

~/ccapp-talk %

A Prompt Is Not All You Need

FROM PROMPTS TO PIPELINES WITH NEMO DATA DESIGNER

Johnny Greco | CCAPP 20th
Anniversary Conference | June 2026

~/ccapp-talk %

What I've Been Up To

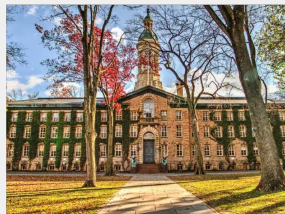
FROM PROMPTS TO PIPELINES WITH NEMO DATA DESIGNER

Johnny Greco | CCAPP 20th
Anniversary Conference | June 2026

CCAPP, IT'S BEEN AWHILE



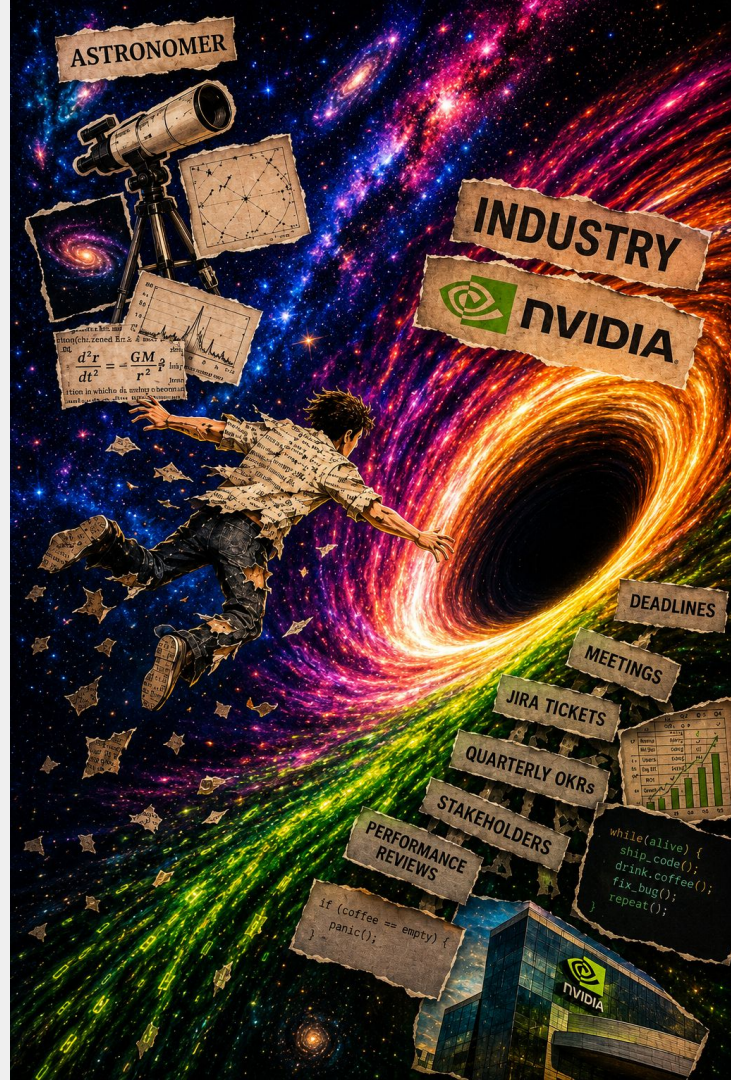
Undergrad
Ohio State
2009 - 2013



Grad school
Princeton
2013 - 2018



Postdoc
CCAPP
2018 - 2021



CCAPP, IT'S BEEN AWHILE

MEASUREMENT OF THE MASS AND STELLAR POPULATION DISTRIBUTION IN M82 WITH THE LBT

JOHNNY P. GRECO, PAUL MARTINI, AND TODD A. THOMPSON

Department of Astronomy, The Ohio State University, Columbus, OH 43210, USA; greco.40@buckeyemail.osu.edu

Received 2012 January 25; accepted 2012 July 24; published 2012 August 31

ABSTRACT

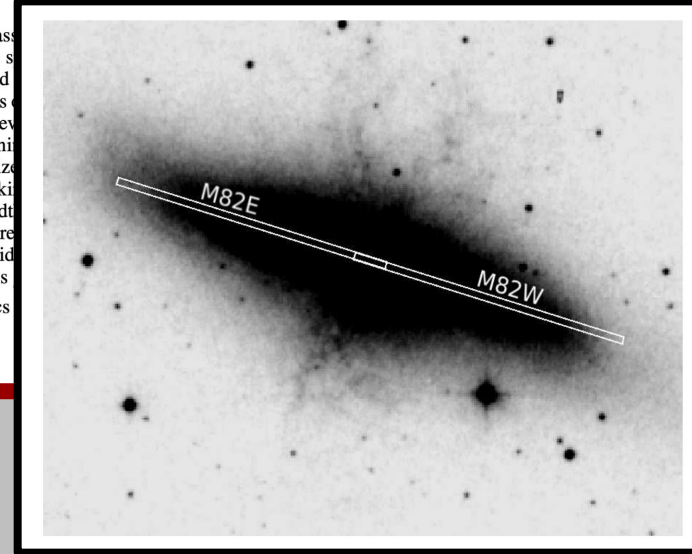
We present a *K*-band spectroscopic study of the stellar and gas kinematics, mass of the archetypical starburst galaxy M82. Our results are based on a single slit through the *K*-band nucleus. We used the ^{12}CO stellar absorption band head rotation curve out to nearly 4 kpc radius on both the eastern and western sides. The rotation curve is flat from 1 to 4 kpc. This stands in sharp contrast to some previous H I and CO emission-line position–velocity diagrams as evidence for a declining rotation curve. The Br γ , H $_2$, and He I emission lines are consistent with, although characterize the stellar kinematics. We derived M82's mass distribution from our stellar kinematics. That its total dynamical mass is $\sim 10^{10} M_{\odot}$. We measured the equivalent width variation in $W_{2.29}$ with radius clearly shows that RSGs dominate the light inside, likely launched from this region, where we estimate that the enclosed mass is

Key words: galaxies: individual (M82) – galaxies: kinematics and dynamics – galaxies

