 - Overcome **highly manual** workflows
 - Enable observation of **emergent patterns**



Top-left: BER KBase project components and applications; top-right: visualization of BNL NNDC's managed nuclide data; bottom: highly successful prostate cancer clinical trial with Actinium-225, one of the isotopes of interest at BNL's Isotope Research and Production Department.

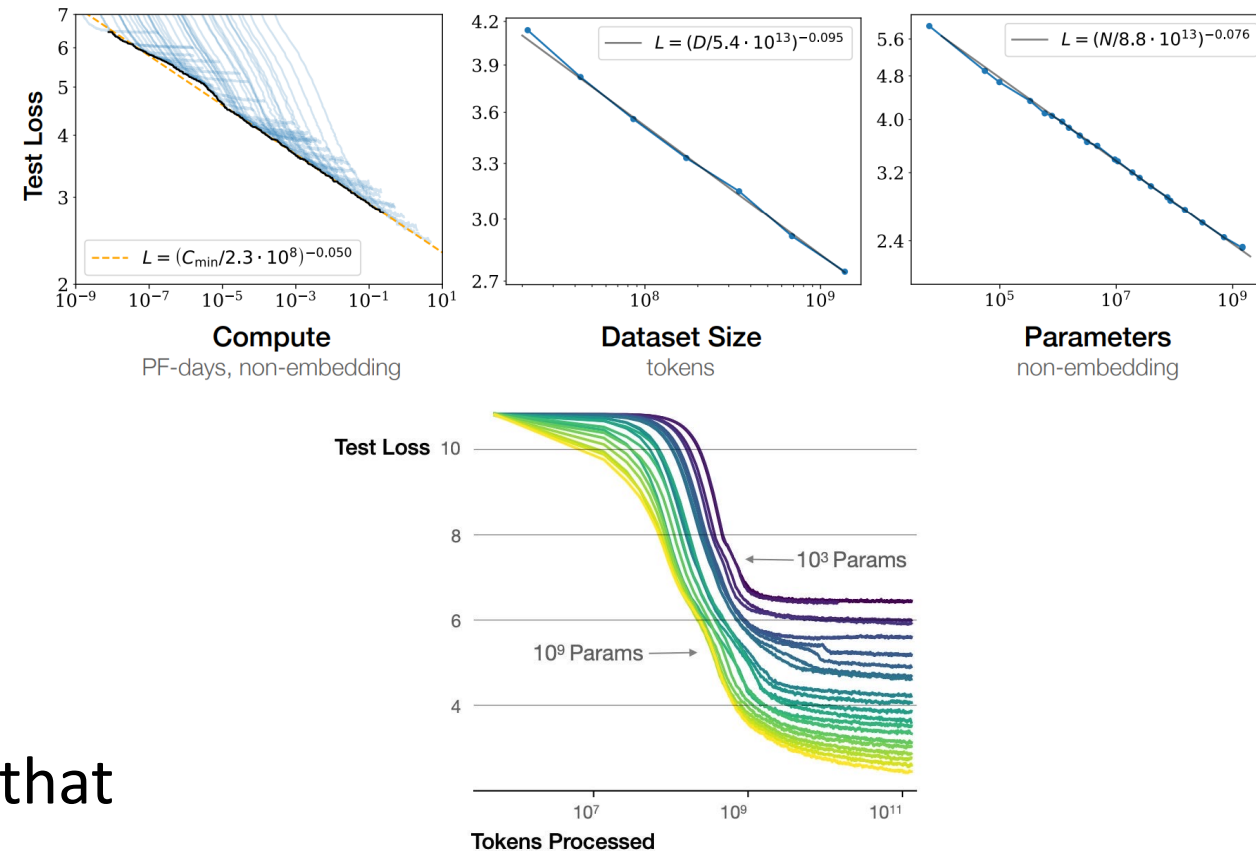
A highly visible success story: Natural Language Processing

- Large Language Models (**LLMs**) have demonstrated great **knowledge retrieval** capabilities; kicked off new AI race
 - Transformer architecture (2017) enabled scale
 - First LLMs in 2018 (GPT-1, BERT)
 - Public attention in 2022 (ChatGPT)
- LLMs are pretrained with simple language modeling objective using **vast text corpora**
 - Limited to *machine readable text*..



