

goto;

GOTO AARHUS 2022

#GOTOaar



WHO AM I



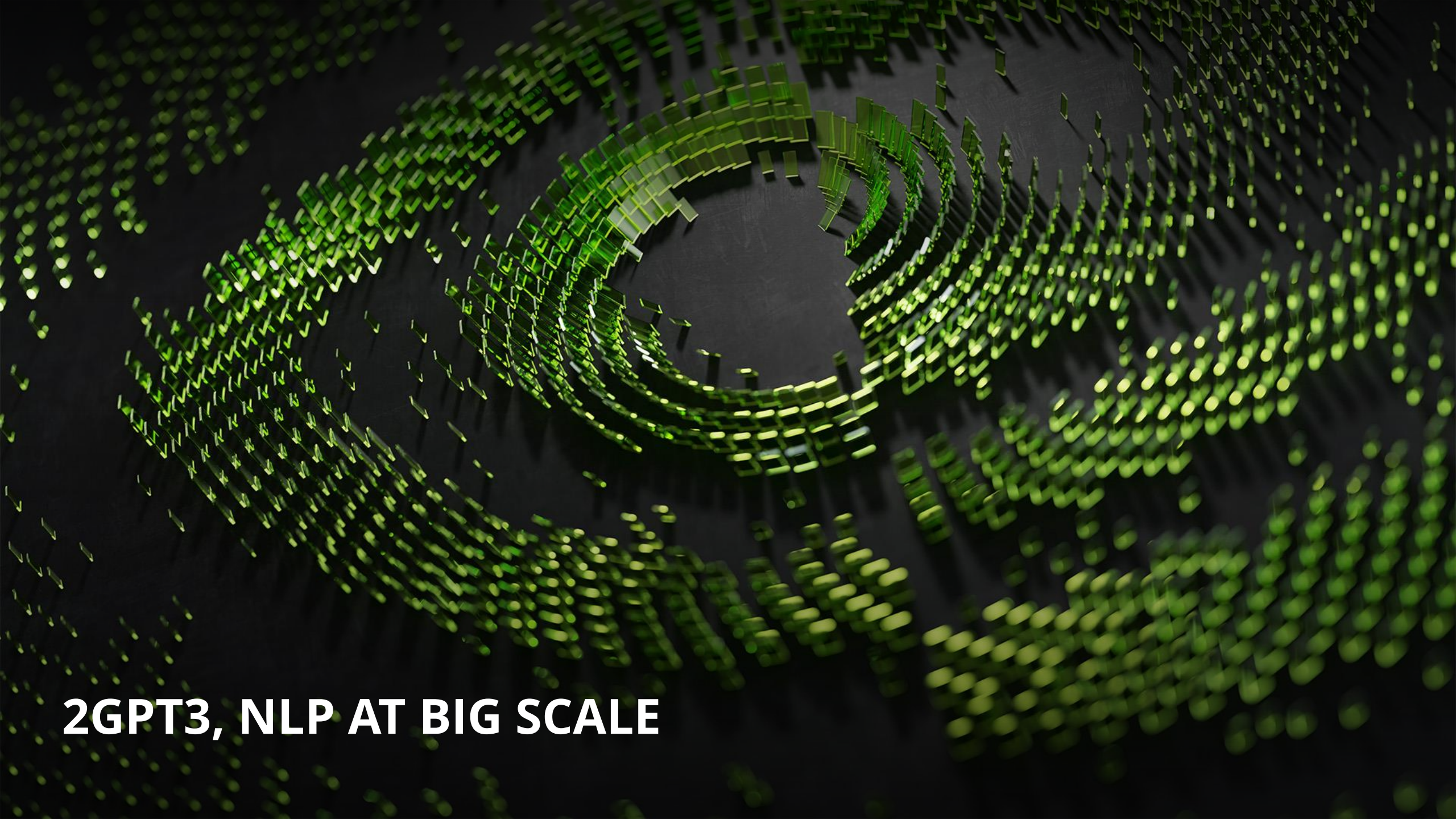
Zenodia Charpy

Senior Deep Learning Data Scientist @ NVIDIA - NeMo | Megatron | bigNLP

My past experience:

Azure : senior data scientist & solution architect

+8 years experience as in-house data scientist & data science consultant

An abstract, high-tech background featuring a complex, repeating pattern of green, three-dimensional rectangular blocks. These blocks are arranged in a way that creates a sense of depth and movement, resembling a stylized, glowing circuit board or a data visualization. The pattern is set against a dark, almost black background, which makes the green elements stand out prominently. The overall effect is one of sophisticated technology and data processing.

2GPT3, NLP AT BIG SCALE

Overview

The 2 GPT3s

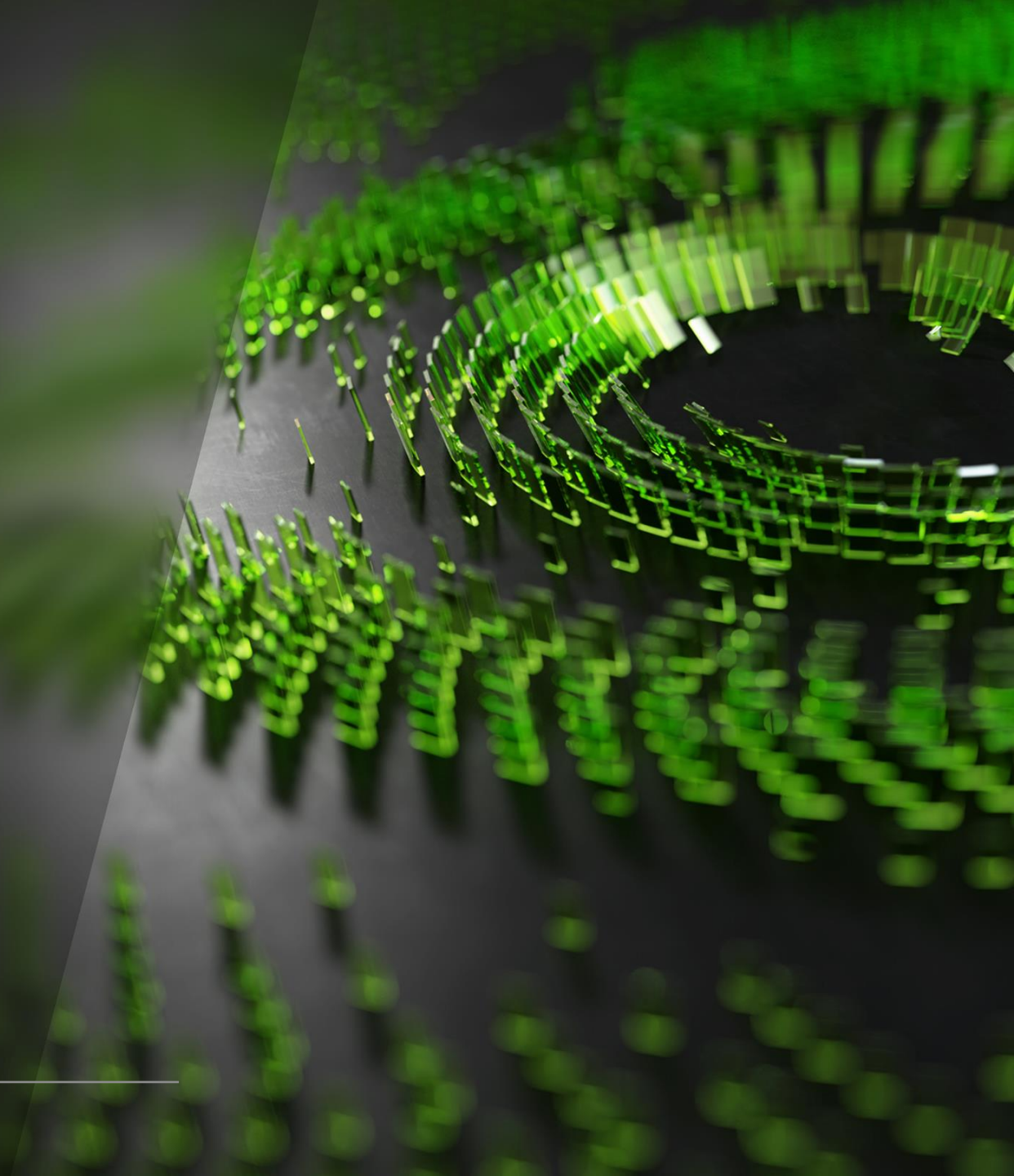
The big 530B Megatron-Turing NLG GPT3 model

Motivation - Why Very Large Language Models

3 basic ingredients

Steps to follow - workflow

QnA – reach out and kick-start



The 2 GPT3

One *English*, the other *Swedish*

	English	Swedish
Data size	Terabytes of data	~40GB
Model size	530B	760M
Compute	4480 GPUs (560 nodes)	8 GPUs (1 nodes)
Time to train	~ 20 days	~ 7 days
Use cases	Mutiple downstream tasks with 0/1/fewshots	Limited to a handfull of cases

<https://arxiv.org/pdf/2201.11990.pdf>

	English	Swedish
Data size	Terabytes of data	~40GB
Model size	530B	760M
Compute	4480 GPUs (560 nodes)	16 GPUs (2 nodes)
Time to train	~20 days	~7 days
Use cases	Mutiple downstream tasks with 0/1/fewshots	Limited to a handfull of cases

530B MEGATRON TURING NLG
- ENGLISH GPT3

A close-up, low-angle shot of a dense array of green, needle-like structures on a dark surface. The structures are arranged in a grid-like pattern, with some in sharp focus in the foreground and others blurred in the background, creating a bokeh effect. The lighting is dramatic, highlighting the texture and color of the green elements against the dark background.

VIRTUAL ASSISTANCE





Customer Service Kiosk

with NVIDIA Omniverse Avatar for Project Tokkio



Large NLP models powers:

- Zero-shot intent and slot classification
- Context-based extractive Q&A
- Conversational phrasing of answer



TOY JENSEN - ASK ME ANYTHING





Expert, Natural Q&A

with NVIDIA Omniverse Avatar
for Project Tokkio

Large NLP models powers:

- Multi-turn Information Retrieval for Q&A



OTHER USE CASES

SOLVE RIDDLE

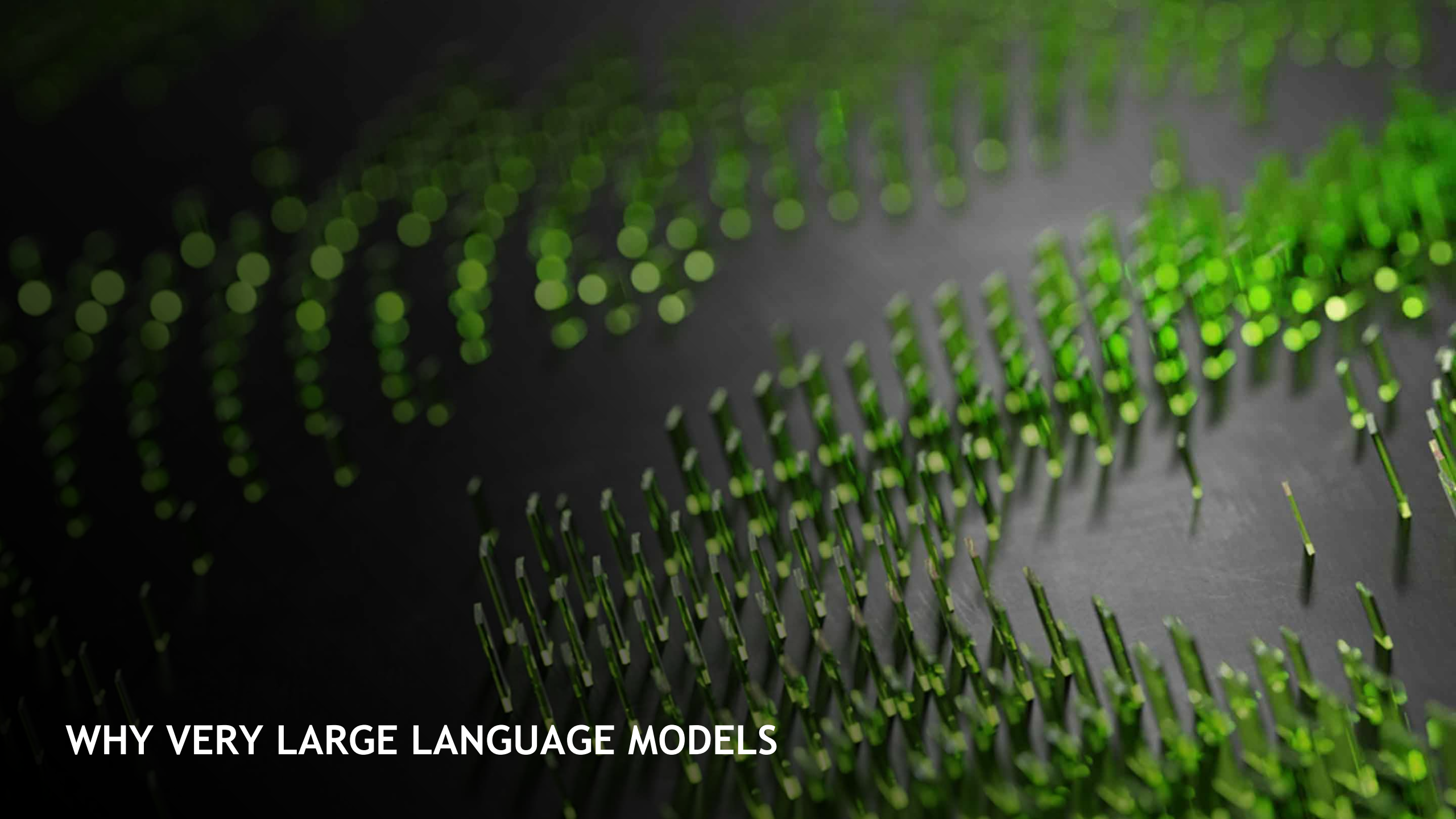
Context → Here is a riddle:
All of us have one, but few get to choose
If you don't know mine, you are not my friend
When it is called, our attention is drawn
Even if we are gone, they are still around