Online-only material: color figures



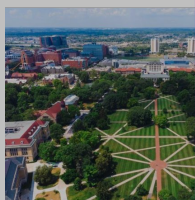
Undergrad
Ohio State
2009 - 2013



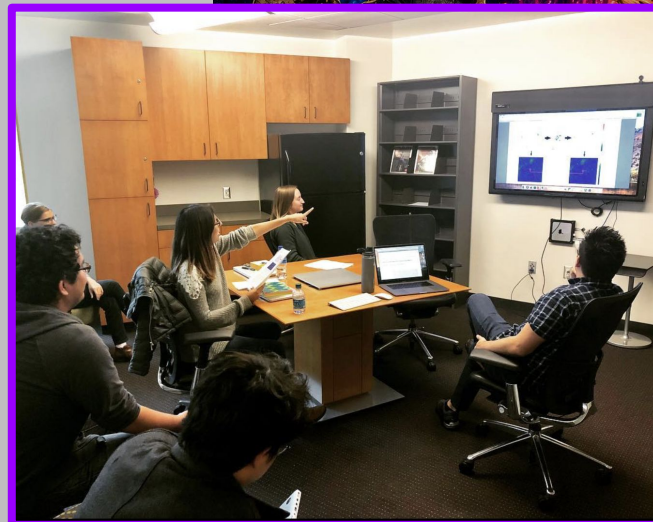
CCAPP IT'S BEEN AWHILE



Postdoc
CCAPP
2018 - 2021



Undergrad
Ohio State
2009 - 2013



CCAPP, IT'S BEEN AWHILE



Undergrad
Ohio State
2013



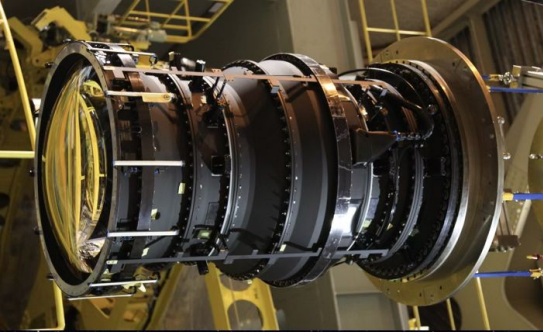
Grad school
Princeton
2018



Postdoc
CCAPP
2018 - 2021

GPT 3 released

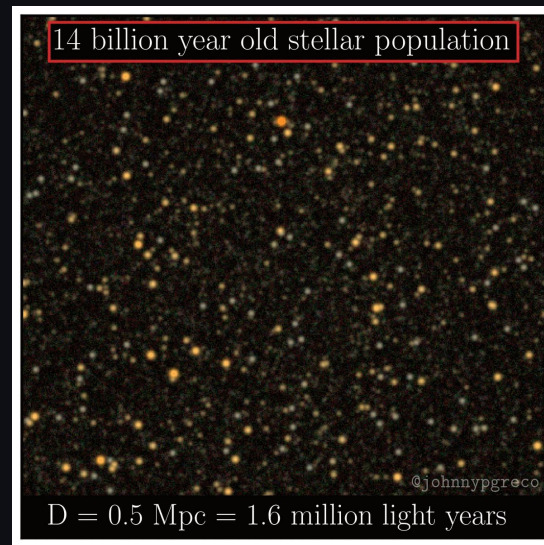
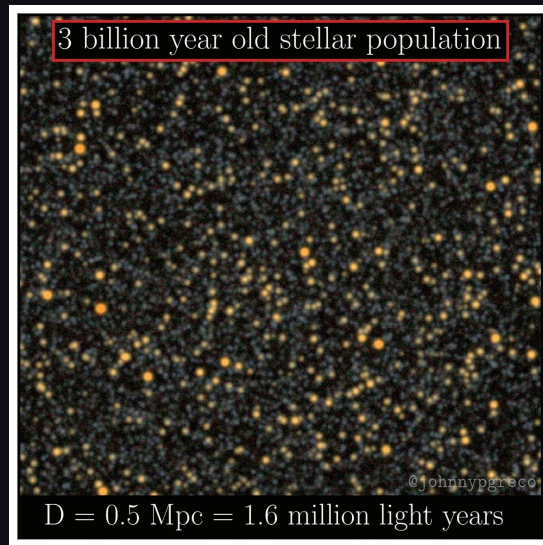
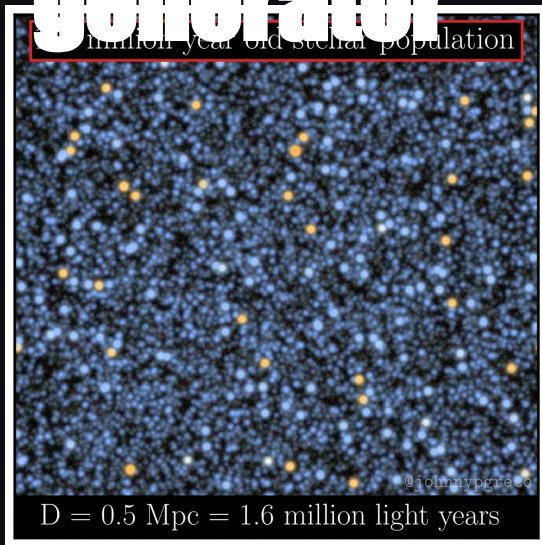