Why Scale Matters for NLP

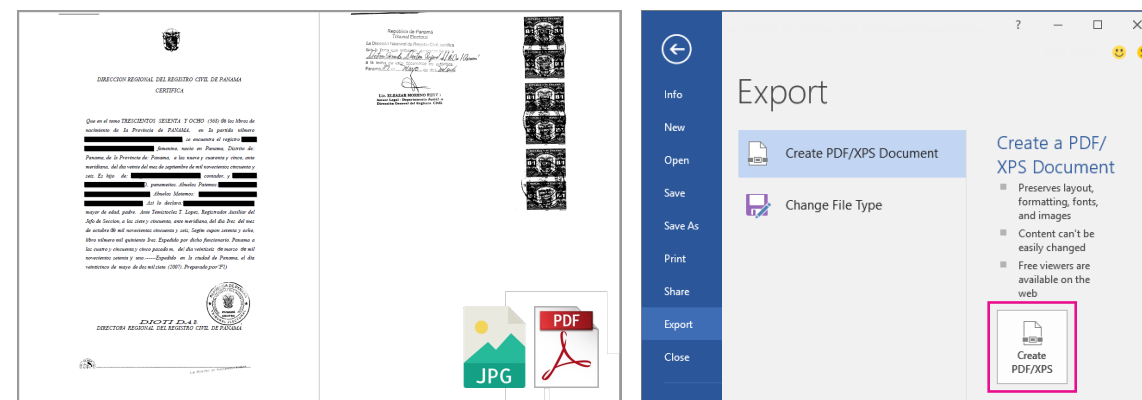
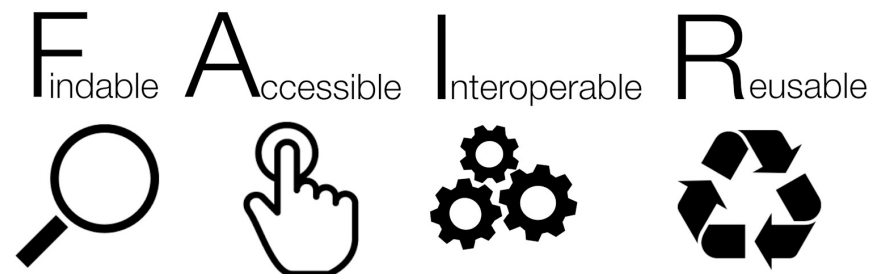
- LLM performance has power-law relationship with each of:
 - Compute time
 - **Training data (# tokens)**
 - LLM parameters
- Similar properties hold for other modalities
 - e.g., “Foundation models”
- Exploiting these properties requires that training data is machine readable
 - *Web text is easy, other sources are not*



Top: Chinchilla scaling laws¹ showing power-law relationships (lower test loss is better, note that plots are log-log); bottom: GPT-3 experiments² showing sudden drop in loss (i.e., performance jump) after sufficient training data.

Documents as scans and PDFs

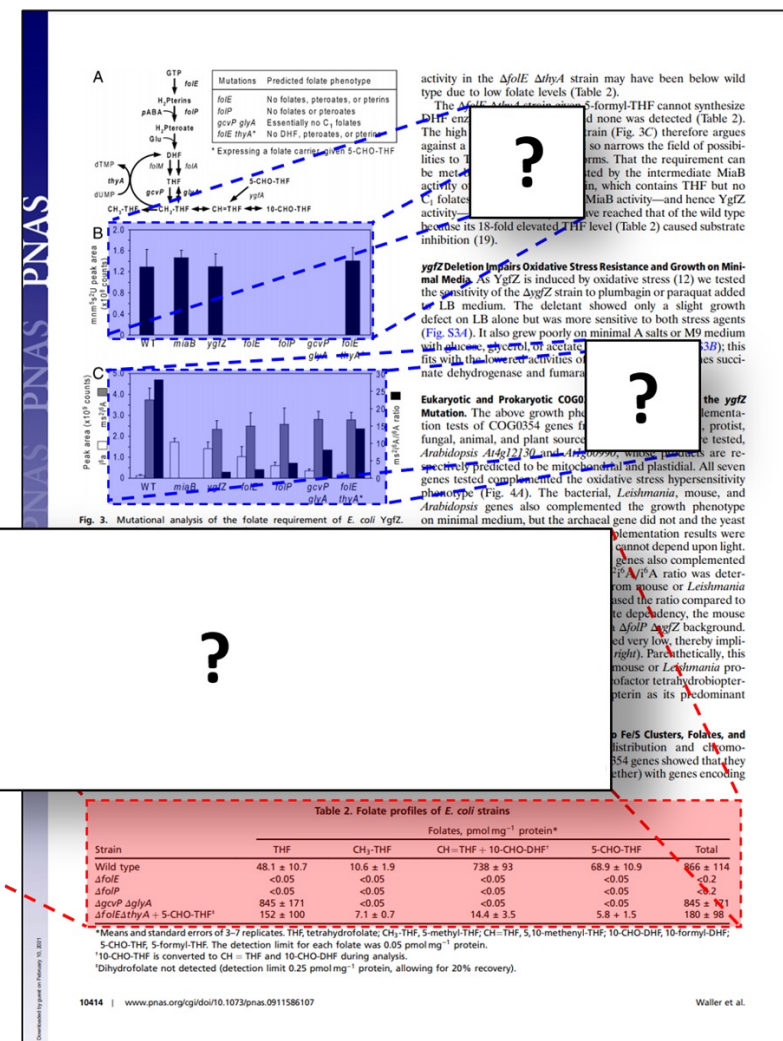
- Increasingly common to encourage or mandate open data & machine-readable documents
 - These standards are still not the norm
 - Huge historical archives*
- Scanned documents are not trivial to process
 - Naïve OCR enables text search, but not organization; **misses other contents**
- Even born-digital PDFs are not actually machine readable..



Top: FAIR data principles are gaining usage but remain far from common; bottom-left: scanned documents (as images or PDFs) are significant component of many archives; bottom-right: digitally sourced PDFs are not a silver bullet.