I think the answer is

CODE GENERATION

Context →

```
def find_3_or_7_divisible_fibs(n):  
    """Find all the Fibonacci numbers below n that are divisible by 3 or divisible by  
       7.  
    """
```

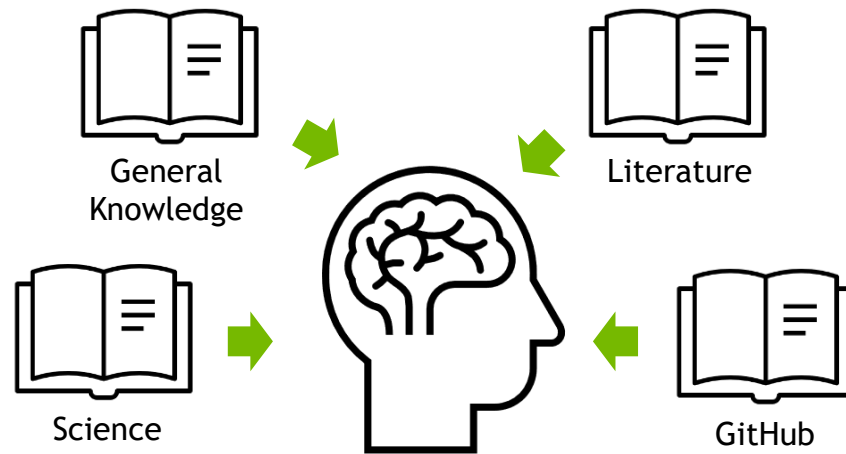



WHY VERY LARGE LANGUAGE MODELS

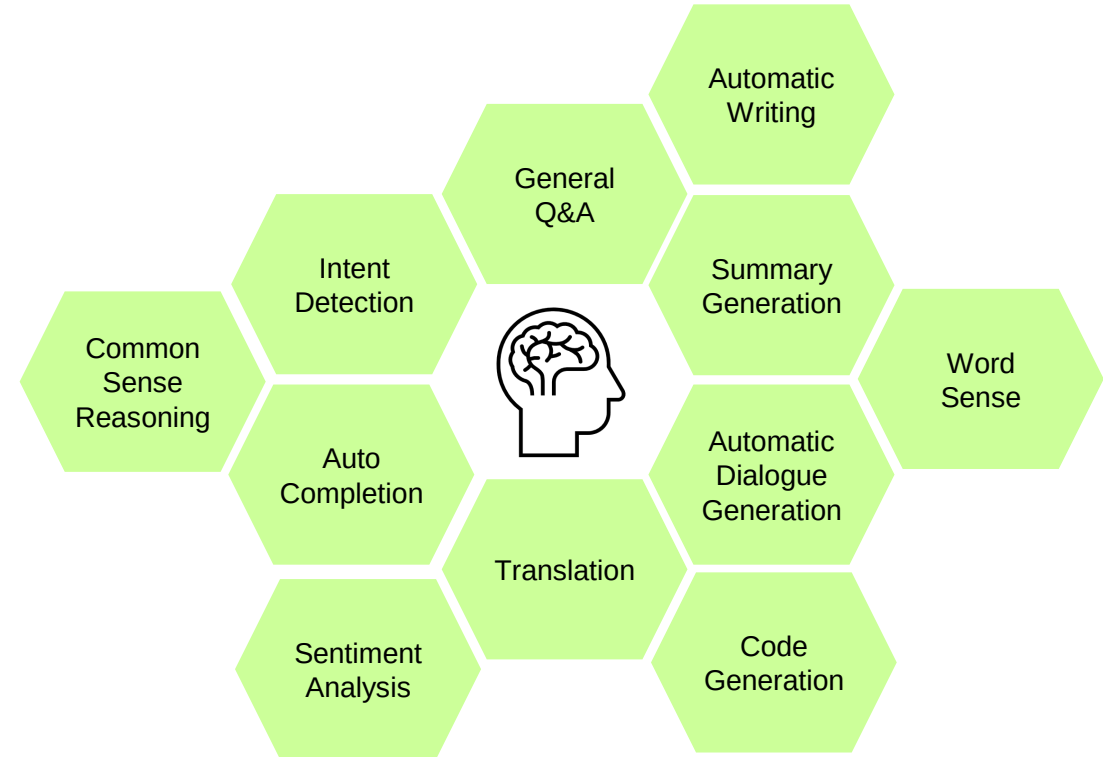
WHY PRE-TRAIN VERY LARGE NLP MODELS

Train Once, Perform different tasks

Step 1: Train a Very Deep/Huge model

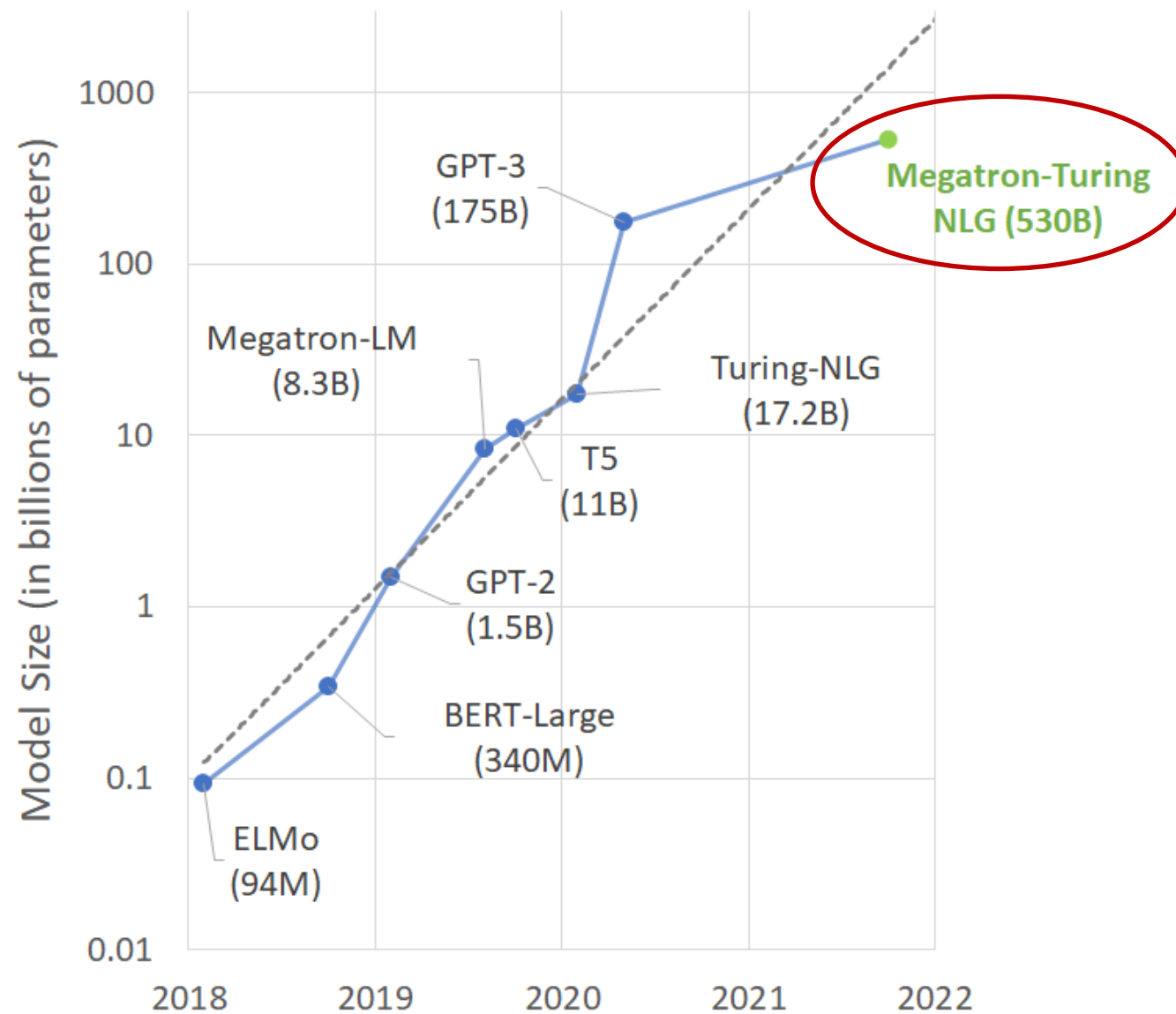


Step 2. Query different tasks with a few labelled data



Huge means Billions of parameters

MODEL SIZE GROWS EXPONENTIALLY



[Link to the article](#)



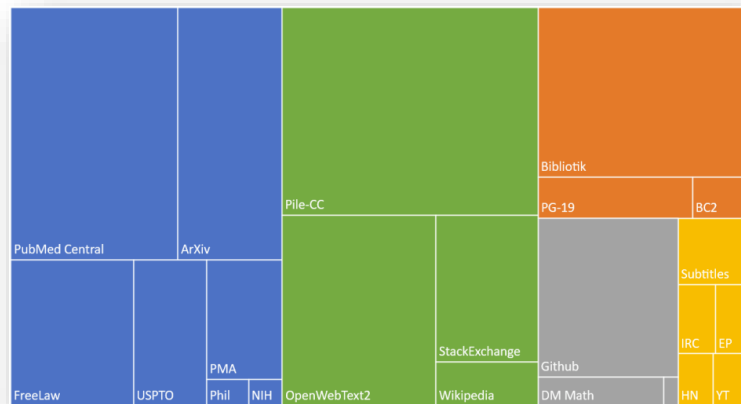
BUT THAT'S ENGLISH
- WHAT ABOUT LOCALIZED LANGUAGE ?

	English	Swedish
Data size	Terabytes of data	~40GB
Model size	530B	760M
Compute	4480 GPUs (560 nodes)	16 GPUs (2 nodes)
Time to train	~20 days	~7 days
Use cases	Mutiple downstream tasks with 0/1/fewshots	Limited to a handfull of cases

