

Large language models

Inference, tokenization and small models



ChatGPT

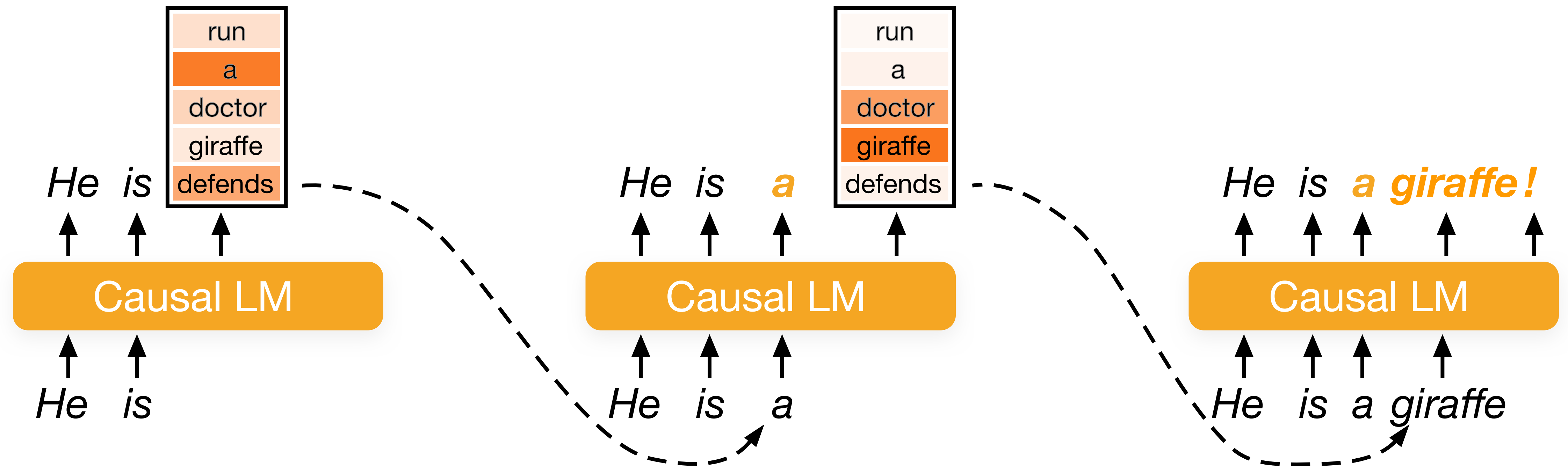
Agenda

Tokenization

Language-specific models

Inference: how to run an LLM efficiently

Generating text with LMs



Parts of a language models

'Heads' of a language model

How a model predicts the next word

Attention mechanism

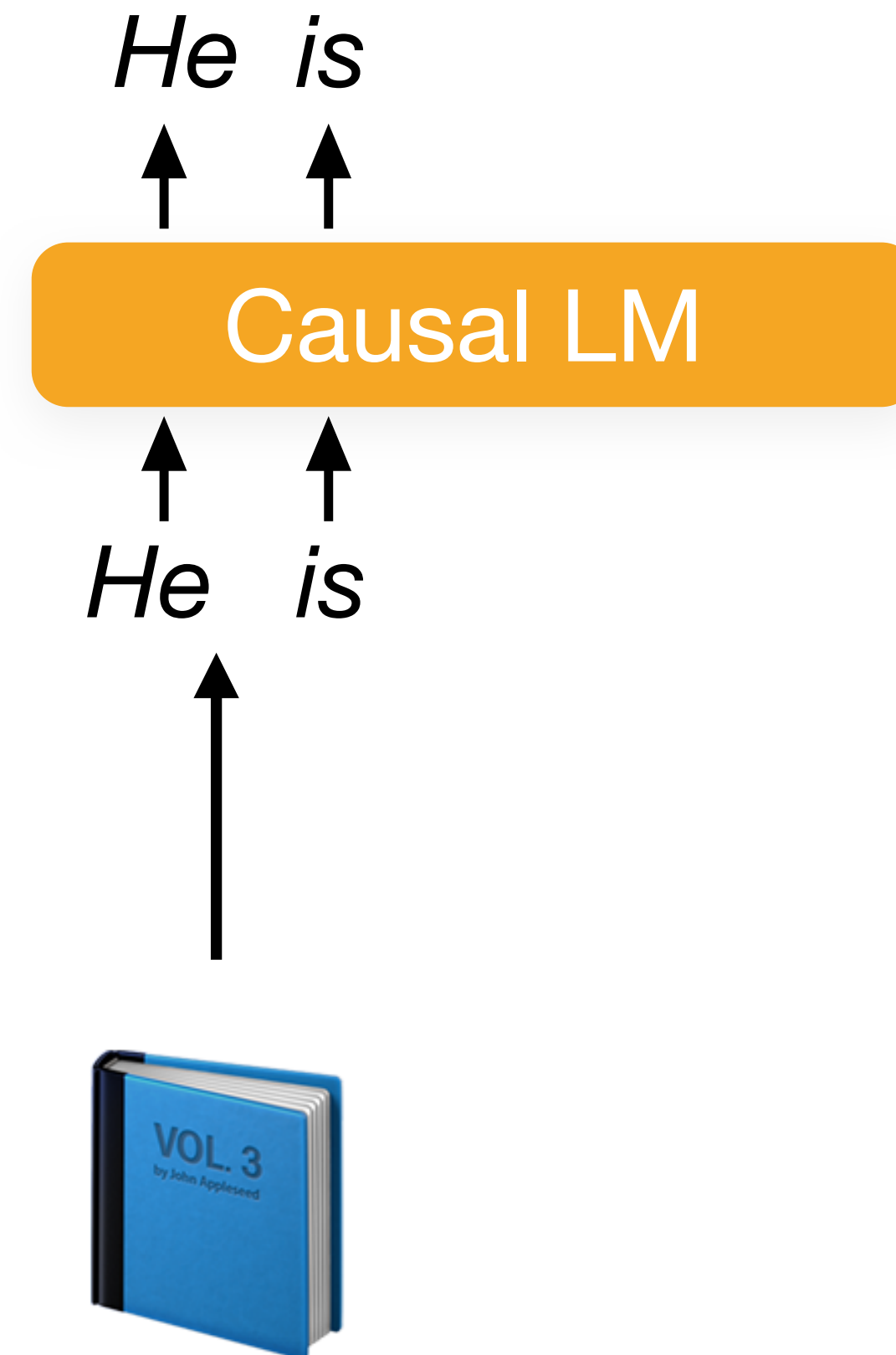
Each word affects the other words

Tokenizer

How a model understands text

Training data

What a model learns



Tokenizing the training data

an example

No, I am not a giraffe.

Tokenizing the training data

an example

No, I am not a giraffe.



No, I am not a giraffe.

Tokenizing the training data

an example

No, I am not a giraffe.



No, I am not a giraffe.



[2822, 11, 358, 1097, 539, 264, 37370, 21223, 13]