- Algorithmically extracting information from PDFs is *exactly as limiting as from OCR scans*

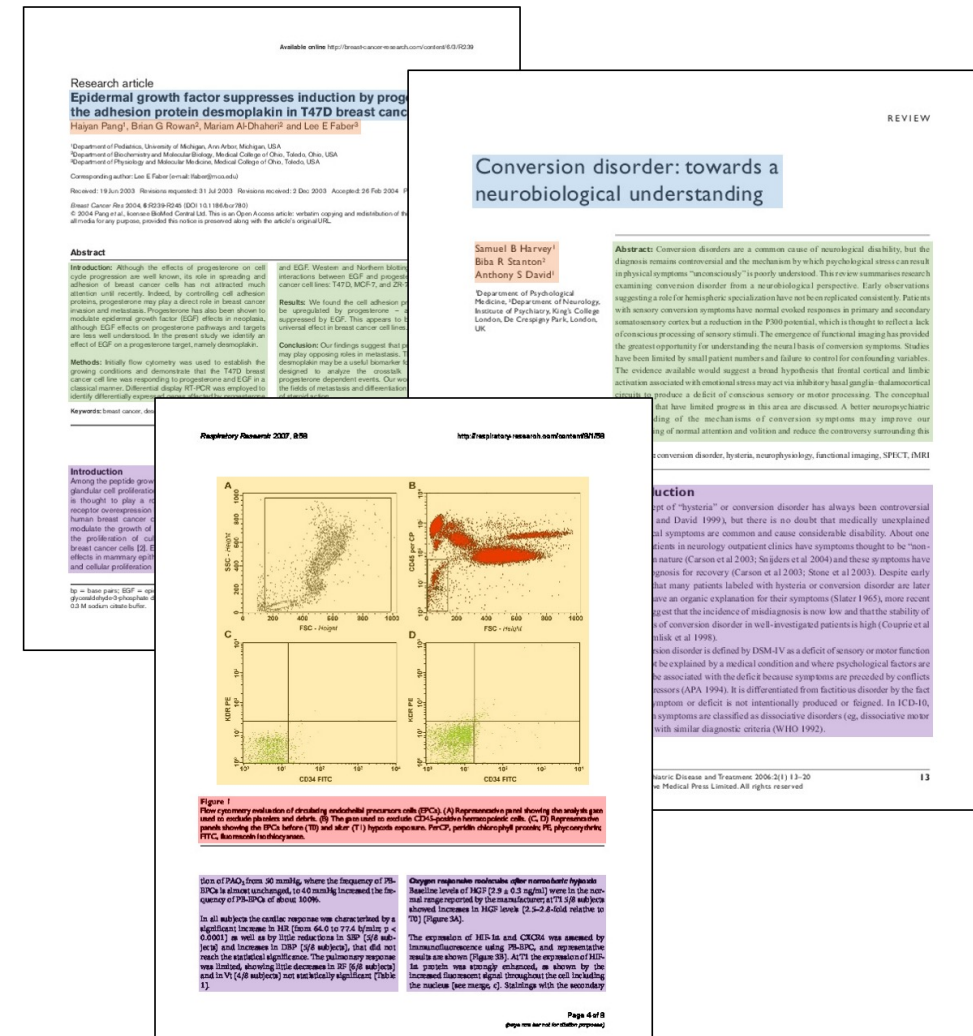
- **Extraction errors** impair scalable information retrieval & analysis



Document Layout Analysis

To **extract contents of different types** (text, tables, figures, captions, etc.) documents must be reliably and automatically segmented.

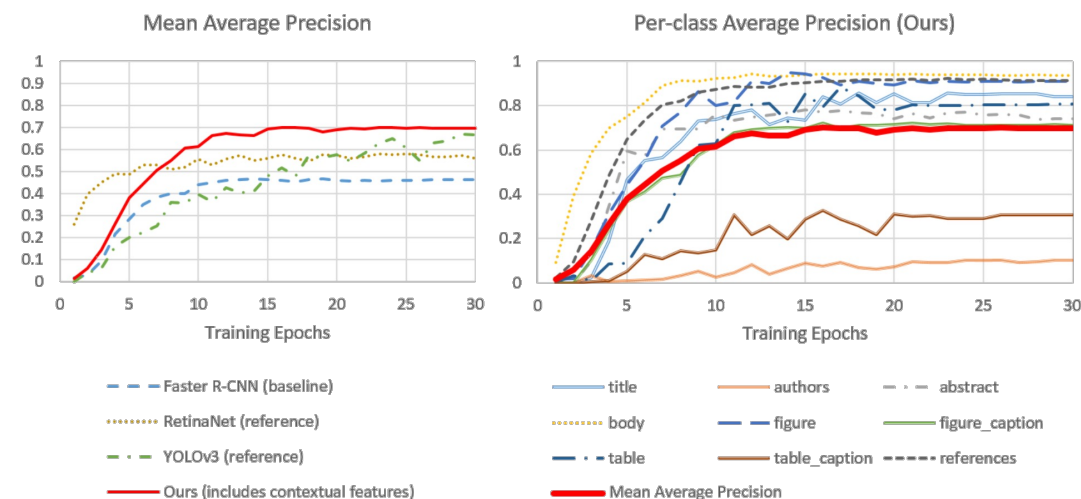
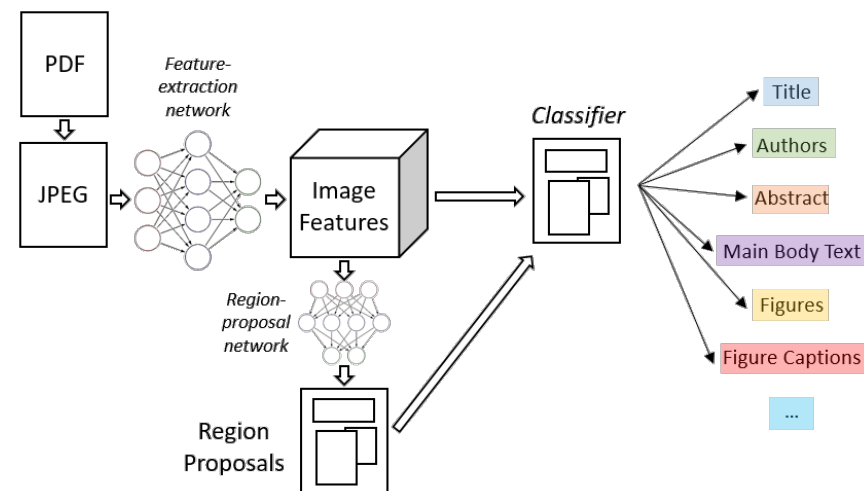
- Large portion of information and knowledge in documents takes **non-text form**
- Article PDF files have **highly varied underlying formatting** (even among digital productions)
- Most extraction tools rely on analyzing document *metadata* and markup; results are often very noisy