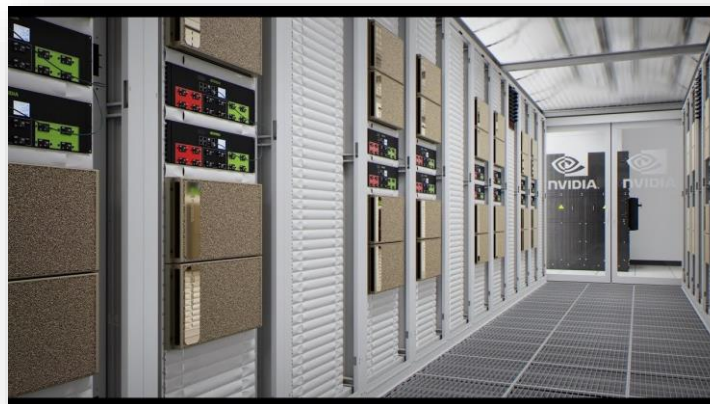
SECURE 3 BASIC INGREDIENTS

3 BASIC INGREDIENTS

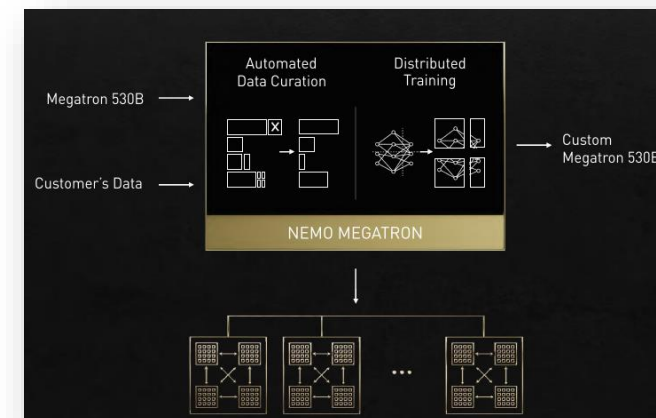
The Data

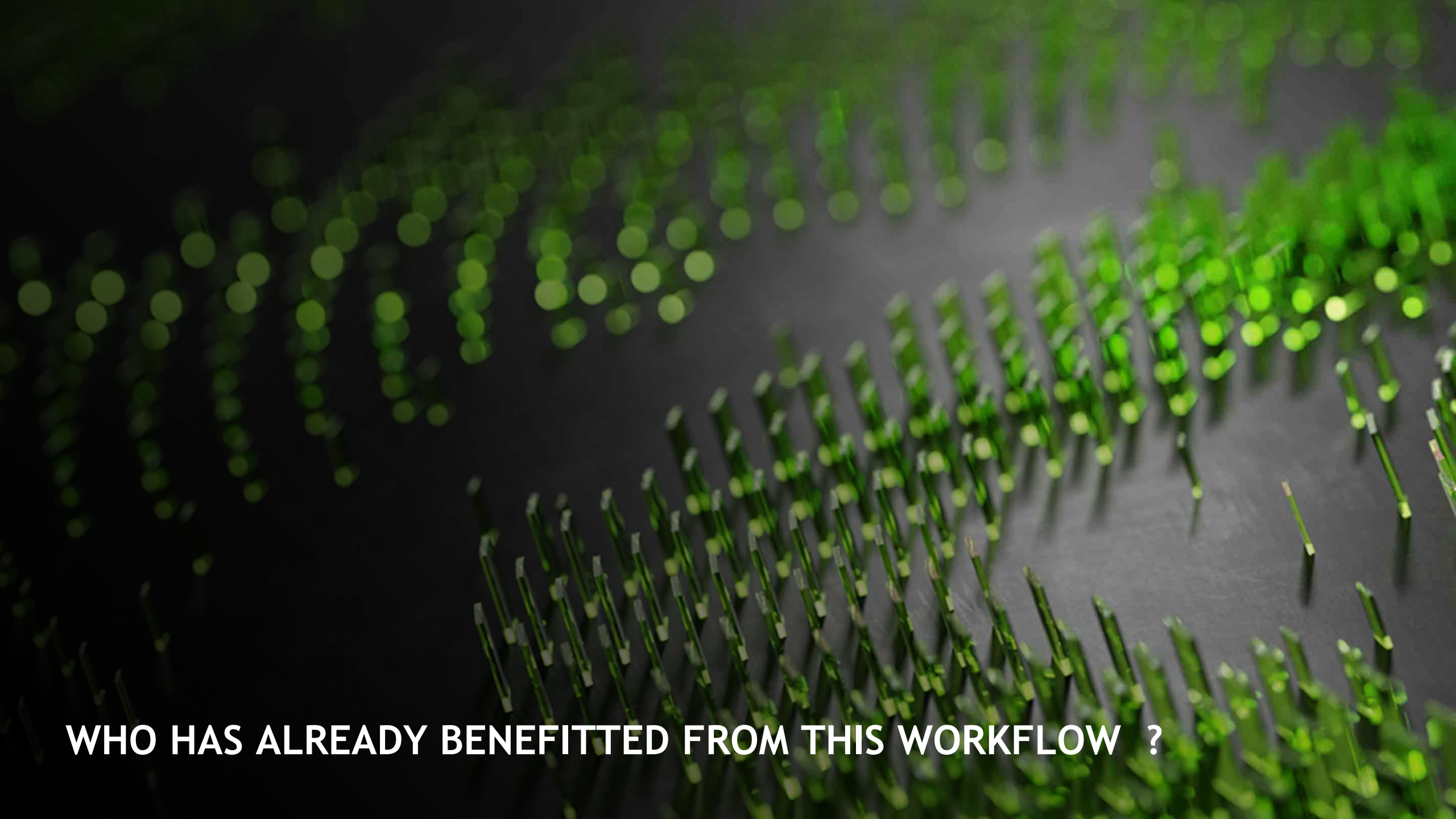


The Compute



The Framework





WHO HAS ALREADY BENEFITTED FROM THIS WORKFLOW ?

WHO HAS ALREADY USED THIS WORKFLOW ?

And trained + deployed their GPT3 models



Published in AI Sweden



Magnus Sahlgren

Jun 8 · 3 min read · [Listen](#)



As the name suggest, **GPT-SW3** is the Swedish counterpart to GPT-3: a large-scale generative pretrained Transformer built on massive amounts of Swedish data with the explicit purpose of being able to generate Swedish text, and to thereby function in zero-shot and few-shot scenarios. The initiative to develop such a model for Swedish is led by the NLU research group at [AI Sweden](#), and includes a number of collaborators, including RISE, the WASP WARA on media and language, and Nvidia. GPT-SW3 already exists in a preliminary version with 3.5 billion parameters, which was trained on a moderately sized corpus of approximately 100 GB of mostly web data. Our current goal is to train a 175 billion parameter model with close to 1 TB of data.

[Link to the article](#)

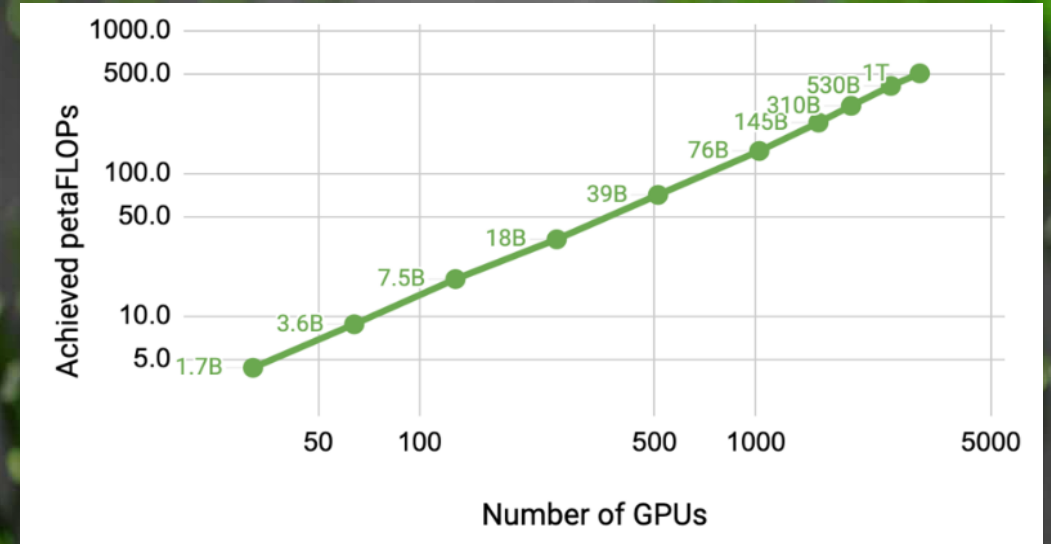
GPT-3 på svenska – de bygger en "superintelligens" som pratar svenska

2022-06-09 06:00

Av: [Simon Campanello](#)

0 kommentarer

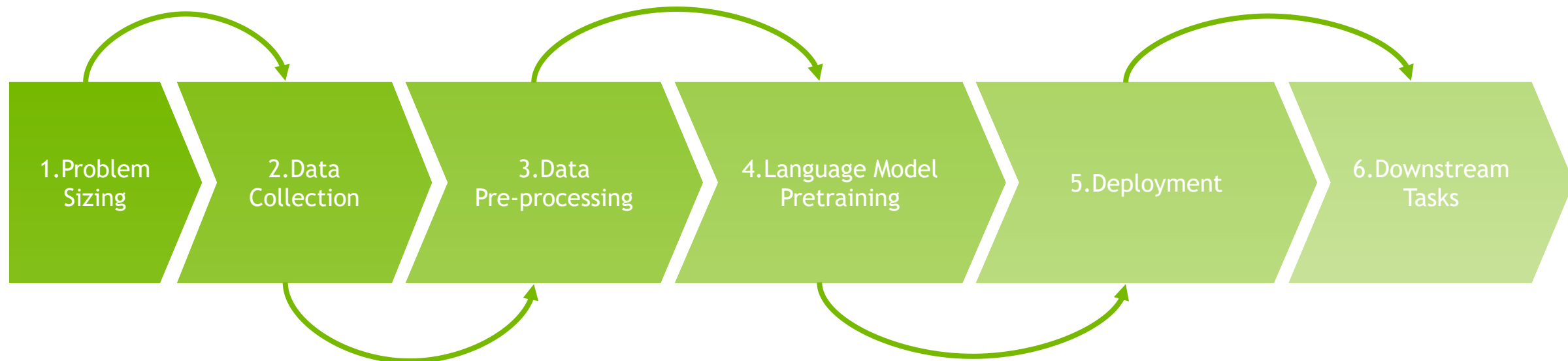


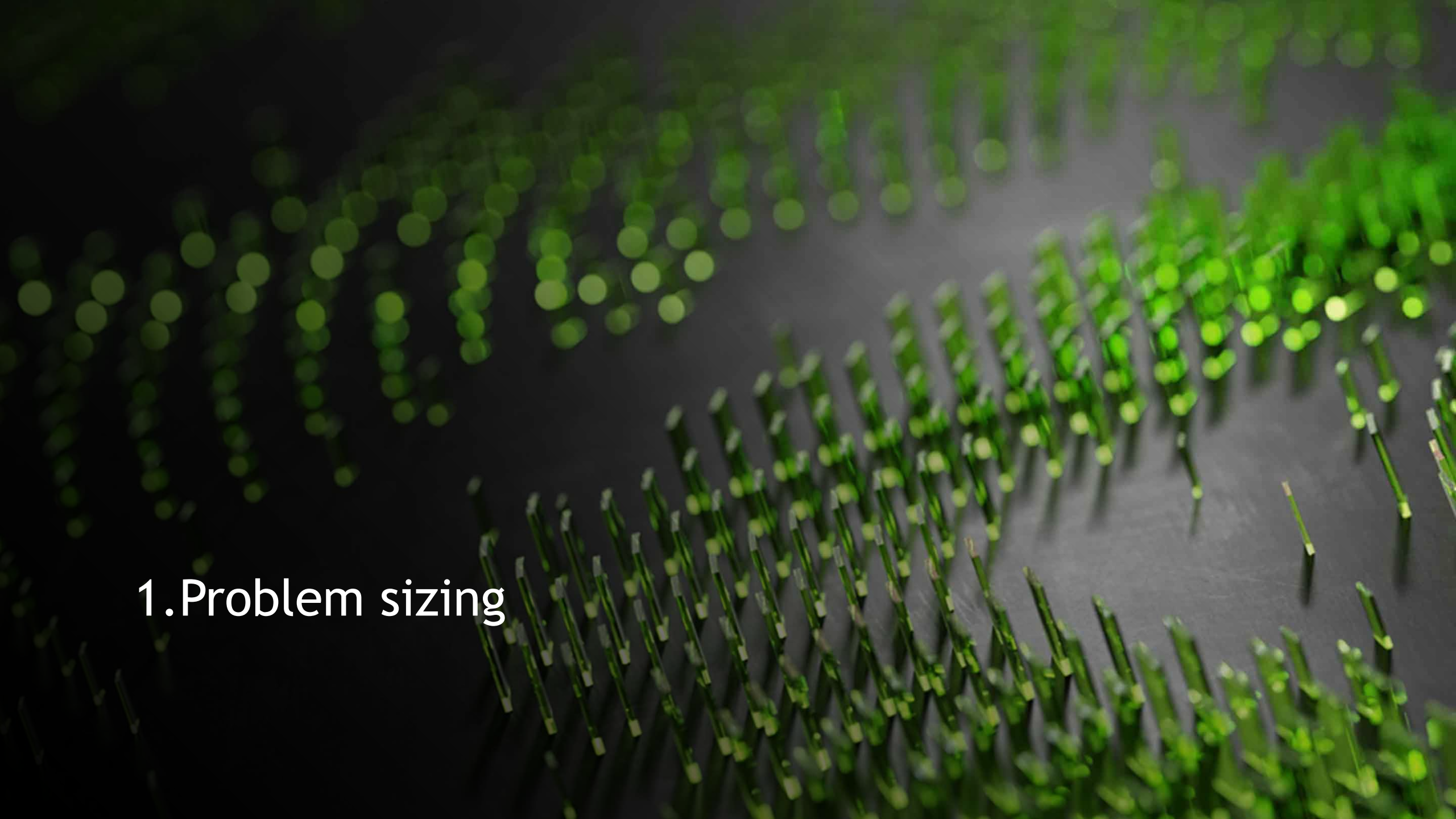


KICK START YOUR JOURNEY - THE WORKFLOWS

THE WORKFLOW - START THE JOURNEY

Generic workflow - works for small to large to very large NLP models





1. Problem sizing

THE ITERATION TIME

Problem sizing - reducing years to weeks/days

GPT-3 Model Parameter Count	Estimated Time to Train			
	5 x SuperPod 100 nodes (days)	3 x SuperPod 60 nodes (days)	1 SuperPod 20 nodes (days)(weeks)	4 x DGX A100 (weeks)(years)
0.5B	0.21	0.21	0.62	0.44
18B	5	8	23	16
175B	43	72	31	3.0

Training

NVIDIA DGX A100 SuperPOD



The background of the slide is a photograph of a dense forest. The trees are tall and thin, with green foliage. The perspective is from a low angle, looking up at the trees, which creates a sense of height and density. The background is dark and out of focus, emphasizing the vertical lines of the trees.