For Old Times' Sake:
Some pretty LSB galaxy gifs

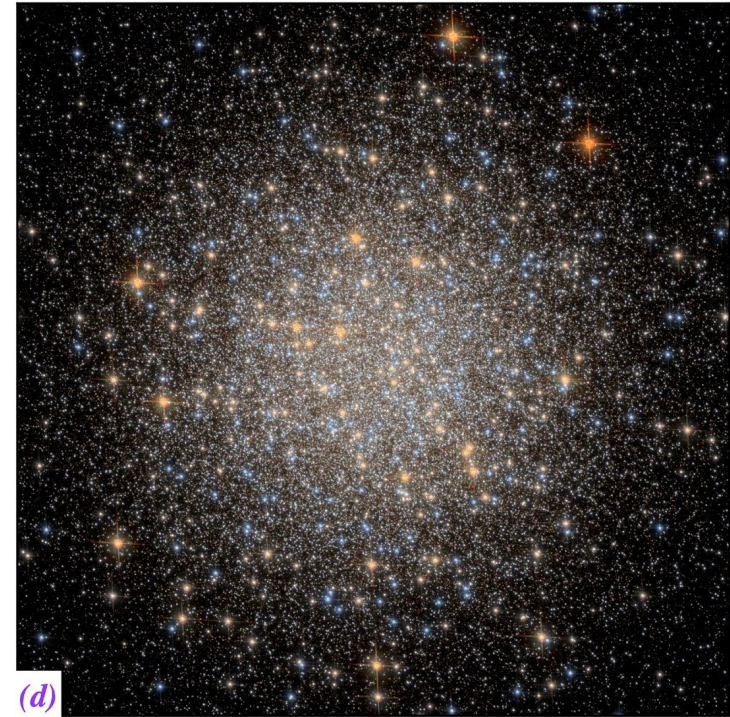
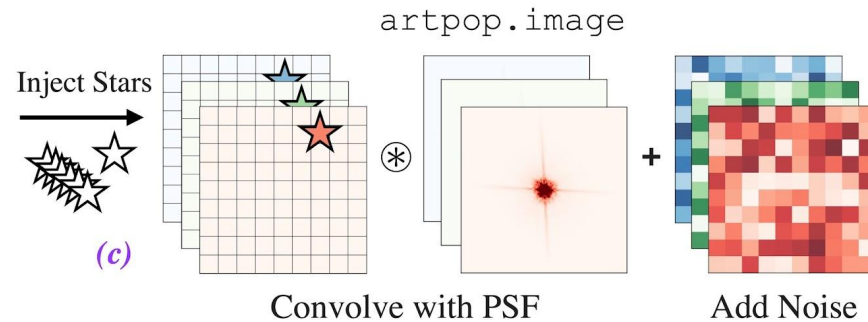
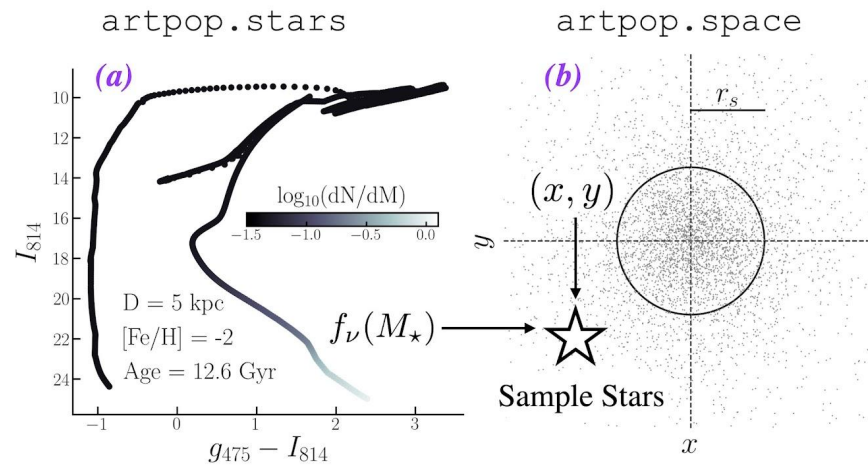


ArtPop – synthetic galaxy generator



```
$ pip install artpop
```


ArtPop Globular Cluster Simulation



(R, G, B) = (F814W, F606W, F475W)

Synthetic Data Goes Mainstream

Synthetic Data

What is synthetic data? Why do we need it?

What is synthetic data?

Synthetic data is data generated or transformed by a process / system rather than collected directly from users in the world / the internet.

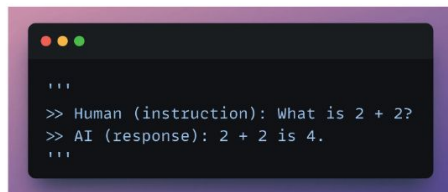
Examples :

- **Simulated scenes** for robotics, driving, or vision
- **LLM-generated text** for instruction tuning



Why synthetic data?

- 🛒 **Cost** — expensive to collect and label
- 🎯 **Coverage** — rare and edge cases are missing
- 👁️ **Bias** — real data reflects what is easy to observe
- 🛡️ **Privacy + Safety** — privacy, safety, copyright
- 🛠️ **Control** — create precise data for your task



Why search (& process) for the perfect dataset when you can design it?

Synthetic Data Goes Mainstream

ChatGPT Released

RL from AI Feedback

Self-Instruct

Textbooks are
all you Need

2022

2023

Anthropic



Computer Science > Computation and Language
(Submitted on 15 Dec 2022)

Constitutional AI: Harmlessness from AI Feedback

Yuntao Bai, Saurav Kadavath, Sandipan Kundu, Amanda Askell, Jackson Kernion, Andy Jones, Anna Chen, Anna Goldie, Azalia Mirhoseini, Cameron McKinnon, Carol Chen, Catherine Olsson, Christopher Olah, Danny Hernandez, Dawn Drain, Deep Ganguli, Dustin Li, Eli Tran-Johnson, Ethan Perez, Jamie Kerr, Jared Mueller, Jeffrey Ladish, Joshua Landau, Kamal Houdou, Kamille Lukosuite, Liane Lovitt, Michael Sellitto, Nelson Elhage, Nicholas Schiefer, Noemi Mercado, Nova DasSarma, Robert Lasenby, Robin Larson, Sam Ringer, Scott Johnston, Shauna Kravec, Sheer El Showk, Stanislav Fort, Tamera Lanham, Timothy Telleen-Lawton, Tom Conerly, Tom Henighan, Tom Zoph, Samuel R. Bowman, Zac Hatfield-Dodds, Ben Mann, Dario Amodei, Nicholas Joseph, Jared Kaplan

As AI systems become more capable, the need for training a harmless AI becomes more pressing. The process involves both a pre-training phase, where the model is trained on a large dataset of text, and a fine-tuning phase, where the model is trained on a dataset of human-written instructions. In the RL phase, we sample from the model's output and use the results to train a preference model from human feedback. We then use this preference model to train a Self-Instruct model, which generates its own instructions. This process is repeated until the model's output is aligned with human preferences.

Computer Science > Computation and Language
(Submitted on 20 Dec 2022 (v1), last revised 25 May 2023 (this version, v2))

Self-Instruct: Aligning Language Models with Self-Generated Instructions

Yizhong Wang, Yeganeh Kordi, Swaroop Mishra, Alisa Liu, Noah A. Smith, Daniel Khashabi, Hannaneh Hajishirzi

Large "instruction-tuned" language models (i.e., finetuned to respond to instructions) have demonstrated a remarkable ability to generalize zero-shot to new tasks. Nevertheless, they depend heavily on human-written instruction data that is often limited in quantity, diversity, and creativity, therefore hindering the generalization of the tuned models. We introduce Self-Instruct, a framework for improving the instruction-following capabilities of pretrained language models by bootstrapping off their own generations. Our pipeline generates instructions, input, and output samples from a language model, then filters invalid or similar ones before using them to finetune the original model on Super-NaturalInstructions, on par with the performance of InstructGPT-001, which was trained with private user data and human annotations. For further evaluation, we curate a set of expert-written instructions for novel tasks, and show through human evaluation that tuning GPT3 with Self-Instruct demonstrates a 33% absolute improvement over the original model on Super-NaturalInstructions. We also release a synthetic dataset of expert-written instructions for novel tasks, and show through human evaluation that tuning GPT3 with Self-Instruct demonstrates a 33% absolute improvement over the original model on Super-NaturalInstructions. We also release a synthetic dataset of expert-written instructions for novel tasks, and show through human evaluation that tuning GPT3 with Self-Instruct demonstrates a 33% absolute improvement over the original model on Super-NaturalInstructions.

at [this https URL](https://arxiv.org/abs/2306.11644).