Document Layout Analysis

Our ML approach¹ is **purely vision-based**

- We adapt an object detection method to exploit the structured nature of scientific articles
- This approach enables processing **tables and figures**, which contain valuable information, yet are inaccessible to most methods
- Our method achieves **94%+ detection precision** on full-text articles, with **14x faster** document processing than prior methods



Automated Table Extraction

Tables are theoretically excellent data containers for machine processing, however **PDF-embedded tables are often very challenging to extract**.

- Standard extraction tools (e.g., *Tabula*) require manual table selection, produce many **extraction errors** (*sample with annotated errors at right*)
- Poor performance and manual intervention make such methods **unscalable**

We pursued two approaches to ML-based table extraction: mixed modality, and image-only

- Mixed modality approach used pipeline of adapted and custom models, used *image and text data*
- Achieved excellent table and cell detection (**98.8** and **96.0 F1**, respectively); but poor *structure prediction*

with neutron or proton numbers 64, 70, and 90 were predicted to be "doubly magic" with respect to tetrahedral symmetry. It was suggested [2] that nuclei displaying effects from this high-rank symmetry would produce a collective structure of negative-parity levels, but that the quadrupole moment, and thus the inband $B(E2)$ transition rate, would approach zero in the $U(5)$ limit. This limit would be most closely reached at the lowest spin of the band. Therefore, a set of rotational levels with "missing" $E2$ transitions at low spin may possibly be viewed as an indication for tetrahedral symmetry.

TABLE II: Branching and $B(E2)/B(E1)$ ratios for bands 2, 4, and 4a in ^{156}Dy . When possible, the branching ratios, λ , were determined from spectra that resulted from coincidence gates placed directly above the state of interest.

I^π	$E_\gamma(E2)$ (keV)	$E_\gamma(E1)$ (keV)	λ	$B(E2)/B(E1)$ ($\times 10^6 \text{ fm}^2$)
				Band 2
9 ⁻	376.3	970.5	0.020(4)	3.16(63)
11 ⁻	449.4	911.6	0.144(7)	7.75(38)
				Band 4
13 ⁻	517.9	868.5	0.43(3)	9.86(69)
15 ⁻	565.6	832.1	1.136(8)	17.7(10)
17 ⁻	611.3	807.9	2.8(2)	22.6(16)
				Band 4a
6 ⁻	271.1	1127.8	0.121(8)	154(10)
8 ⁻	363.1	1045.7	0.639(8)	151(2)
				Band 4a
12 ⁻	479.1	901.2	5.58(22)	211(8)
14 ⁻	491.5	790.7	20.5(10)	460(22)

Broken super/subscript Merged columns Wrong cell placement

Noisy extraction results with Tabula (even after fine manual alignment)

tetrahedral symmetry. The quadrupole moments of several of the negative-parity states of interest could be deduced. In addition, band mixing calculations allowing for an estimate of quadrupole moments were performed for the lowest negative-parity bands in the "tetrahedral doubly-magic" ^{154}Gd and ^{160}Yb [6] nuclei. In each of these cases, the quadrupole moments of the low-spin, negative-parity sequence were determined to be consistent with those measured in the ground-state bands of ^{154}Gd and ^{160}Yb . Differences in these ratios are a factor of 10-100 between the odd- and even-spin negative-parity bands. Therefore, the question has been raised as to whether these bands can be interpreted as "partner" bands with such large disparities in the $B(E2)/B(E1)$ ratios. However, as noted in Refs. [6, 21-23], these bands are likely octupole vibrations based on different K values, where the odd-spin sequence is assigned the $K = 0$ prin-

Automated Table Extraction

Image-only approach uses a Transformer-based object detection model (DETR²)