2. Collect data - part 1/3

- Training data



DATA COLLECTION

Training data: Quantity vs. Quality



OSCAR

Open Super-large Crawled Aggregated corpus

<https://oscar-corpus.com/>

Composition of the Pile by Category

■ Academic ■ Internet ■ Prose ■ Dialogue ■ Misc

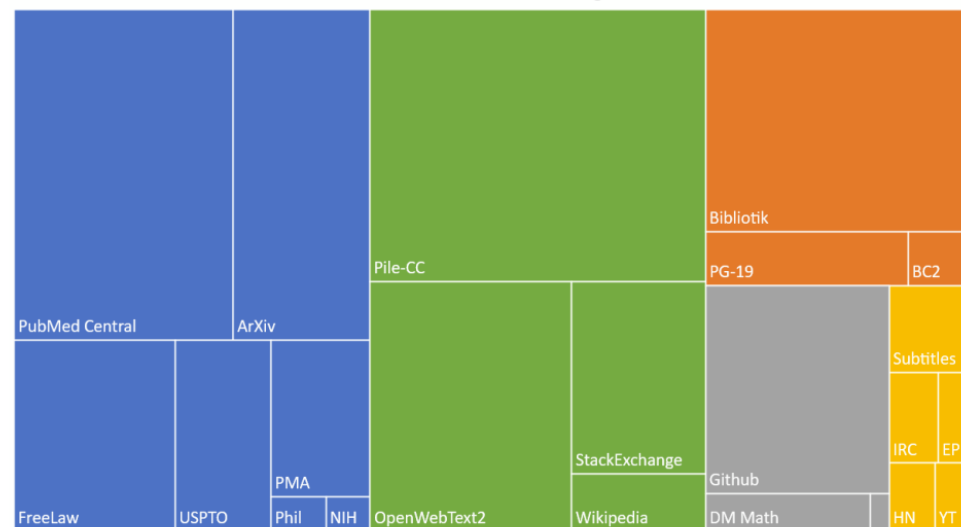


Figure 1: Treemap of Pile components by effective size.



An 800GB Dataset of Diverse Text for Language Modeling

<https://pile.eleuther.ai/>

WEIGHING AGAINST QUALITY VS QUANTITY

Weighing mechanism is specific to the datasets used for training

Dataset	Tokens (billions)	Weights (%)	Epochs
Books3	25.7	14.3	1.5
OpenWebText2	14.8	19.3	3.6
Stack Exchange	11.6	5.7	1.4
PubMed Abstracts	4.4	2.9	1.8
Wikipedia	4.2	4.8	3.2
Gutenberg (PG-19)	2.7	0.9	0.9
BookCorpus2	1.5	1.0	1.8
NIH ExPorter	0.3	0.2	1.8
Pile-CC	49.8	9.4	0.5
ArXiv	20.8	1.4	0.2
GitHub	24.3	1.6	0.2
CC-2020-50	68.7	13.0	0.5
CC-2021-04	82.6	15.7	0.5
RealNews	21.9	9.0	1.1
CC-Stories	5.3	0.9	0.5

The background of the slide is a photograph of a dense forest of tall, thin, green reeds or grasses. The reeds are in sharp focus in the foreground, showing their segmented structure and vibrant green color. They are set against a dark, blurred background, creating a strong sense of depth and texture. The lighting appears to be coming from the side, highlighting the edges of the reeds.

2. Collect data - part 2/3

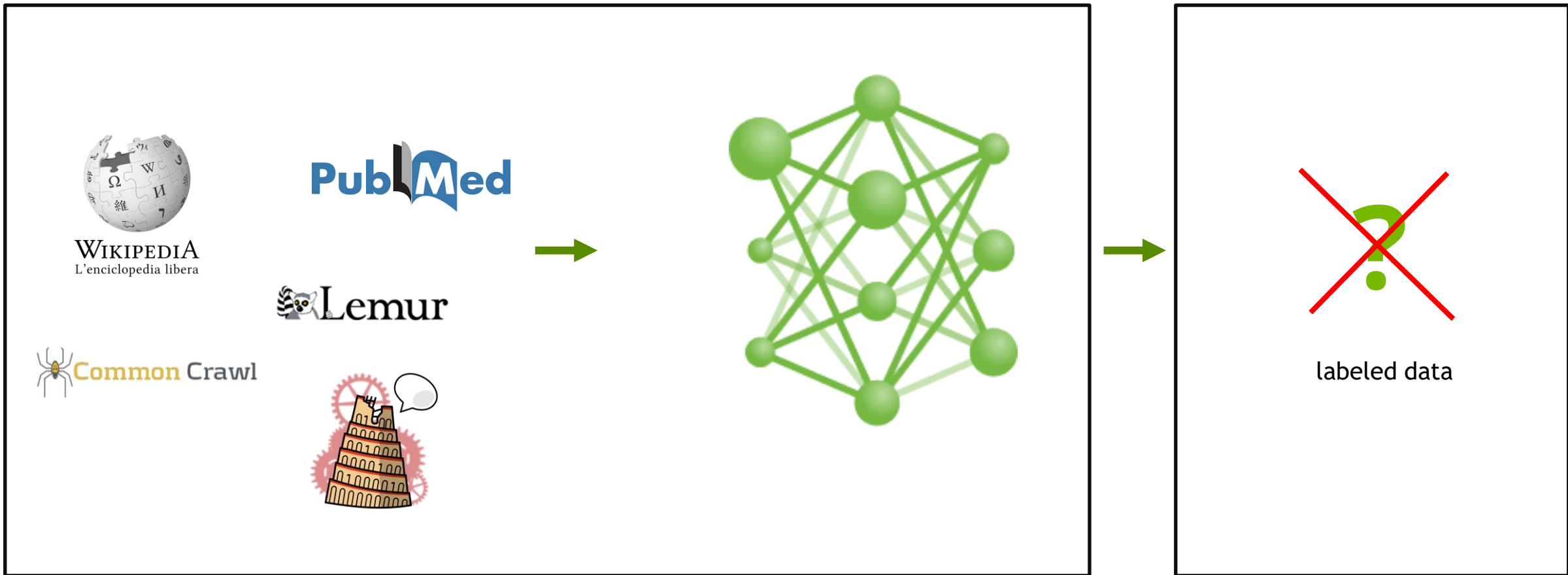
- Data for downstream tasks

DATA COLLECTION

Annotating and creating data for downstream tasks

Train the GPT model unsupervised

Evaluate



DATA + **LABEL** COLLECTION

Annotation requires Tools and Expertise



The screenshot shows the 'SPRÅKBANKENTEXT' website. The header includes the University of Gothenburg logo and the word 'Svenska'. A navigation bar contains links for NEWS AND EVENTS, BLOG, RESEARCH, TOOLS (highlighted), RESOURCES, FAQ, ABOUT US, and CONTACT US. Below the navigation bar, a breadcrumb trail reads 'Home / Tools / Sparv / Annotations by Sparv'. The main content area is titled 'Annotations by Sparv' and contains a sidebar with links to 'Sparv Pipeline', 'Sparv's user manual', 'Annotations by Sparv' (selected), 'Web service (API)', and 'Web Sparv'. The main text explains that the page provides an overview of analyses available within the Sparv Pipeline and Sparv plugins. It defines 'annotations' as names used in the corpus config file and 'annotators' as names of annotation functions. It also lists available analyses for contemporary Swedish, such as sentence segmentation with PunktSentenceTokenizer, and provides details like description, model, method, and annotations for each.

Stockholm-Umeå corpus 3.0

The Stockholm-Umeå Corpus (SUC) is a collection of Swedish texts from the 1990's, consisting of one million words in total. The corpus is balanced, meaning that it contains various text types and stylistic levels. The texts are annotated with part-of-speech tags, morphological analysis and lemma (all that can be considered *gold standard data*), as well as some structural and functional information.

Version 1.0 was developed in co-operation between Gunnel Källgren at Stockholm University and Eva Ejerhed at Umeå University and was made available in 1997 by the department of linguistics at Stockholm University.

Version 2.0 was made available in 2006 by Sofia Gustafson-Capková and Britt Hartmann at the department of linguistics at Stockholm University. It contains the same texts as SUC 1.0 but is extended with some annotation. Additionally, SUC 2.0 contains bonus materials. TigerSUC is SUC 2.0 converted to TIGER-XML by Martin Volk. StorSUC is additional SUC material of four million words. *Version 3.0* is available since 2012. It contains improved annotations, and unannotated texts with seven million words. (For the TigerXML-version, Suc2c, Suc2d, and the DTDs we still refer to version 2.0.)

The background of the slide is a close-up, low-angle photograph of a green, textured surface. The texture appears to be composed of many small, vertical, needle-like or hair-like structures, possibly a plant or a digital display. The lighting is dramatic, with the green elements glowing against a dark, almost black background. The focus is sharp in the foreground and slightly blurred in the background, creating a sense of depth.

2. Collect data - part 3/3
- Legal/ethical/responsible AI concerns

DATA COLLECTION

Establish Censoring Criteria



1. Data source licence implications
2. Data characteristics and coverage
3. Sufficiency of data pre-processing pipeline
4. Addressing issues of dataset bias
5. Addressing issues of dataset fairness
6. Addressing to ethical concerns
7. Addressing to legal concernss
8. Blacklisting & toxicity removal
9. Composition and weighing
10. ...

The background of the slide is a photograph of a dense forest of tall, thin, green reeds or grasses. The plants are in sharp focus in the foreground, showing their segmented structure and vibrant green color. They are set against a dark, blurred background, creating a sense of depth and texture.

3.Data preprocessing - part 1/3

- Filter / clean / deduplicate



DATA PRE-PROCESSING

Filtering + Deduplicate + Sentence Boundary

```
root@a0ca60d991fa: /worksj x data_cleaning_demo.ipynb Python 3 C
Step 1 - detect and filter undesired language in the raw text corpus

[1]: import random
      from langdetect import detect
      f=open('./bland_langs.txt','r')
      lines=f.readlines()
      len(lines)

[1]: 1824

[ ]: rn=random.randint(0,len(lines)-1)
      print("sentence = ", lines[rn])
      lang=detect(lines[rn])
      print("detected language is = ", lang)

step 2 - deduplicate based on similarity threshold

sanity check, fetch previously labelled two documents with groundtruth ( i.e duplicated or not )

check against threshold, set to 0.85 and determine these two documents are duplicated or not !

[ ]: import itertools
      from lsh import cache, minhash # https://github.com/mattilyra/lsh

      # a pure python shingling function that will be used in comparing
      # LSH to true Jaccard similarities
      def shingles(text, char_ngram=5):
          return set(text[head:head + char_ngram] for head in range(0, len(text) - char_ngram))
```

Demo: Data pre-processing toy example for the Swedish language

Languages filtering + Deduplication + Sentence boundary (filter short sentences/single sentence document)

Link to [data cleaning example notebook](#)



3.Data preprocessing - part 2/3

- Train your own tokenizer



TRAIN OWN GPT TOKENIZER

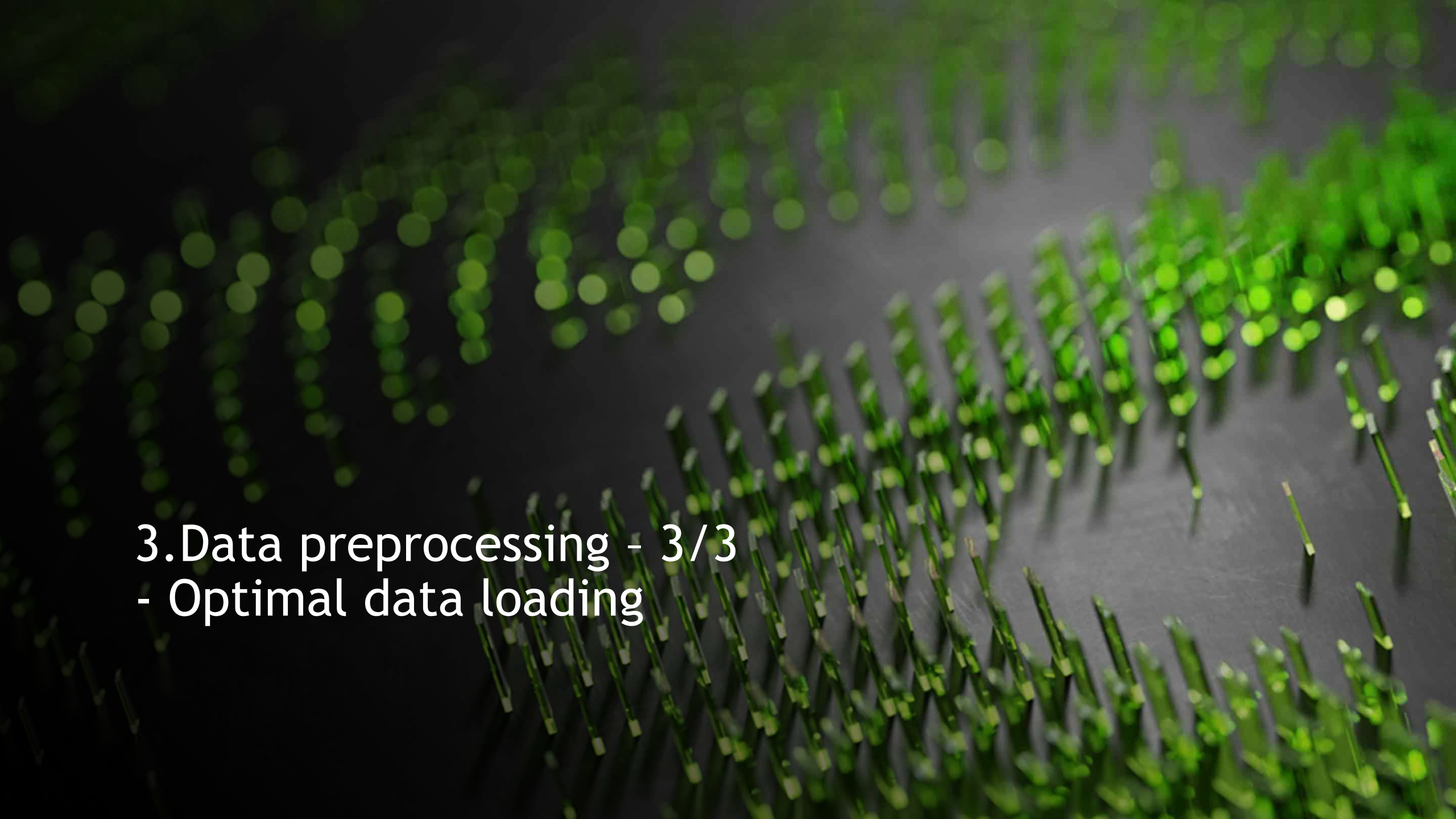
Train a GPT tokenizer on your own raw corpus

```
root@a0ca60d991fa: /workspace/ SVdata/gtc# python trainGPTTokenizer.py --infile ./sv10thousand.txt --bpe_path ./32k/ --vocab_size 32000

I

trainGPTTokenizer.py
17 parser.add_argument('--pretrained_gpt_dir', default=None, type=str, help='path to pretrained gpt vocab and
merge files, default None')
18 parser.add_argument('--vocab_size', default=None, type=int,
19                       help='specify the vocab_size when training HF GPTBPE for own language usually 16k/32k
/48k/64k')
20 args = parser.parse_args()
21 tokenizer = Tokenizer(models.BPE())
22 if args.load_pretrained :
23     if args.pretrained_gpt_dir is not None :
24         print("loading gpt2bpe english vocab and merge \n")
25         vocab_file=args.pretrained_gpt_dir+'/gpt2-vocab.json'
26         merge_file=args.pretrained_gpt_dir+'/gpt2-merges.txt'
27         tokenizer.model = models.BPE.from_file(vocab_file, merge_file)
28     else:
29         print("please provide path to the pretrained gpt vocab and merge file !")
30
31 tokenizer.pre_tokenizer = pre_tokenizers.ByteLevel()
32 tokenizer.decoder = ByteLevelDecoder()
33
34 print("include minimal special token end of text ")
35 special_tokens= ["<|endoftext|>"]
36
37 # Set the training hyperparameters
38 trainer = trainers.BpeTrainer(
39     vocab_size=args.vocab_size,
40     special_tokens=special_tokens,
41     initial_alphabet=pre_tokenizers.ByteLevel.alphabet()
42 )
43 # Train it with either files or an iterator:
44 tokenizer.train([args.infile], trainer=trainer)
45 print("Trained vocab size: {}".format(tokenizer.get_vocab_size()))
46 # You will see the generated files in the output.
47 print("saving trained BPE model to : ", args.bpe_path)
48 tokenizer.model.save(args.bpe_path)
49 print("model saved ! \n\n\n")
50 print("testing ... \n\n\n")
51 test_txt="Har någon funderat på varför man inte får inomhustemperaturens kurva synlig i grafen? Är det någon
som frågat Thermia? Skulle det inte vara väsentligt att kunna kolla historiken på den då man skall ställa in
```