Computer Science > Computation and Language
(Submitted on 20 Jun 2023 (v1), last revised 2 Oct 2023 (this version, v2))

Textbooks Are All You Need

Suriya Gunasekar, Yi Zhang, Jyoti Aneja, Caio César Teodoro Mendes, Allie Del Giorno, Sivakanth Gopi, Mojan Javaheripi, Piero Kauffmann, Gustavo de Rosa, Olli Saarikivi, Adil Salim, Shital Shah, Harkirat Singh Behl, Xin Wang, Sébastien Bubeck, Ronen Eldan, Adam Tauman Kalai, Yin Tat Lee, Yuanzhi Li

We introduce phi-1, a new large language model for code, with significantly smaller size than competing models: phi-1 is a Transformer-based model with 1.38 parameters, trained for 4 days on 8 A100s, using a selection of "textbook quality" data from the web (68 tokens) and synthetically generated textbooks and exercises with GPT-3.5 (18 tokens). Despite this small scale, phi-1 attains pass@1 accuracy 50.6% on HumanEval and 55.5% on MBPP. It also displays surprising emergent properties compared to phi-1-base, our model before our finetuning stage on a dataset of coding exercises, and phi-1-small, a smaller model with 350M parameters trained with the same pipeline as phi-1 that still achieves 45% on HumanEval.

Comments: 26 pages; changed color scheme of plot, fixed minor typos and added couple clarifications
Subjects: [arXiv:2306.11644](https://arxiv.org/abs/2306.11644) [cs.CL], Artificial Intelligence (cs.AI), Machine Learning (cs.LG)
Cite as: [arXiv:2306.11644v2](https://arxiv.org/abs/2306.11644) [cs.CL] for this version
(or [arXiv:2306.11644](https://arxiv.org/abs/2306.11644))
<https://doi.org/10.48550/arXiv.2306.11644>

Synthetic Data Goes Mainstream

gretel

MOSTLY AI

hazy

TONIC
THE FAKE DATA COMPANY

datacebo

 datagen

Synthetic Data Goes Mainstream



→ Applied Scientist from 2023 - 2025

Working for a Startup

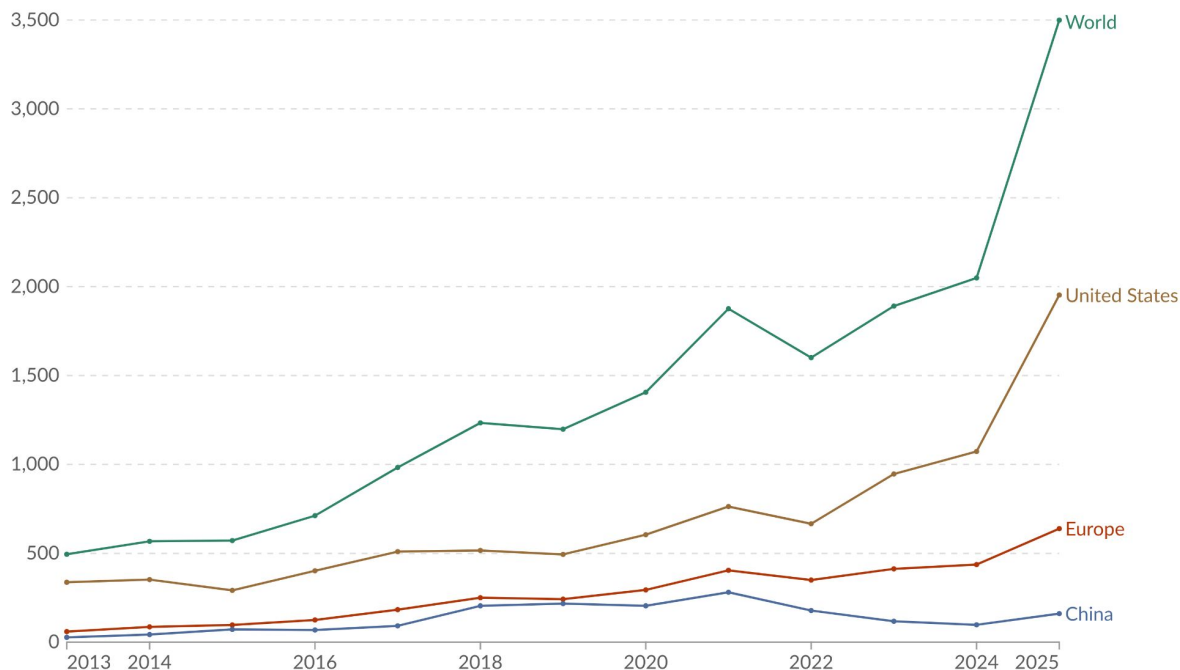


Working for a Startup

Newly-funded artificial intelligence companies

Our World
in Data

Newly-funded AI companies in each year that received an investment of more than \$1.5 million (not adjusted for inflation).