- Treats structure prediction task as a specialized image task
 - Previous approach was sequence to sequence text prediction (with LSTM or Transformer)
- Pretrained on ~1M machine-annotated tables
 - Fine-tuned on tables from ~100 LaTeX-PDF document pairs
- Improves structure prediction from **78.1 to 98.5** (*grid table similarity metric*)

group	ASD (n = 11)			TD (n = 8)			Mann-Whitney test	
	Q ₁	Median	Q ₃	Q ₁	Median	Q ₃	U	p
Android robot								
Feel enjoyable	4	5	8	2	3.5	4.75	16.50	0.02*
Feel embarrassed	3	5	7	2.5	5	7.5	42.00	0.90
Feel stressed	1	3	5	2.25	3.5	6	36.50	0.55
Feel bored	1	2	3	2	3	5.75	27.00	0.18
Visually simple robot								
Feel enjoyable	5	6	9	2	5	5.75	17.50	0.03*
Feel embarrassed	1	4	6	1	3	3	31.50	0.31
Feel stressed	1	3	5	1	2.5	3.75	39.00	0.72
Feel bored	1	2	4	1.25	3	3.75	35.00	0.18

Legend:
 Row (odd) Column (odd) Spanning cell Projected row header
 Row (even) Column (even) Column header

[illegible]

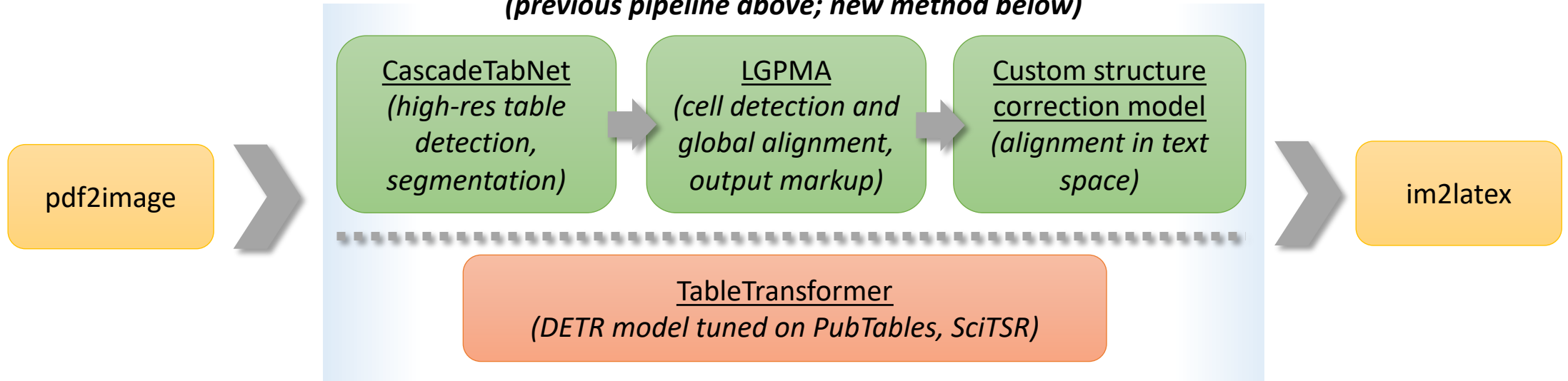
E_i	J^π	E_f	J_f^π	E_γ	BR_γ
0	$1/2^-$				
60.1	$7/2^+$				
129.6	$9/2^+$	60.1	$7/2^+$	68.8	
376.6	$3/2^+$	0	$1/2^-$	376.6	
545.4	$7/2^+$	376.6	$3/2^+$	168.9	23 (2)
		129.6	$7/2^+$	415.4	100
		60.1	$7/2^+$	486.0	16 (2)
607.0	$5/2^+$	376.6	$3/2^+$	230.3	
		60.1	$7/2^+$	547.0	
683.3	$9/2^+$	60.1	$7/2^+$	623.3	
704.9	$11/2^{(+)}$	129.6	$9/2^+$	575.2	
823.9	$13/2^+$	704.9	$11/2^{(+)}$	118.9	< 5
		129.6	$9/2^+$	694.3	100
876.9	$9/2^+$	60.1	$7/2^+$	816.8	
958.5	$11/2^+$	544.4	$7/2^+$	413.3	
		129.6	$9/2^+$	828.7	
1125.7	$11/2^+$	876.9	$9/2^+$	248.7	
		683.3	$9/2^+$	442.4	
1159.4	$(13/2^+)$	704.9	$11/2^{(+)}$	454.2	
		129.6	$9/2^+$	1030.0	
1388.3	$(13/2^+)$	1125.7	$11/2^+$	262.6	
1452.	$(15/2^+)$	683.3	$9/2^+$	768.8	
1474.0	$15/2^+$	1159.4	$(13/2^+)$	314.8	12 (2)
		823.9	$13/2^+$	650.0	100
		704.9	$11/2^{(+)}$	769.4	55 (4)
1573.6	$(15/2^+)$	958.5	$11/2^+$	615.1	
1664.5	$(17/2^+)$	823.9	$13/2^+$	840.9	
2130.6	$(15/2^+)$	823.9	$13/2^+$	1306.1	
2352.8	$(17/2^+)$	2130.6	$(15/2^+)$	222.6	
		1474.0	$(15/2^+)$	878.4	