Demo: Train a GPTBPE Tokenizer on a toy Swedish data
Link to [train GPTBPE Tokenizer example notebook](#)



3.Data preprocessing - 3/3

- Optimal data loading



PREPROCESS TO BINARY FILES

Improve data loading performance during training time

Name	Last Modified
old_ckpt	11 days ago
zh_glue	11 days ago
bk_text_sentence_test_indexmap_80mns_512msl_0.10ssp_1234s.npy	23 days ago
bk_text_sentence_train_indexmap_8000000mns_512msl_0.10ssp_1234s.npy	23 days ago
bk_text_sentence_valid_indexmap_800080mns_512msl_0.10ssp_1234s.npy	23 days ago
bk_text_sentence.bin	23 days ago
bk_text_sentence.idx	23 days ago
demo_gpt_sentence_text_sentence.bin	an hour ago
demo_gpt_sentence_text_sentence.idx	an hour ago

mmap :
process data into xxx.idx + xxx.bin for
optimized memory mapping to improve
performance



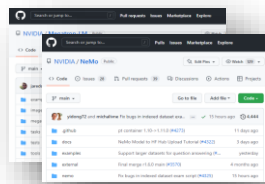
4. Training considerations



CONSIDERATION FOR TRAINING LAUNCHES

Manual vs automatic : obtain optimal training/inference configurations

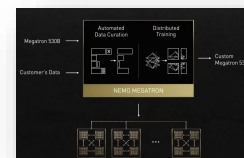
► DIY(Megatron-LM)



1. Manual runs to obtain optimal training config
2. Manually develop inference route for pretrained models
3. Manually debug with limited support



► Enterprise ready(NeMo Megatron)



1. Get automatic recommendation of optimal training config from the HP tool
2. Support by default training and Inference per model size in the template
3. Get enterprise-grade expertise support

Control



Manual work

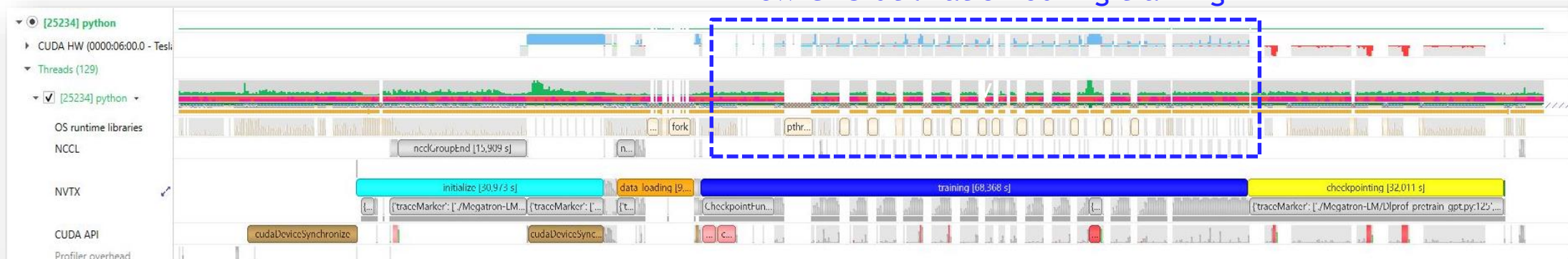




DIY- MANUALLY TUNE TO ENSURE PERF VIA PROFILING

Profile training runs | improve iteratively

Naive config



Improved config





DIY - OBTAIN GOOD CONFIG FOR MODEL PRETRAINING

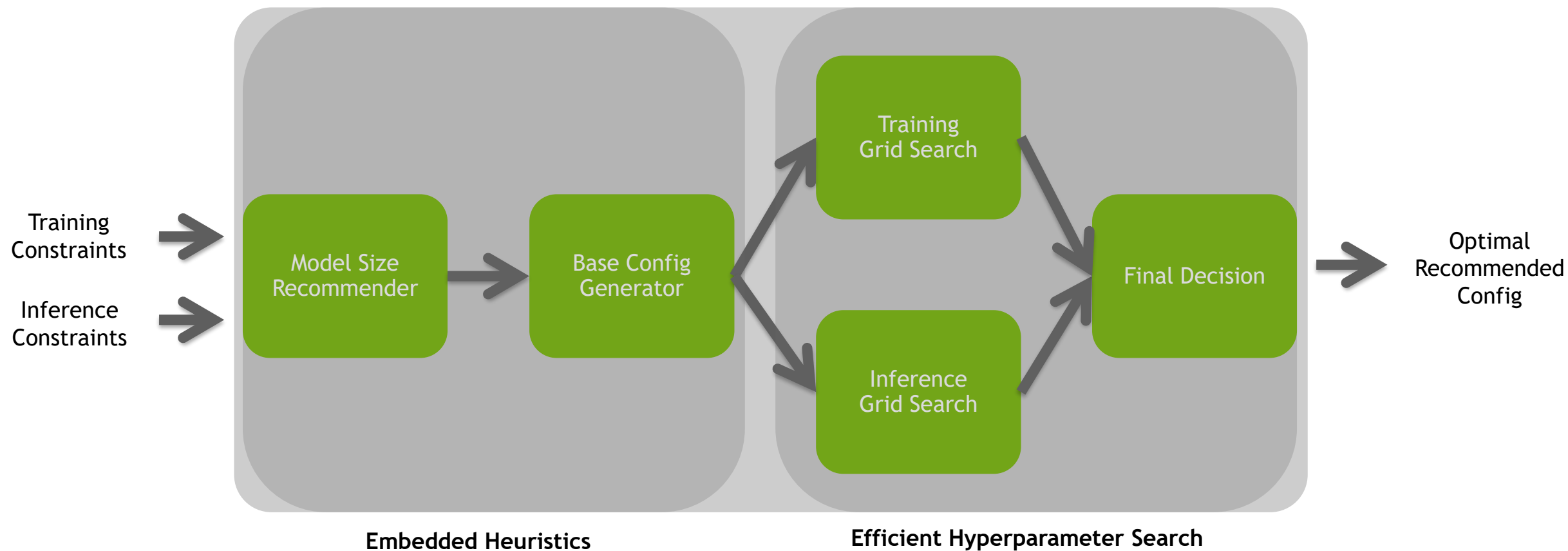
Training configuration performance obtained via manual experiments

based on 100		Activation	4							312	A100 FP16 TFLOPS
		No Activation	3							112	V100 PCIe FP16 TFLOPS
elapsed time per				Find good training config						125	V100 SXM
Speed ms	N Gpus	Number of Nodes	Model	Micro	DP Size	TP Size	PP Size	PP		NVIDIA	Lower bound utilisation
31579	8	2	18	4	1	4	4	128	156.514	Teraflop util %	50.16%
44438	8	2	18	6	1	4	4	128	166.835		53.47%
31532	8	2	18	4	1	4	8	128	156.747		50.24%
44688	8	2	18	6	1	4	8	128	165.902		53.17%
62078	8	2	18	4	1	4	8	256	159.236		51.04%
15889	8	2	18	8	1	4	8	32	155.533		49.85%
16355	8	2	18	4	1	4	8	64	151.102		48.43%
31191	8	4	39	4	1	4	8	128	167.118		53.56%
45146	8	4	39	6	1	4	8	128	173.191		55.51%
87563	8	4	39	6	1	4	8	256	178.588		57.24%
173505	8	4	39	6	1	4	8	512	180.257		57.77%
43119	8	8	52	4	1	4	16	256	160.871		51.56%
43143	8	8	52	4	1	4	16	256	160.782		51.53%
43120	8	8	52	4	1	4	16	256	160.868		51.56%
43154	8	8	52	4	1	4	16	256	160.741		51.52%
62427	8	8	52	6	1	4	16	256	166.673		53.42%
80920	8	8	52	8	1	4	16	256	171.444		54.95%
121531	8	8	52	6	1	4	16	512	171.231		54.88%
158302	8	8	52	8	1	4	16	512	175.275		56.18%
43112	8	8	52	4	1	4	32	256	160.897		51.57%



ENTERPRISE READY - PRETRAINING WITH RECOMMENDED CONFIG

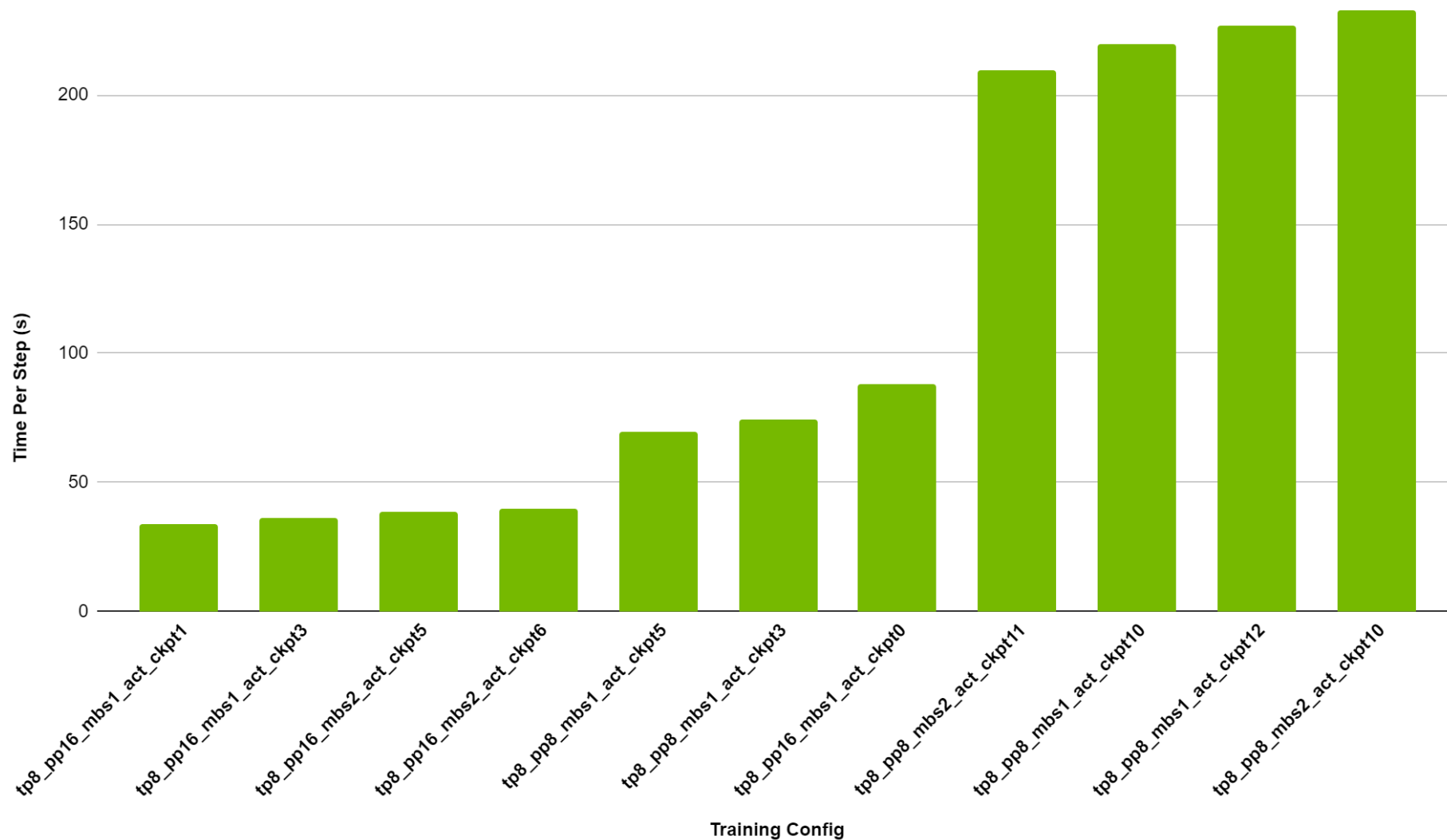
NVIDIA NeMo Megatron makes training easy with recommended config & perf guarantee