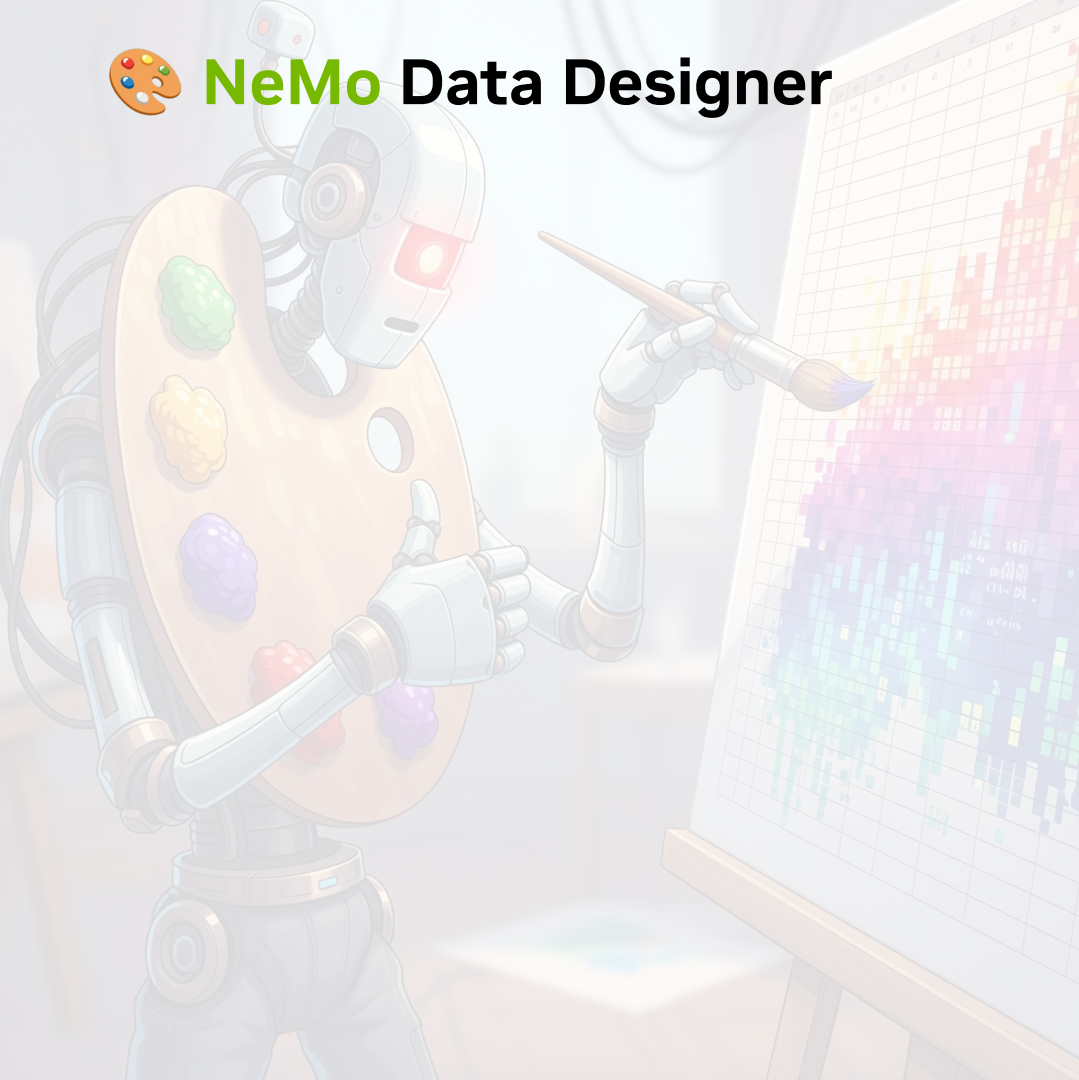


Data source: Quid via AI Index Report (2026)

OurWorldinData.org/artificial-intelligence | CC BY



NeMo Data Designer



What is it?

A **general framework** for generating high-quality **synthetic data from scratch** or based on **your own seed data**.

★ General framework

- Use-case agnostic
- Modular and extensible

★ From scratch

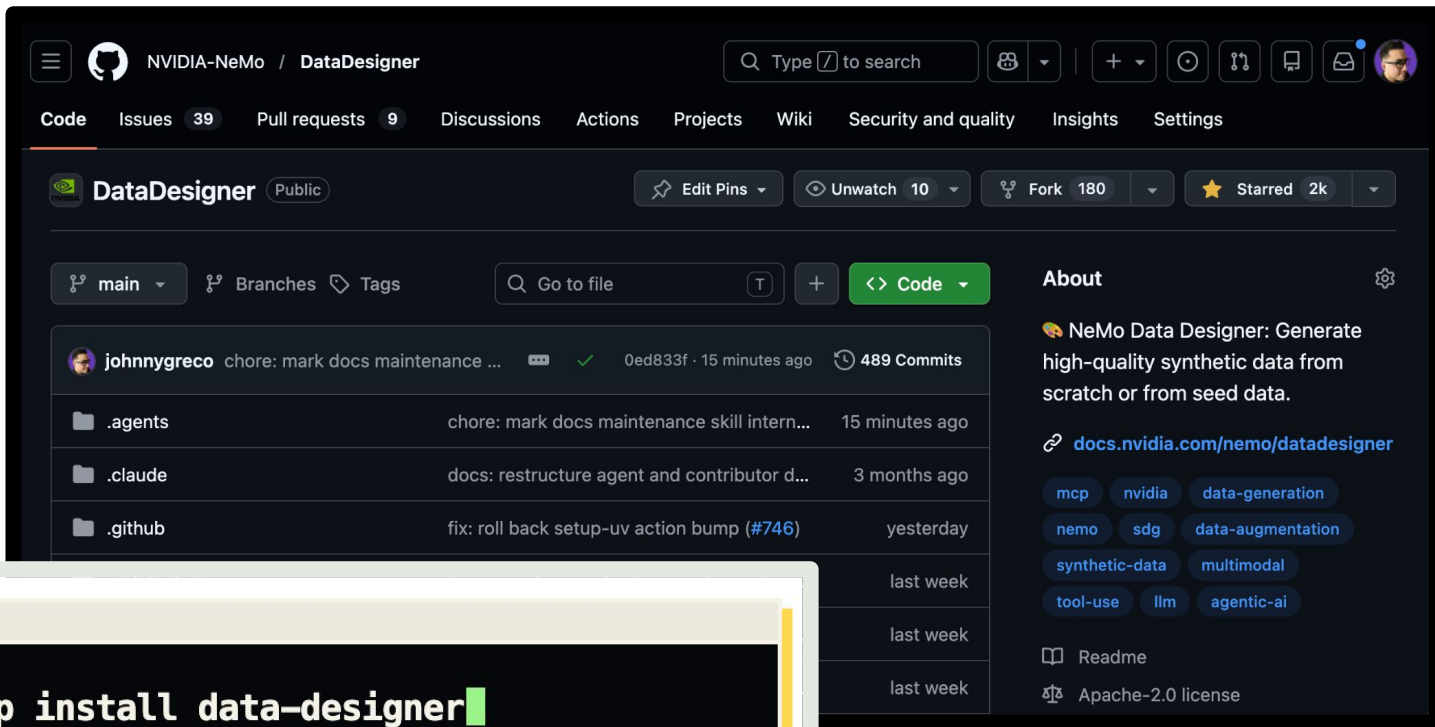
- Declarative config → *high-quality* dataset