Automated Table Extraction

PDF pre-processing

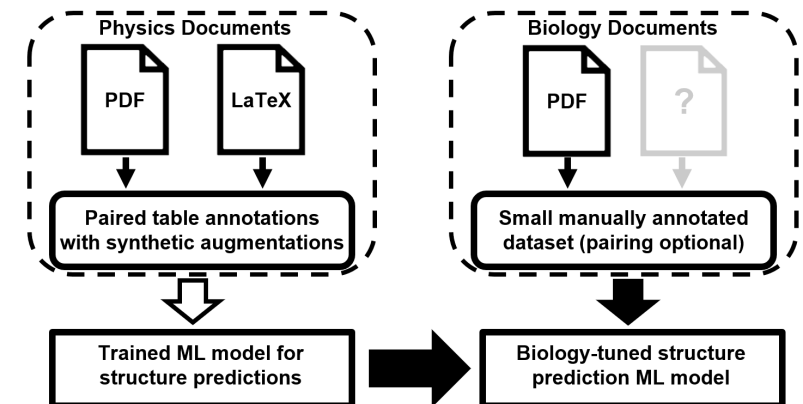
ML-based table detection, segmentation, and structure parsing
(previous pipeline above; new method below)

OCR, post-processing



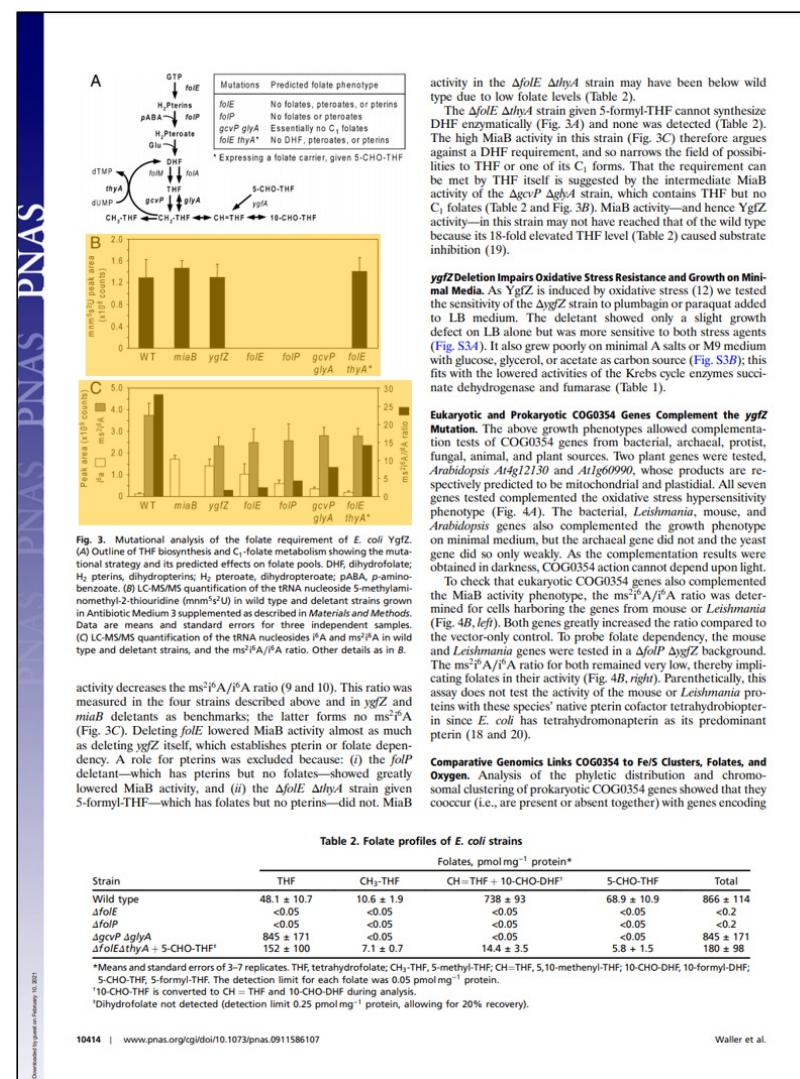
Structure prediction is a key challenge, particularly for domains without structure-resolving annotation pairs (e.g., LaTeX source)