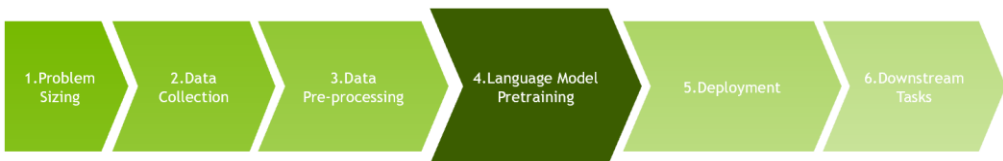




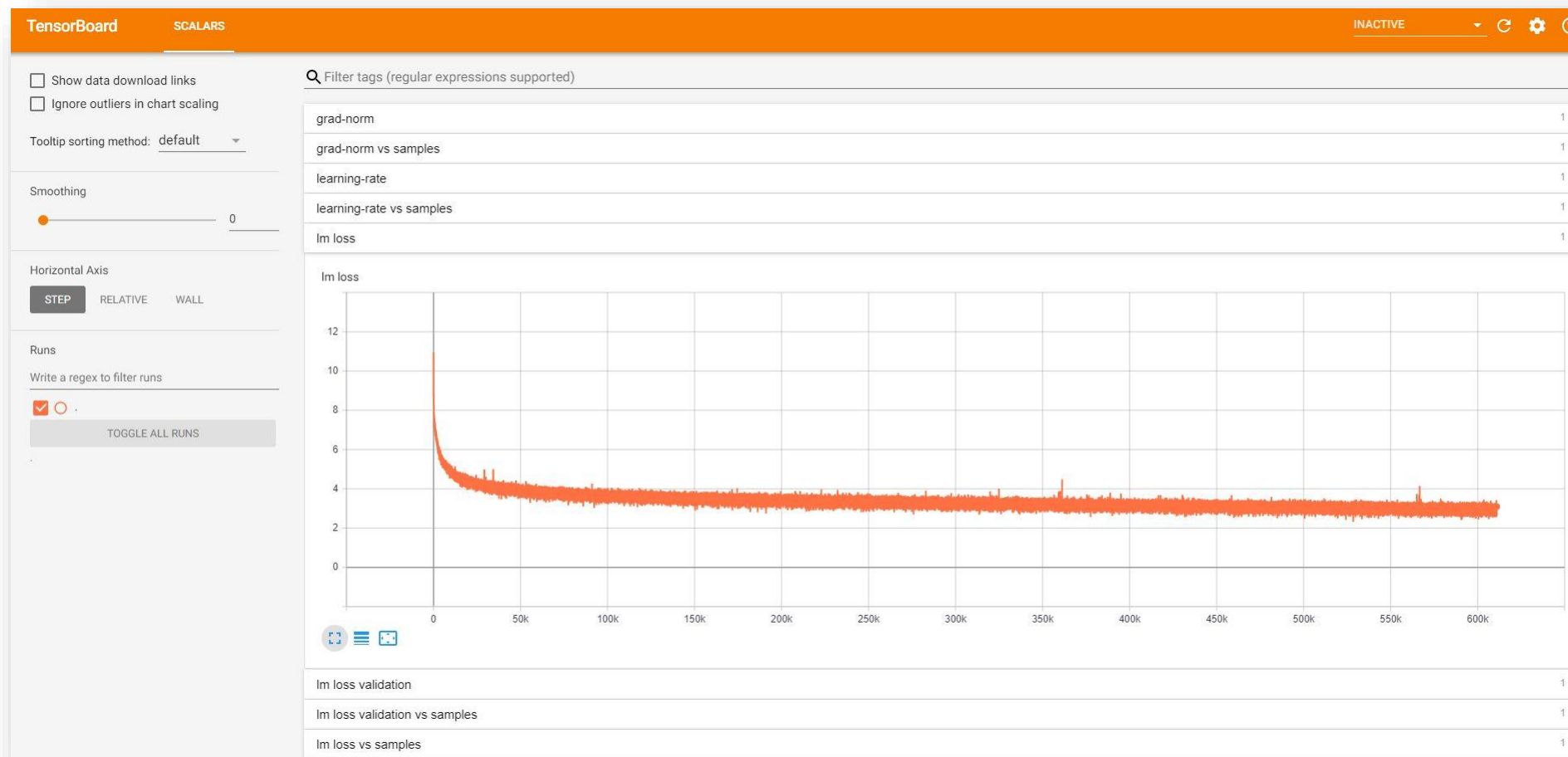
ENTERPRISE READY - PRETRAINING WITH RECOMMENDED CONFIG

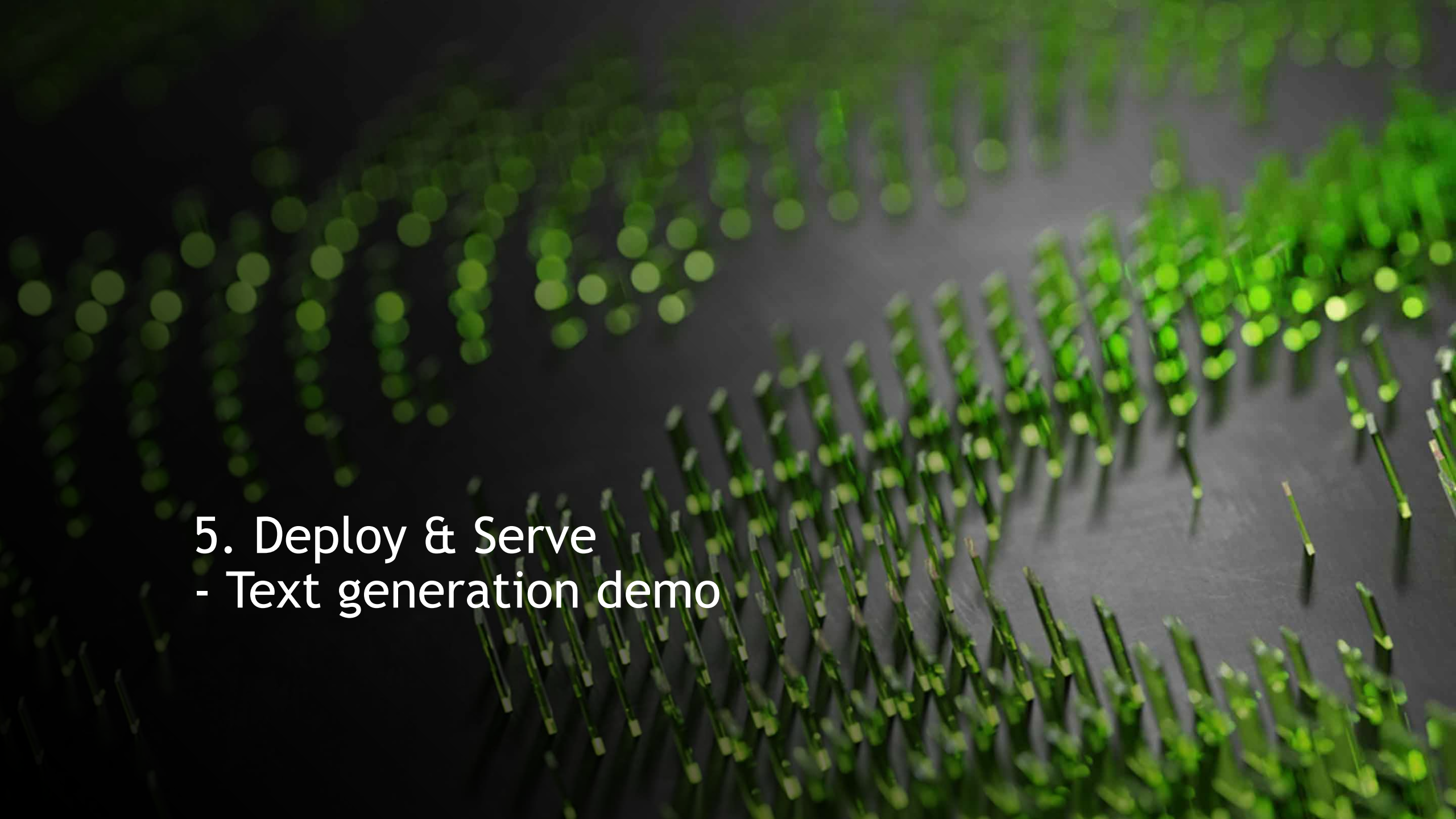
NVIDIA NeMo Megatron makes training easy with ready-made template config & perf guarantee





MONITOR TENSORBOARD





5. Deploy & Serve

- Text generation demo

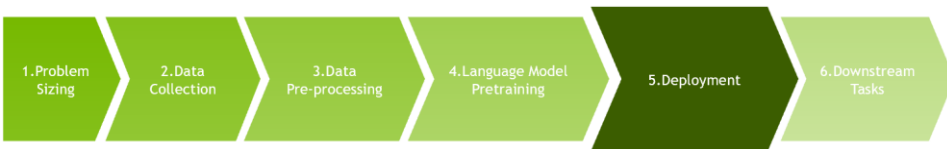


```
localhost:2345/lab? 110% ☆
File Edit View Run Kernel Tabs Settings Help
+ + + + +
/ SVdata / gpt2bpe /
Name
32k
48k
56k
gpt2-merges.txt
gpt2-vocab.json
SVGPT_32k_text_d...
SVGPT_32k_text_d...

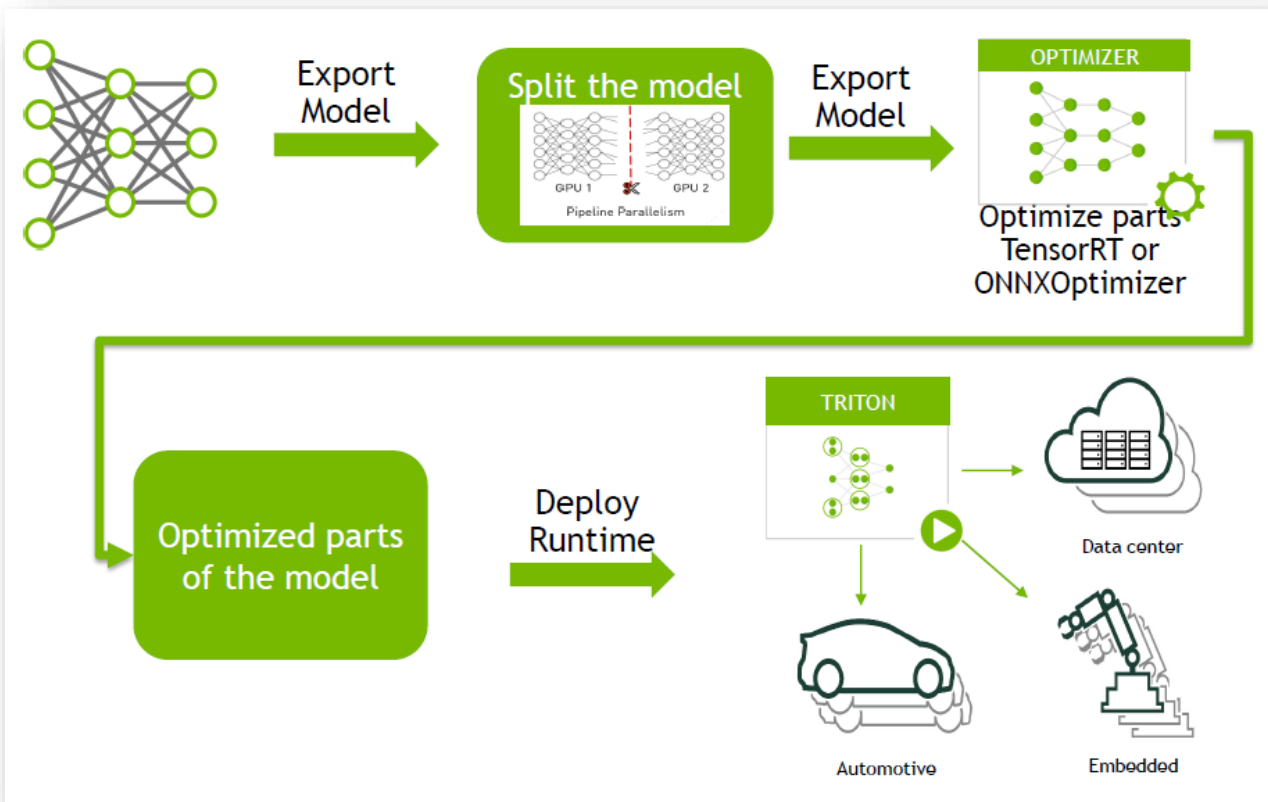
root@3c633147b13b: /works Step3_ViewGenerated_text.i Step2_760MGPTGenerateTe
use_contiguous_buffers_in_ddp ..... False
use_cpu_initialization ..... None
use_one_sent_docs ..... False
virtual_pipeline_model_parallel_size ..... None
vocab_extra_ids ..... 0
vocab_file ..... /workspace/SVdata/gpt2bpe/32k/vocab.json
weight_decay ..... 0.01
world_size ..... 1
----- end of arguments -----
setting number of micro-batches to constant 1
> building GPT2BPETokenizer tokenizer ...
> padded vocab (size: 32000) with 0 dummy tokens (new size: 32000)
> initializing torch distributed ...
> initializing tensor model parallel with size 1
> initializing pipeline model parallel with size 1
> setting random seeds to 1 ...
> initializing model parallel cuda seeds on global rank 0, model parallel rank 0, and data parallel rank 0 with model parallel seed: 2719 and data parallel seed: 1
> compiling dataset index builder ...
make: Entering directory '/workspace/megatron/data'
make: Nothing to be done for 'default'.
make: Leaving directory '/workspace/megatron/data'
>>> done with dataset index builder. Compilation time: 0.038 seconds
WARNING: constraints for invoking optimized fused softmax kernel are not met. We default back to unfused kernel invocations.
> compiling and loading fused kernels ...
Detected CUDA files, patching ldflags
Emitting ninja build file /workspace/megatron/fused_kernels/build/build.ninja...
Building extension module fused_mix_prec_layer_norm_cuda...
Allowing ninja to set a default number of workers... (overridable by setting the environment variable MAX_JOBS=N)
ninja: no work to do.
Loading extension module fused_mix_prec_layer_norm_cuda...
>>> done with compiling and loading fused kernels. Compilation time: 1.028 seconds
building GPT model ...
> number of parameters on (tensor, pipeline) model parallel rank (0, 0): 729897984
loading checkpoint from /workspace/tools/gpt3_iter2mil/ at iteration 2000000
checkpoint version 3.0
successfully loaded checkpoint from /workspace/tools/gpt3_iter2mil/ at iteration 2000000
Avg s/batch: 0.9077534675598145
root@3c633147b13b: /workspace#
```