★ Your own seed data

- Domain-grounded data generation

NeMo Data Designer: *What is it?*

Open Source Library: <https://github.com/NVIDIA-NeMo/DataDesigner>

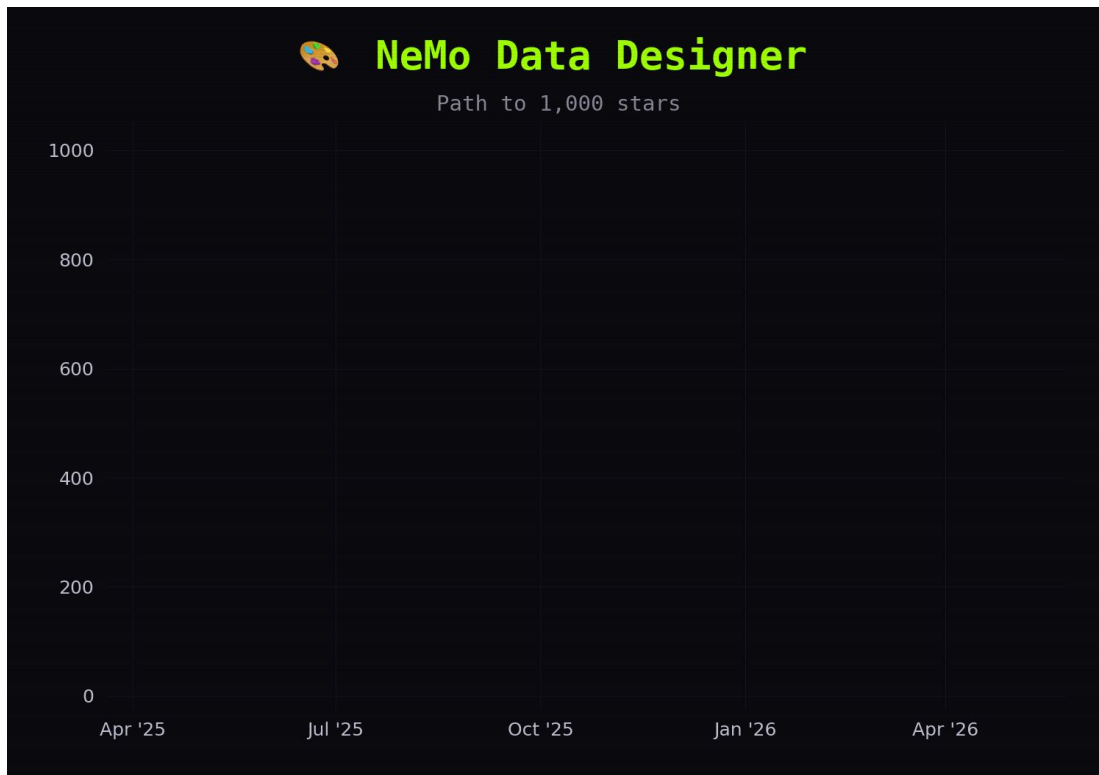


The screenshot displays the GitHub repository for NVIDIA-NeMo DataDesigner. The repository is public and has 180 forks and 2k stars. The main content area shows a list of recent commits by user johnnygreco, including updates to the .agents, .claude, and .github folders. The right sidebar provides an overview of the project, stating it is a tool to generate high-quality synthetic data from scratch or from seed data. It includes tags for various features like mcp, nvidia, data-generation, nemo, sdg, data-augmentation, synthetic-data, multimodal, tool-use, llm, and agentic-ai. Links to the README and the Apache-2.0 license are also present.

```
$ pip install data-designer
```

Greco et al. in prep.

NeMo Data Designer: *What is it?*



2.5+ trillion tokens processed and *accelerating!*



NeMo Data Designer: *Why do we need it?*



```
1 records = []
2
3 for _ in range(1_000_000):
4     records.extend(
5         [
6             {"customer": llm.generate("Generate a customer profile."),
7             {"product": llm.generate("Generate a product that the customer would buy."),
8             {"review": llm.generate("Generate a review for the product."),
9         ]
10    )
11
12 synthetic_dataset = pd.DataFrame.from_records(records)
```

Is this Synthetic Data Generation?



NeMo Data Designer: *Why do we need it?*

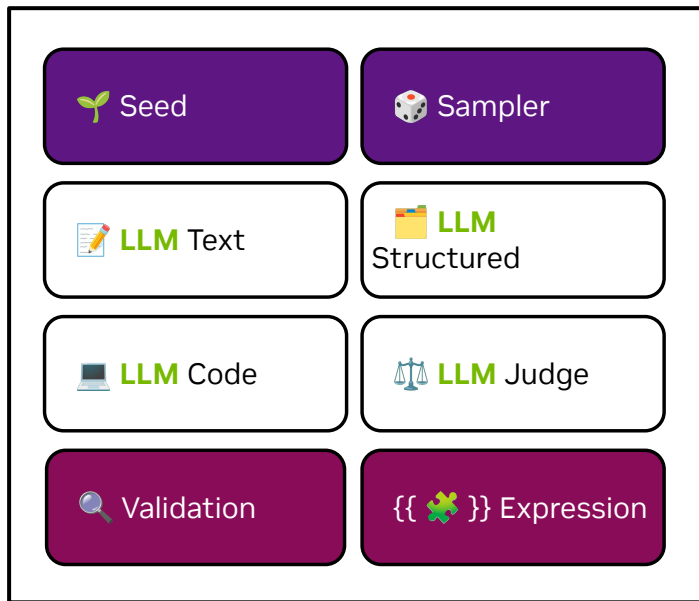
What's missing?

- Diversity, diversity, diversity!
- Correlations
- Multimodality
- Steerability
- Validation / verification
- Reproducibility
- **Diversity!**

```
1 records = []
2
3 for _ in range(1_000_000):
4     records.extend(
5         [
6             {"customer": llm.generate("Generate a customer profile.")},
7             {"product": llm.generate("Generate a product that the customer would buy.")},
8             {"review": llm.generate("Generate a review for the product.")},
9         ]
10    )
11
12 synthetic_dataset = pd.DataFrame.from_records(records)
```