- Addressable to some degree with domain adaptation



Reverse-engineering Data Charts

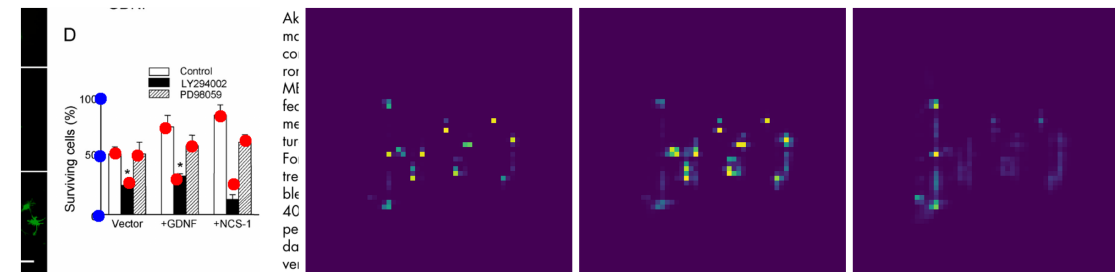
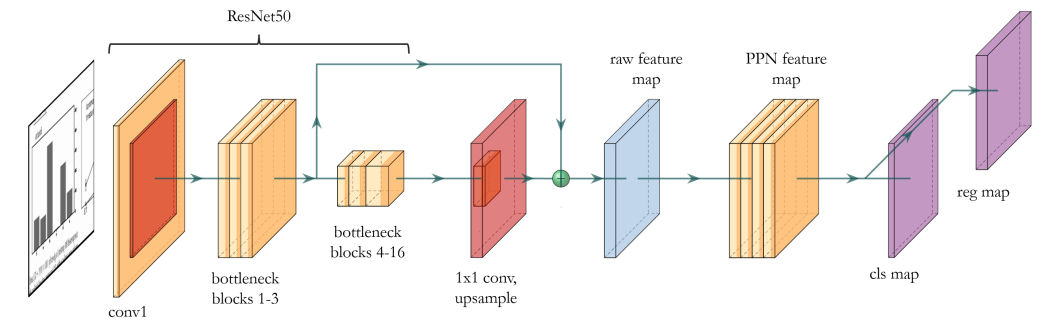
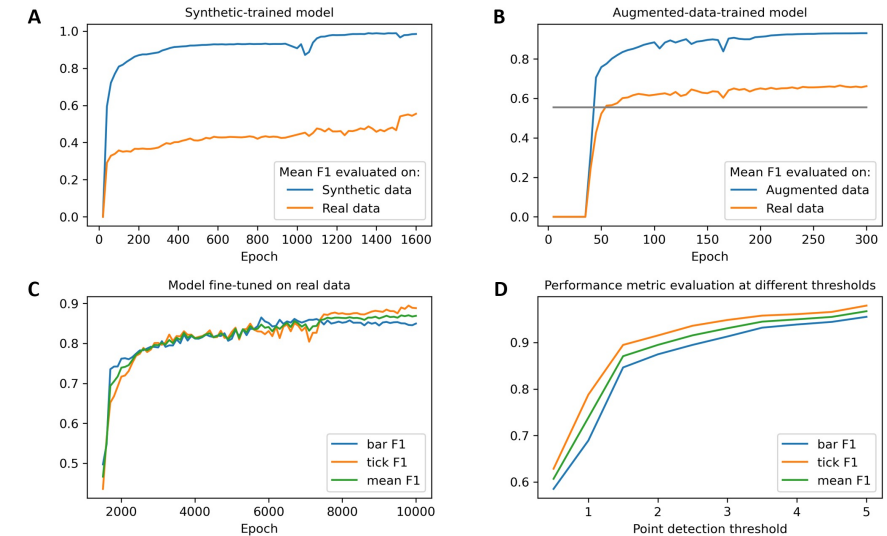
- Automatically extracting data chart contents could **enable new large-scale analysis** and knowledge discovery tools
 - Charts/plots are currently **inaccessible to automated document processing**
- Identifying charts in articles is straightforward
 - Can use whole document layout analysis tool, or simple object detection
 - >95% mAP** with standard object detection models (e.g., YOLOv3³)
- Imperfect alignment, need to oversample; so **value extraction must be robust to artifacts**



Extracting Values in Charts

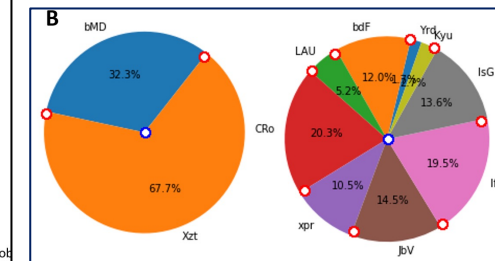
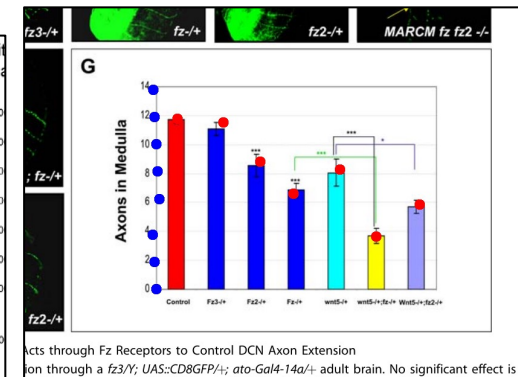
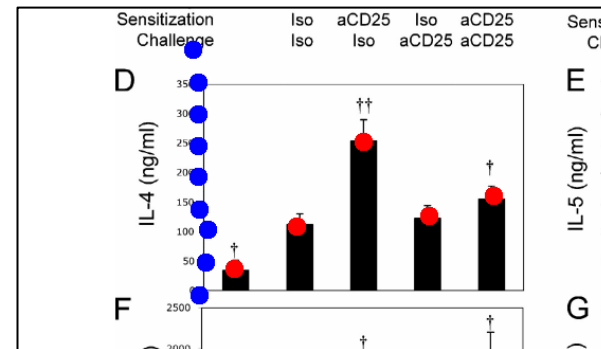
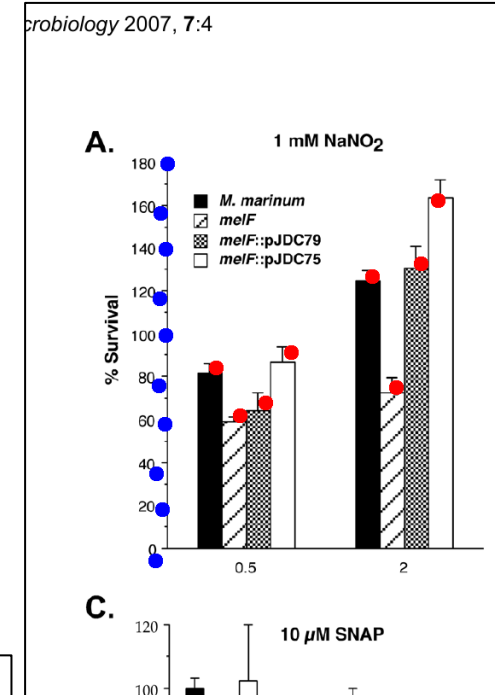
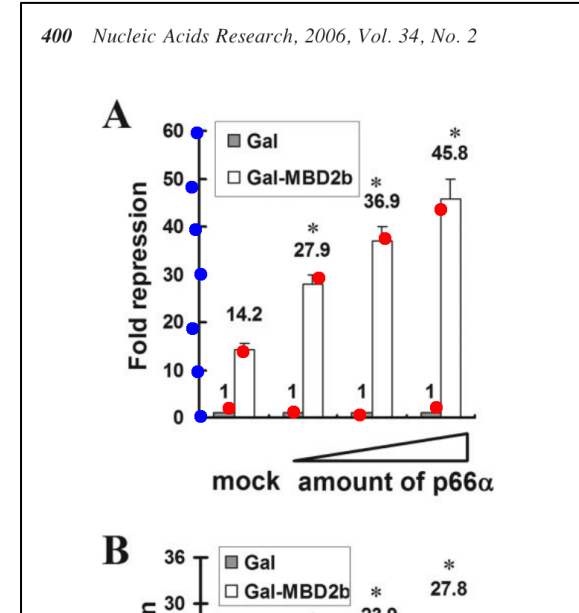
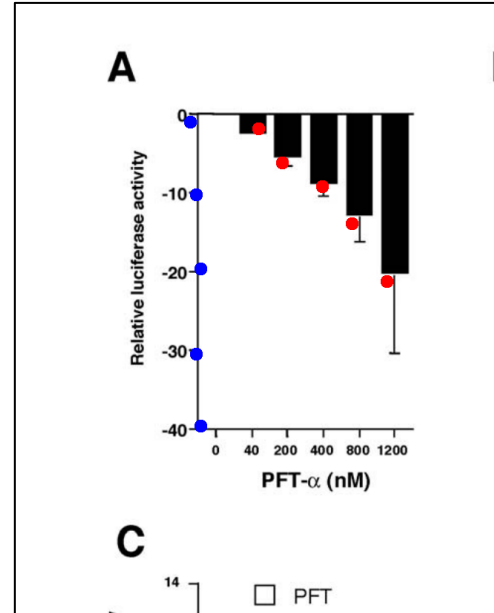
Point Proposal Network (PPN): an extensible model for robust data plot value extraction