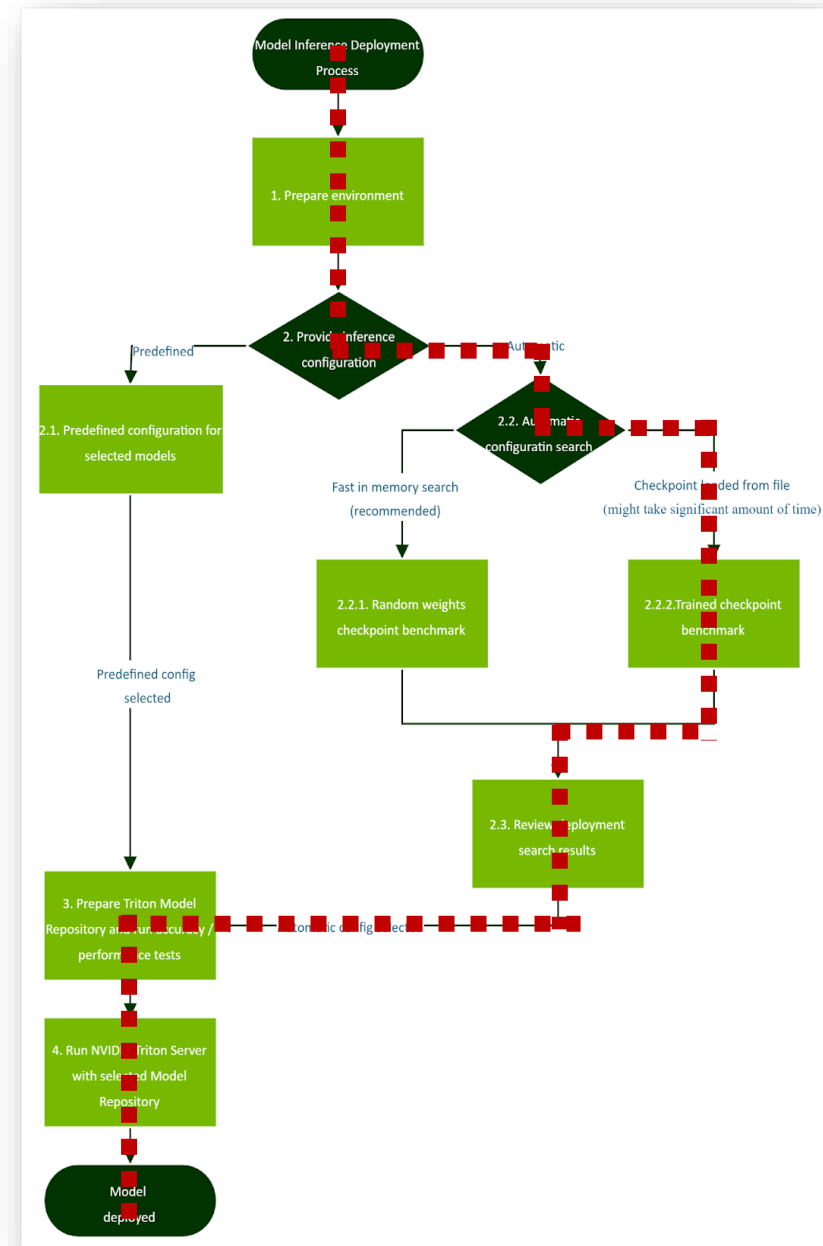

5. Deploy & Serve



DIY(Megatron-LM) – Manually develop deployment route



Enterprise ready(NeMo Megatron) - support by default automtaic deployment route





P-tuning for downstream tasks

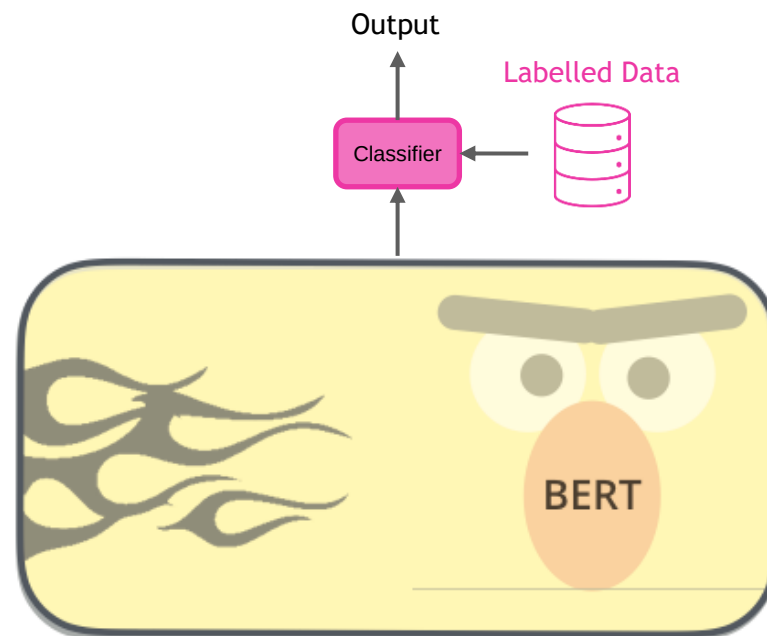
NLP APPROACH (CIRCA 2019)

Fine tune per specific task

Step 1: Pre-training a Encoder



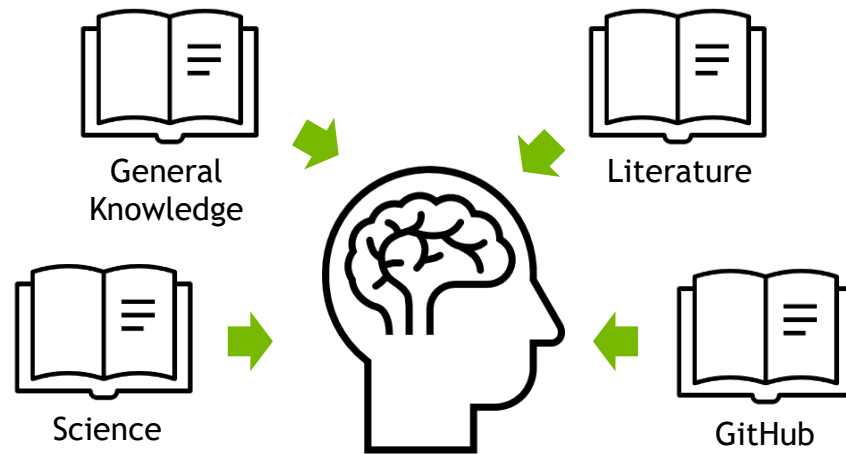
Step 2. Fine tune for a specific task



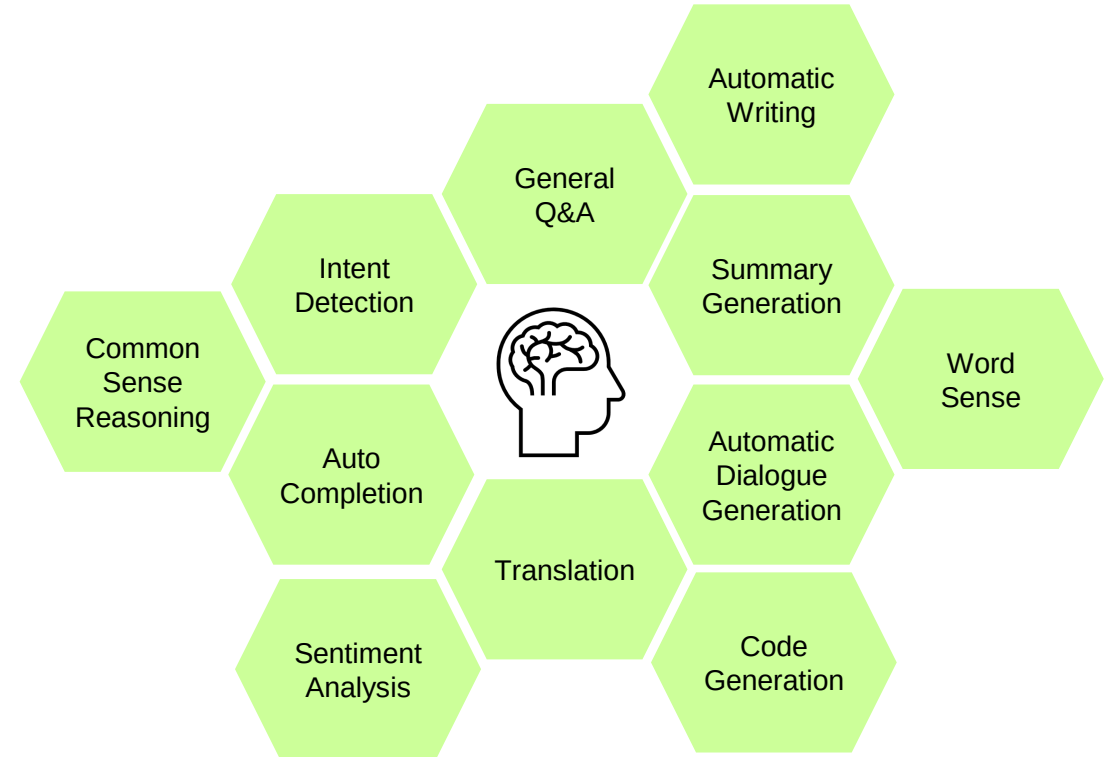
NEW NLP APPROACH (CIRCA 2021)