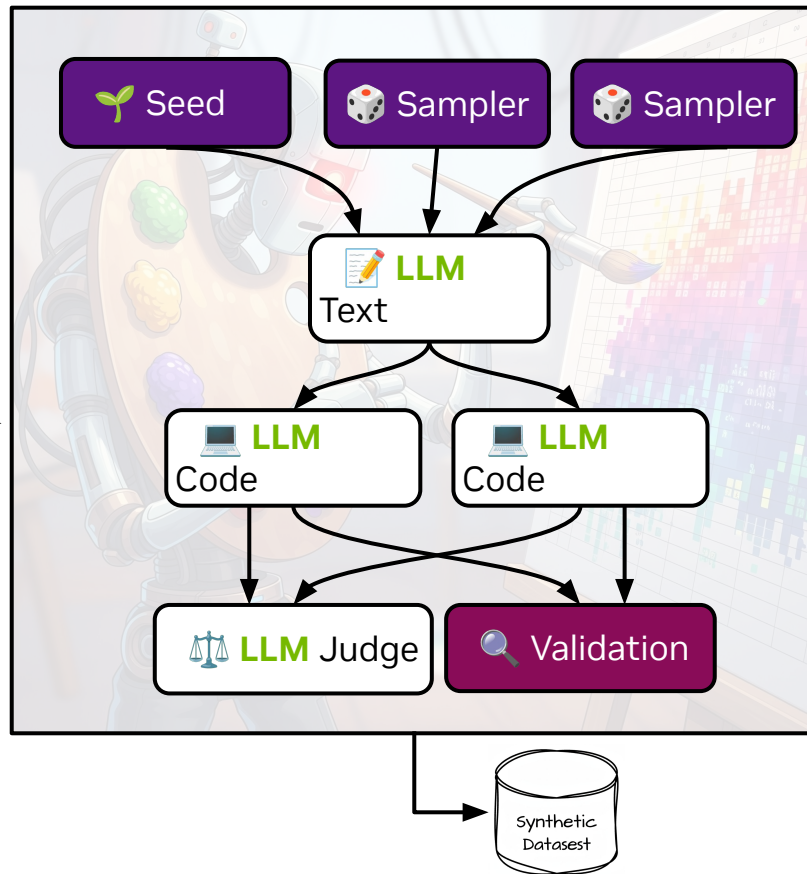
NeMo Data Designer: *How does it work?*

Column Types



Create Dataset

Dataset Builder workflow



NeMo Data Designer: Agent **skill** demo

Dev Notes

Dev Notes

🔮 Ask a question | 📄 Copy page | 📖 View as Markdown | ⋮ More actions

Welcome to NeMo Data Designer Dev Notes — in-depth guides, benchmark write-ups, and insights from the team building NeMo Data Designer.

Nemotron-Personas



JUN 1, 2026

Designing Nemotron-Personas

The compound-AI pipeline behind Nemotron-Personas seeding Nemotron training. Now open-sourced: fork it, add your data, ship personas for a...

Yev Meyer +1

Prompt Sensitivity

Small changes in prompt.
Same intent. Same answer.

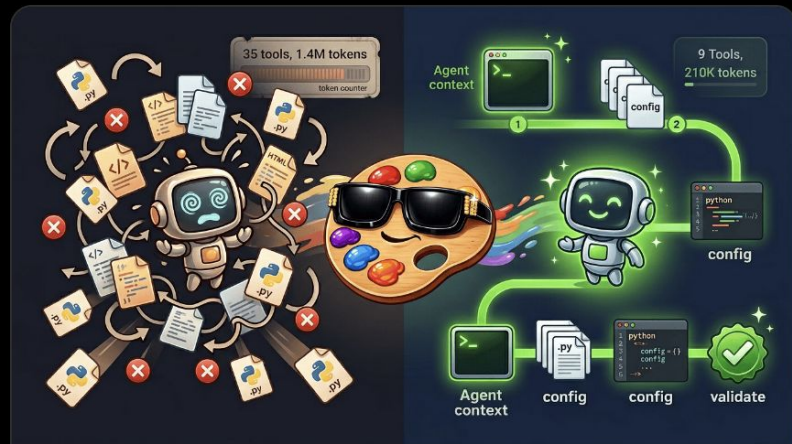


MAY 28, 2026

Mitigating Prompt Sensitivity

A Data Designer pipeline that generates thousands of regex-paired prompt preambles to manufacture format robustness — six diversity samplers, four...

Dhruv Nathawani

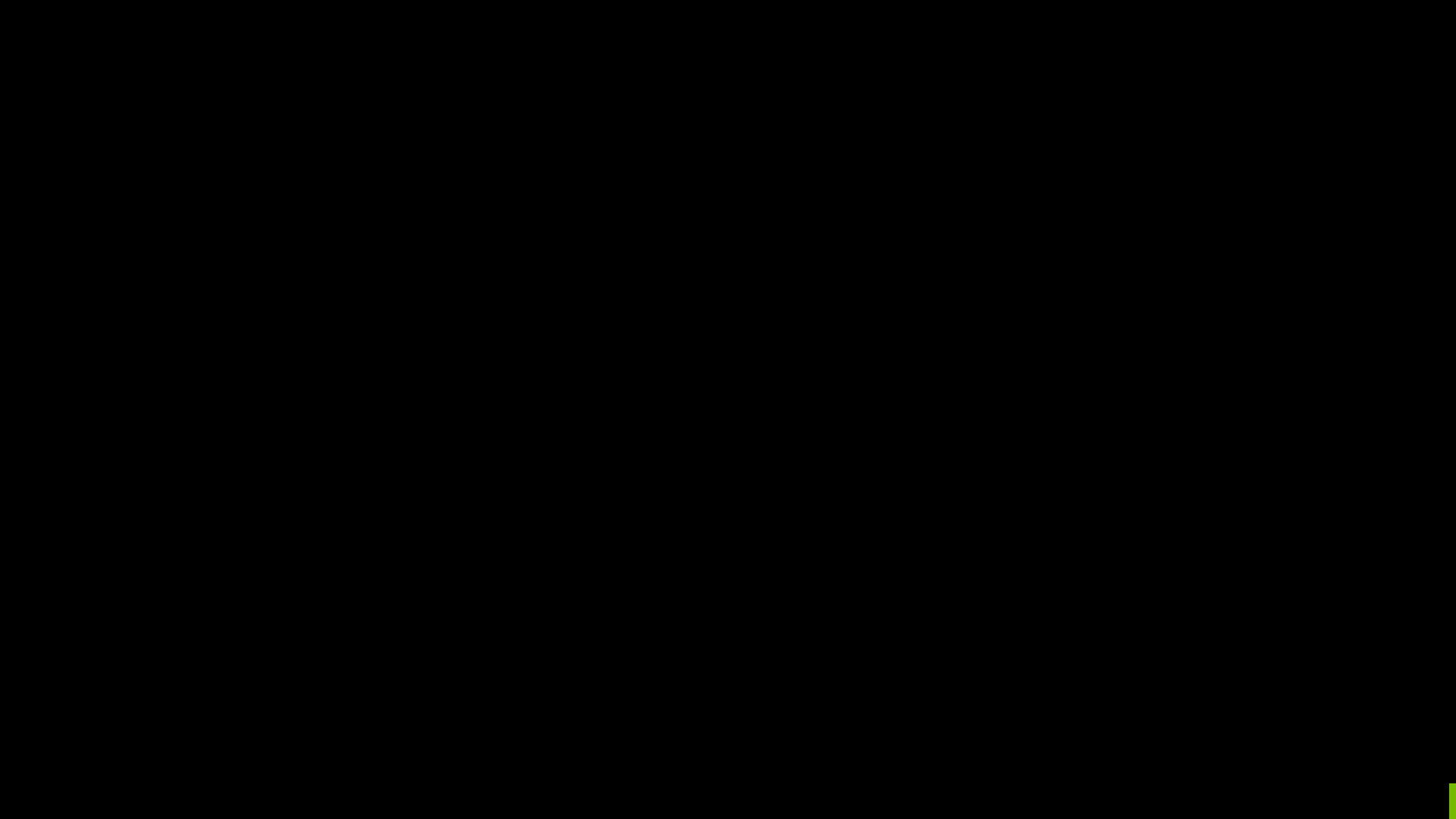


MAR 24, 2026

Data Designer Got Skills

A CLI and skill workflow that lets agents drive Data Designer end-to-end — leaner context, fewer tool calls, the same output.