Byte-pair encoding (BPE) merges frequent tuples

Ġ s t a n d a a r d t a r i e f

Byte-pair encoding (BPE) merges frequent tuples

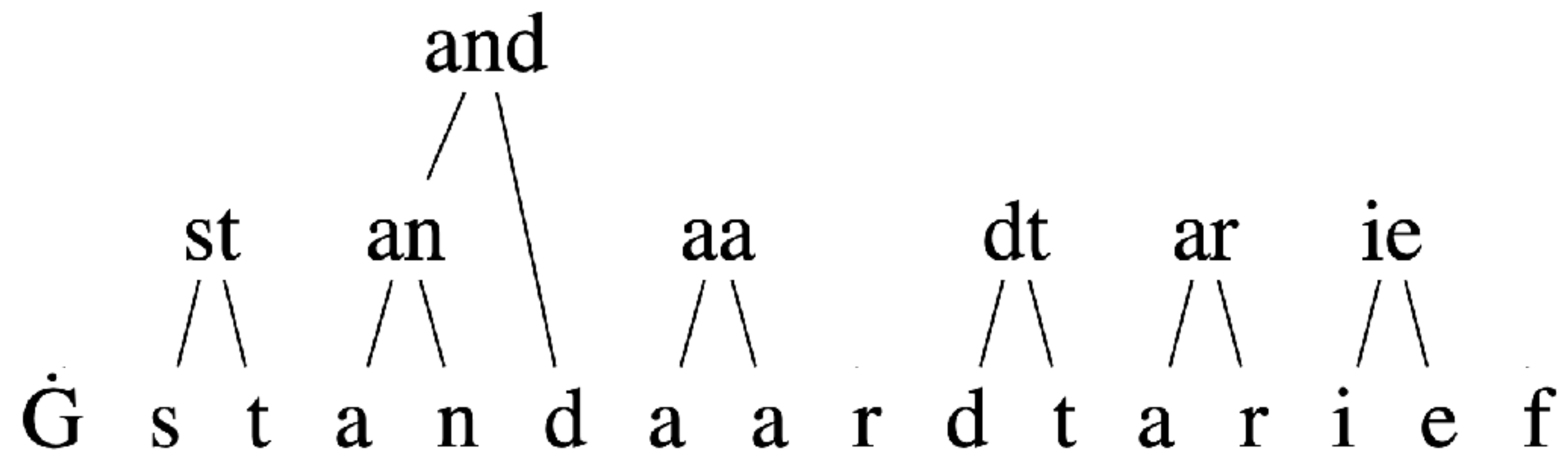
an
/\

Ġ s t a n d a a r d t a r i e f

Byte-pair encoding (BPE) merges frequent tuples

st an aa dt ar ie
/ \ / \ / \ / \ / \ / \
G s t a n d a a r d t a r i e f

Byte-pair encoding (BPE) merges frequent tuples



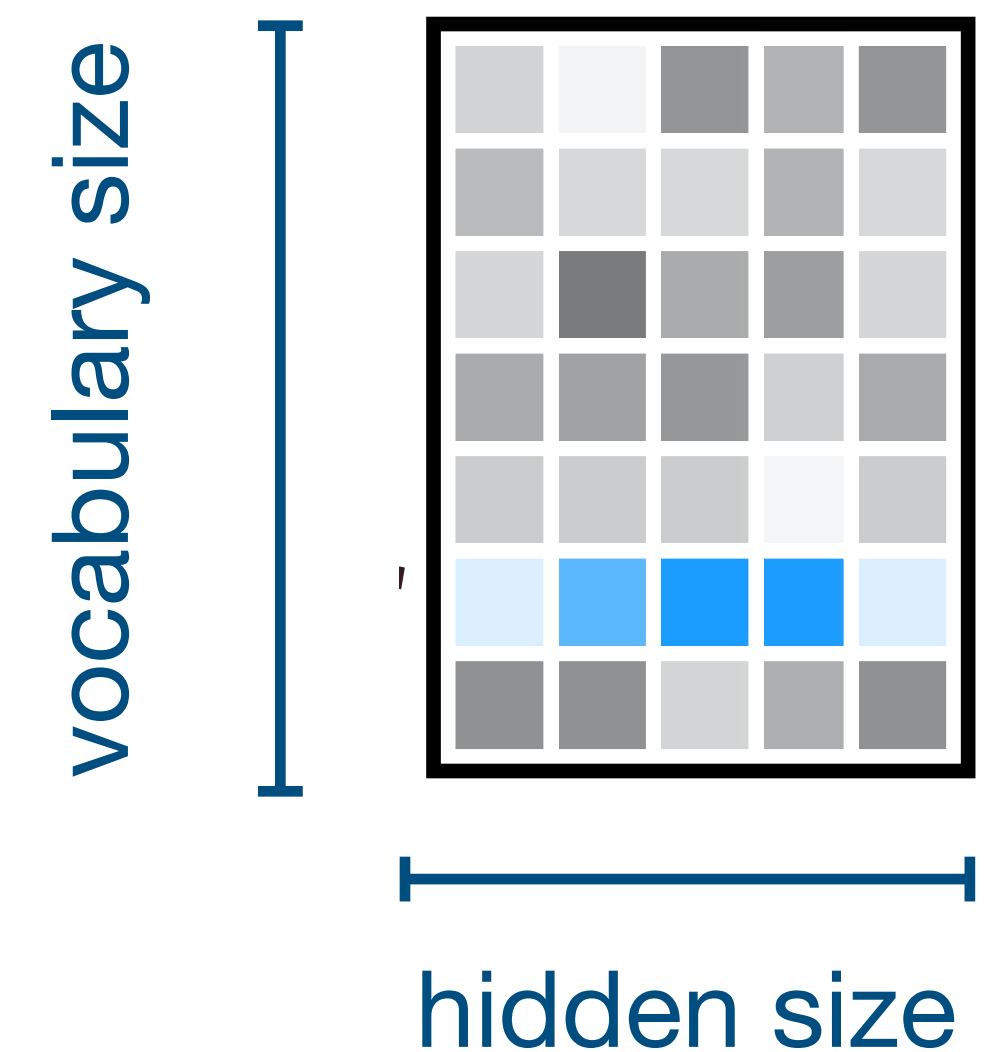
Few non-English words are tokens

Token types for words in English do not match, so the tokenizer falls back to non-representative tokens types.

Language-specific models

Why?

- Vocabulary size is limited
 - Not every word in every language can be a token
 - Embeddings need to be stored (16bits x vocab x h)



Inference