Train Once, Perform different tasks

Step 1: Train a Very Deep/Huge model



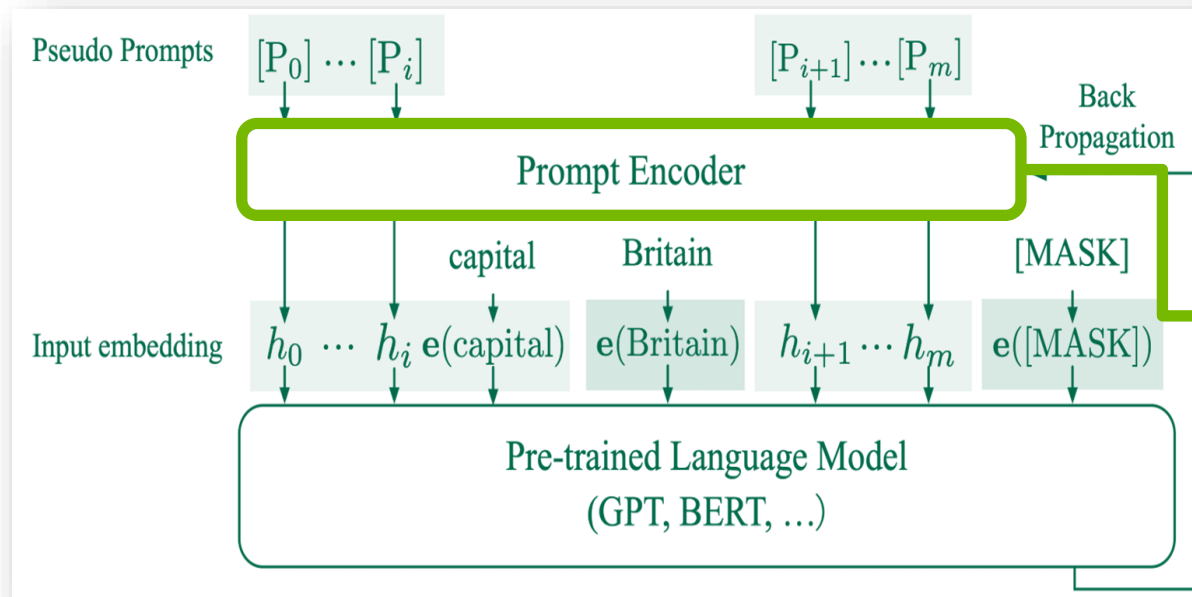
Step 2. Query different tasks with a few labelled data



Huge means Billions of parameters

P-TUNING METHOD

Vanilla P-tuning

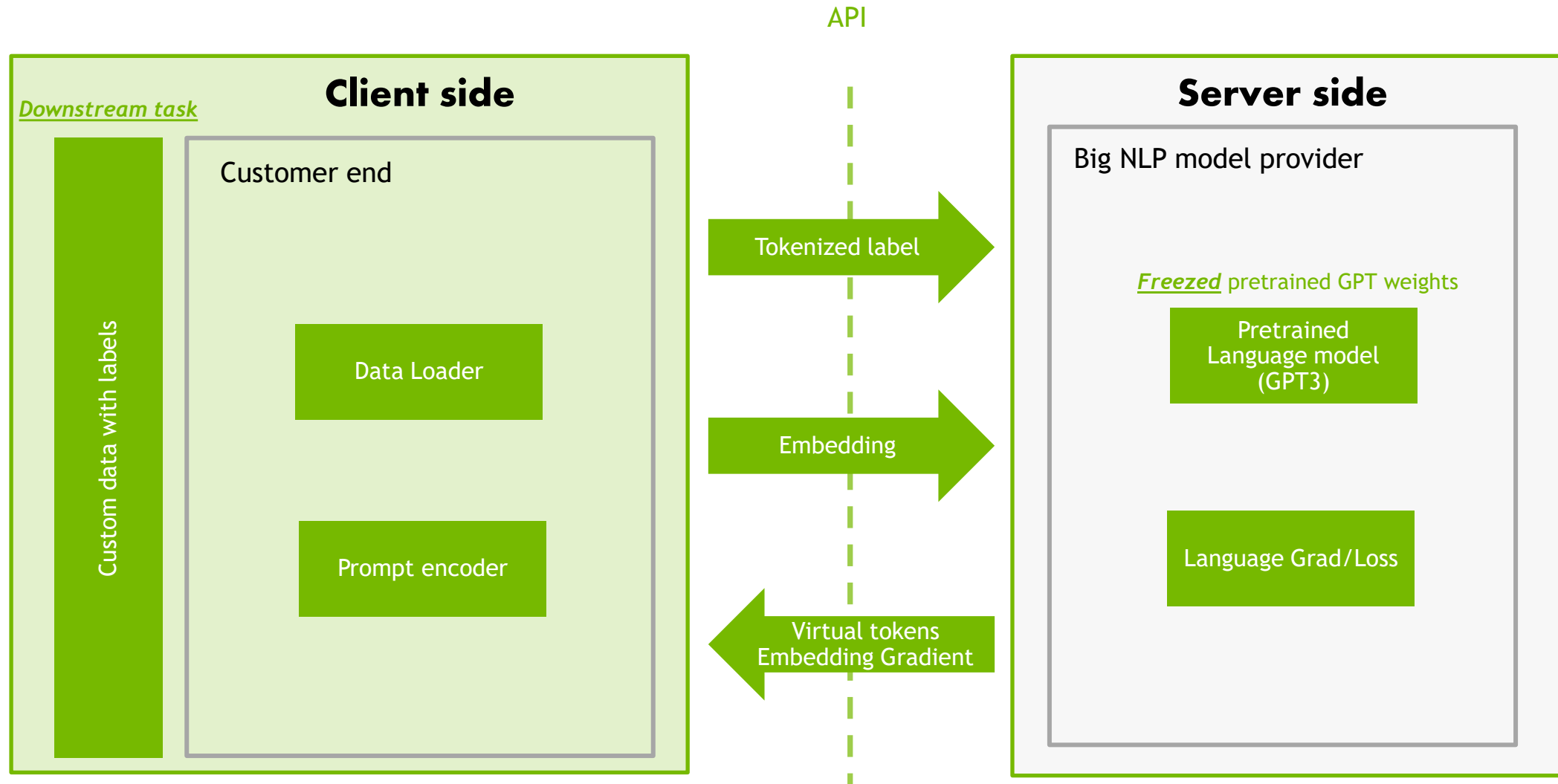


	Name	Type	Params
0	model	MegatronGPTModel	3.6 B
1	word_embeddings	VocabParallelEmbedding	154 M
2	prompt_table	PromptTable	0
3	prompt_encoder	PromptEncoder	132 M
132 M	Trainable params		
3.6 B	Non-trainable params		
3.7 B	Total params		
7,383.245	Total estimated model params size (MB)		

Liu, Zheng et al, GPT Understands, Too 2020

BRING YOUR OWN DATA VIA P-TUNING

Well-engineered architecture in NeMo (future release in NeMo Megatron)



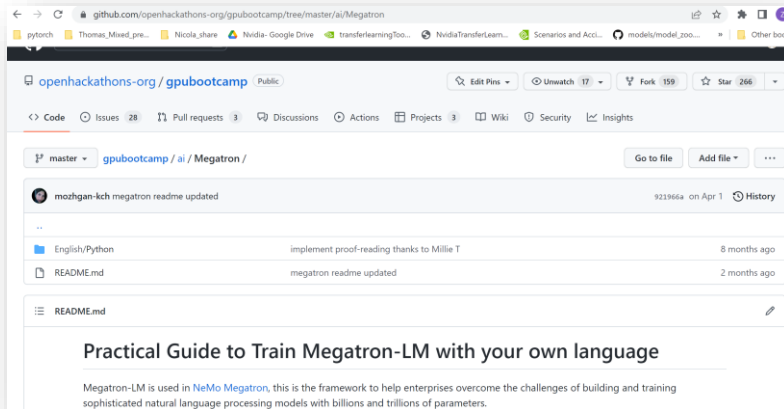


BUT HOW DO I TRY-OUT/KICK-START FOR REAL ?

MULTIPLE WAYS TO TRY-OUT/KICK-START

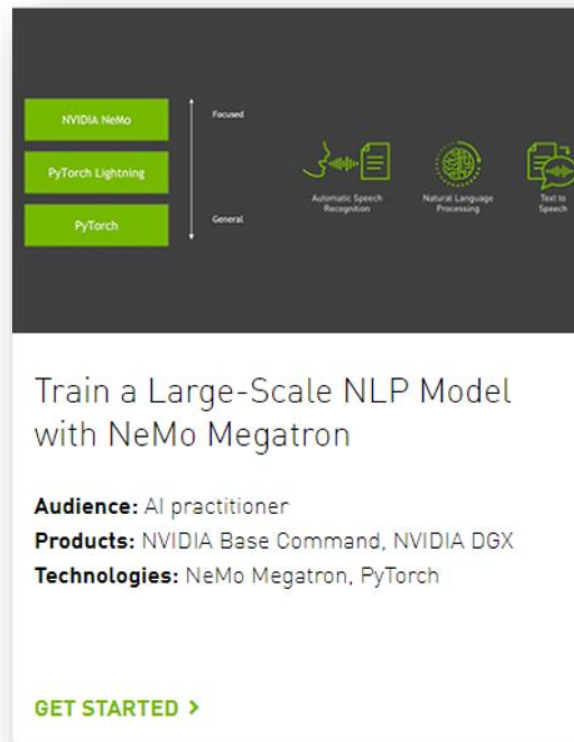
Contact NVIDIA

Request a bootcamp



<https://github.com/openhackathons-org/gpubootcamp/tree/master/ai/Megatron>

Register LaunchPad



<https://www.nvidia.com/en-us/data-center/launchpad/>

Join NeMo Megatron EA

NeMo Megatron Early Access

NeMo Megatron is a new capability in the NeMo framework that allows developers to effectively train and scale language models to billions of parameters. With NeMo Megatron, you can train different variants of GPT-3 models and scale them to multiple nodes on superpods. This deep learning software stack is optimized for NVIDIA DGX A100 SuperPODs using NVIDIA InfiniBand to provide efficient on-premises compute for training and inferring complex workloads.

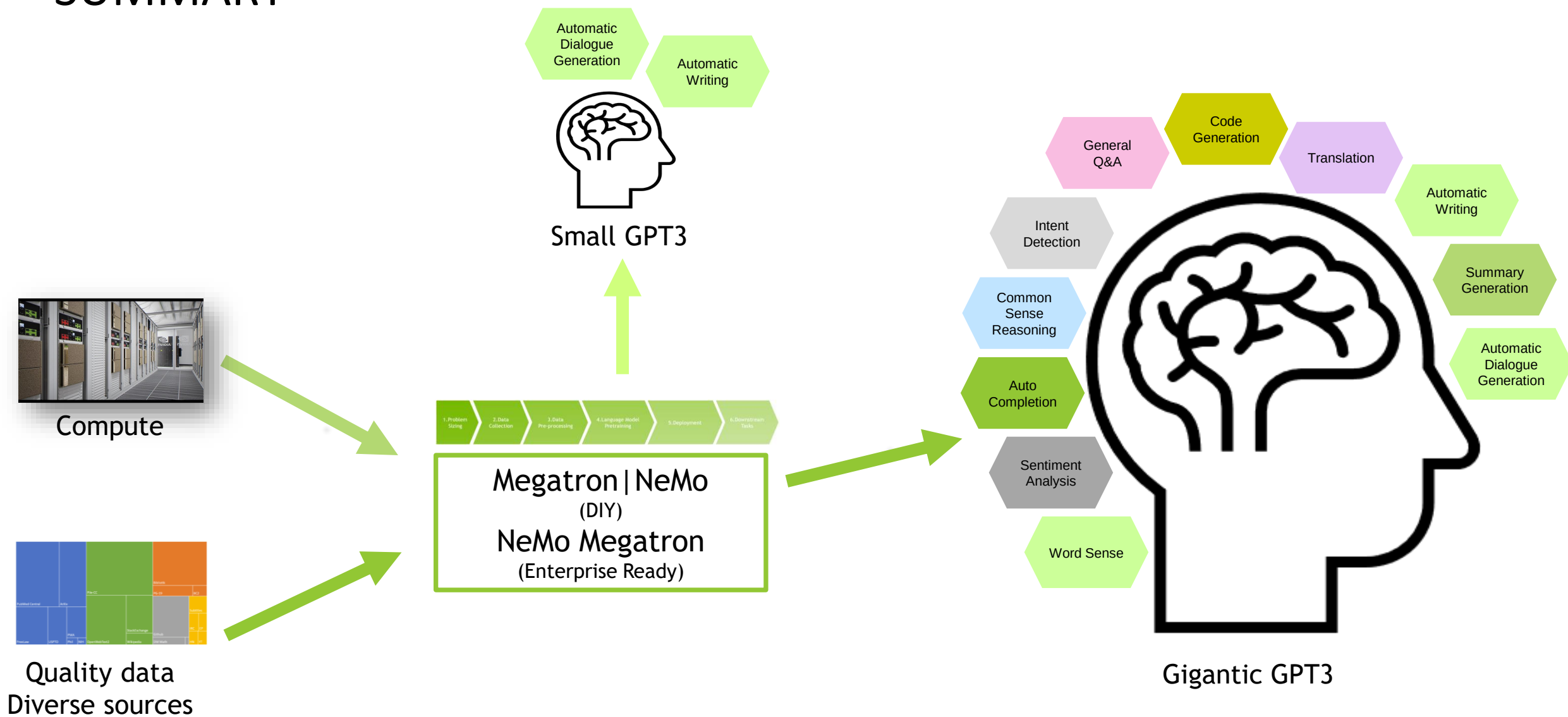
Early access to NeMo Megatron is limited to enterprises that want to train and deploy GPT-3 style models on NVIDIA A100 SuperPOD to perform zero-shot tasks such as answering deep domain questions, translating languages, comprehending and summarizing complex documents, etc.

<https://developer.nvidia.com/nemo-megatron-early-access>



SUMMARY

SUMMARY



REACH OUT



Rasmus Bisgaard
rbisgaard@nvidia.com



Zenodia Charpy
zcharpy@nvidia.com

DON'T FORGET TO **RATE THE SESSIONS**

#GOTOaar

Rate a minimum of **5 sessions** and
claim your **reward** at the
Registration Desk at the Trifork Hall