- Directly predicts location and type of salient data points
- **Extensible to new chart types** (e.g., pie charts, line graphs, etc.) and chart features (e.g., error bars)
- Synthetic data generation enables efficient model pre-training
- **87.05 F1** on *complex* scientific charts



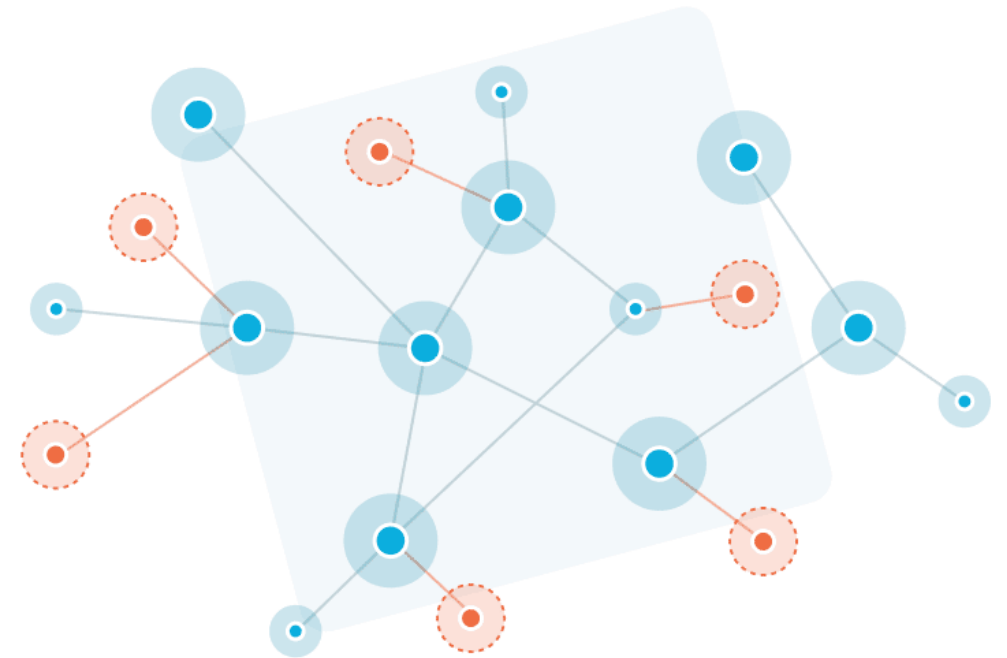
PPN Results and ongoing work

- Working on extensions to new chart elements and chart types
 - e.g., error bars, line graphs, stacked charts
- Building end-to-end tool
 - Integrating with OCR models (with minor modifications)
 - Label alignment and output generation using positional attention



Vision

- Scalable information retrieval is only beginning
- Build knowledge graphs at scale
 - Explore and interact with NLP
 - Transform and visualize previously disparate data
 - Include programmatic interfaces (APIs) for accessing curated databases, and to feed back into them



csoto@bnl.gov