Johnny Greco




```
* Build the project

the two accuracy=2 records.

• Bash(.venv/bin/python -c "
    import pandas as pd, textwrap, json...
    =====
    [0] types of galaxies | hobbyist amateur astronomer
    Q: What exactly are ultra-diffuse galaxies and why do they appear so large and faint compared to
    ... +74 lines (ctrl+o to expand)

The judge's reasoning is spot-on and well-calibrated – these aren't over-corrections:
- [0] correctly catches that the answer overstates UDG sizes as "comparable to the Milky Way" (most are dwarf-scale).
- [1] correctly flags genuinely wrong numbers in the answer ("tens of metres" tides, "several ms/year") and unsupported speculation.

So the harsher scores are real errors in the answers, surfaced by the judge – exactly the filtering signal you want for training data.

Summary of the review pass

All three issues are resolved:
```

Issue	Before	After
Mode collapse	Europa in 3/10, two duplicate black-hole questions	No repeated examples; questions span "what-if", observing, recent-mission (OSIRIS-REx), misconception angles
Audience mismatch	Beginner got grad-level transit Q; student got dubious infrared-bar Q	Beginners now ask simple real questions (ISS passes, "what if the Sun went supernova"); depth tracks audience
Lenient judge	29/30 scores were 4–5	Scores now span 2–5 and correctly catch real factual errors

```
Updated preview:

file:///Users/johnnygreco/demo-data-designer/artifacts/preview_results_20260609_232026/sample_records/sample_records_browser.html

The question_angle helper sampler does its job and is dropped from the final schema, so your output is the clean 6 columns: topic, subtopic, audience, question, answer, answer_quality.

This is in good shape. Want any further tweaks, or are you happy to finalize and scale it up?

* Sautéed for 3m 24s

> [ ] finalize it and give me the create command

  >> bypass permissions on (shift+tab to cycle) · - for agents
```



[index: 0]

Generated Columns

Name	Value
topic	Galaxies & Cosmology
subtopic	types of galaxies
audience	hobbyist amateur astronomer
question	What exactly are ultra-diffuse galaxies and why do they appear so large and faint compared to typical dwarf galaxies?
answer	**Ultra-diffuse galaxies (UDGs) are “milky-halo”-like systems that are surprisingly large—often comparable to the Milky Way in diameter—but so faint that they are hard to see even with modest telescopes.** Think of them as the cosmic equivalents of a huge, fluffy balloon that’s almost completely deflated: it stretches over many light-years, yet the amount of luminous material inside is tiny, so most of the “light” is just scattered wisps of stars spread thinly across the sky. Because their surface brightness is extremely low (often < 1 % of a typical dwarf galaxy’s), they only stand out when observed with sensitive equipment or in the data from large sky surveys. What makes them appear so big and faint? **Very low star-formation efficiency.** They started out as normal dwarf galaxies, but processes such as strong tidal stripping, early bursts of star formation that blew gas out, or a lack of fresh gas have left them with only a sparse sprinkling of stars. **Huge dark-matter halos.** Dynamical measurements show that many UDGs sit inside surprisingly massive dark-matter halos. The halo’s gravity can keep the galaxy stretched out over several kiloparsecs, inflating its size even though the visible stellar component is minuscule. **Low surface density.** Surface brightness depends on how many stars are packed into each square arcsecond. UDGs have far fewer stars per unit area than “regular” dwarfs, so their integrated light spreads out over a large area, making each pixel look faint. In short, UDGs look large and faint because they are **low-surface-brightness, star-poor galaxies** that reside in extended dark-matter envelopes. They’re not simply “failed” massive galaxies; rather, they’re **successful** in forming dark halos but less efficient at turning that material into bright stars, giving them the “big, ghostly” appearance that challenges our usual notions of galaxy size and luminosity.

LLM-as-a-Judge: answer_quality

accuracy	clarity	audience_fit
score: 2 reasoning: The answer correctly explains that ultra-diffuse galaxies are low surface-brightness, star-poor systems residing in extended dark-matter halos and notes mechanisms such as low star-formation efficiency and tidal effects. However, it overstates their size by claiming they are often comparable in diameter to the Milky Way, which is misleading for most UDGs (they are typically dwarf-scale, only a minority reach >5 kpc). Aside from this inaccuracy, the scientific mechanisms described are largely precise and up to date.	score: 4 reasoning: The response is well-structured, uses clear headings and bullet points, and employs accessible analogies (e.g., “fluffy balloon”) that help a hobbyist grasp the concepts. The explanation flows logically from definition to appearance to underlying causes, though a few phrasing choices (“milky-halo-like”) could be smoother. Overall it is easy to follow for the target audience.	score: 4 reasoning: The tone and vocabulary are appropriate for an amateur astronomer with a curiosity for deeper concepts: they avoid excessive technical jargon while still providing meaningful scientific context. The depth is balanced—neither oversimplified nor overly technical—making it a strong match for the stated audience.



```
~/ccapp-talk % exit
```

thanks, everyone!

NeMo Data Designer - Moving the Needle

Data Designer was critical in developing the pre/post training datasets for Nemotron models

Goal	SDG Tokens	Results
Search Tool Calling Accuracy for Nemotron Super v3 (this workshop)	7000 records, 0.087B tokens	0% to 31.28% on BrowseComp
Multi-Turn Tool Call Accuracy for Nemotron Nano, Super, Ultra v3	2M Records	2.5% to 27% on multi-turn accuracy
Structured Output Generation for Nemotron Nano, Super, Ultra v3	0.07B Tokens	+6.3% on JSONSchemaBench (80.2% → 86.5%) +8.0% across StructEval-T (TOML, JSON, YAML, CSV and XML)
Structured Query Language Performance Nemotron Nano, Super, Ultra v3	96.5k records	26.77% to 41.80% BIRD-bench
Cross Domain Scientific Reasoning (InfiniByte) for Nemotron Nano, Super, Ultra v3	14.9B Tokens	Validated SDG by showing a 10% improvement on Qwen-2.5 on SciCode Benchmark
STEM performance for Nemotron Nano, Super, Ultra v3 <i>This contributed to 2% of all pretraining data, making it the 3rd-highest-contributing SFT dataset in Super pretraining</i>	140B Tokens	• HumanEval Plus (+3% absolute) • MMLU-Pro (+4% absolute) • Math Reasoning (5% absolute) • MBPP (+3% absolute)
Nemotron-Personas A collection of multilingual synthetic persona datasets that support sovereign AI development across many countries and regions. Used to produce more diverse and rich data across variety of applications	Billions	• Safety data • Instruction-following data • Chat data • Multi-turn tool calling data