

## Large Language Models

### Doelgroep cursus Large Language Models

De cursus Large Language Models is bedoeld voor ontwikkelaars, data scientists en technisch professionals die Large Language Models (LLMs) willen begrijpen en toepassen.

### Voorkennis cursus Large Language Models

Voor deelname aan de cursus is basiskennis van Python en machine learning vereist. Bekendheid met neurale netwerken natural language processing is aanbevolen.

### Realisatie training Large Language Models

De cursus wordt gegeven door deskundige trainer en bestaat uit een afwisseling van theorie en praktijk. Demonstraties en case studies met LLMs verduidelijken de leerstof.

### Certificaat Large Language Models

Na succesvolle afronding van de cursus ontvangen deelnemers een certificaat van deelname aan Large Language Models.

Duur: 2 dagen

Prijs: € 1499

[Open Rooster](#)



Large Language Models



## Inhoud Large Language Models

De cursus Large Language Models van SpiralTrain biedt een diepgaand overzicht van grote taalmodellen (LLMs), van transformer architecturen tot geavanceerde onderwerpen zoals fine-tuning, veiligheid, deployment en praktische toepassingen. Je leert werken met open source modellen, tools en innovatieve AI-oplossingen.

### Introductie tot LLMs

Deze module geeft een overzicht van LLM's, met uitleg over transformers, attention en tokenisatie. Modellen als GPT, BERT en T5 worden besproken, net als schaalwetten, trainingsdoelen en het verschil tussen pretraining en fine-tuning. Open source vs gesloten modellen komt ook aan bod.

### Modelarchitecturen

Deelnemers leren over decoders, encoder-decoders en modellen zoals LLaMA en PaLM. Er wordt ingegaan op trainingspipelines, precisie (FP16, quant), en tools zoals Hugging Face en Deepspeed. LoRA, QLoRA en instructietuning worden geïntroduceerd.

### LLMs Trainen

Deze module behandelt het voorbereiden van data en het fine-tunen van modellen. Onderwerpen zijn tokenizer setup, SFT, adapters, overfitting voorkomen en evaluatie. Benchmarking en alignment worden toegelicht met Hugging Face tools.

### LLM Deployen

Deelnemers leren hoe LLMs in productie worden gebracht. Inference-optimalisatie, quantisatie, distillatie en cloudhosting komen aan bod. Ook wordt gekeken naar LangChain, vector stores, caching en kostenbeheersing bij schaalvergroting.

### Veiligheid en Bias

Focus ligt op het herkennen en beperken van bias in modellen. Er is aandacht voor auditing, filtering, privacy, uitlegbaarheid, en juridische aspecten. Ook adversarial prompts, red teaming en ethiek worden behandeld.

### Toepassingen van LLMs

De cursus sluit af met toepassingen in o.a. softwareontwikkeling, recht, zorg en onderwijs. Voorbeelden zijn RAG-systemen, AI-agents en plugins. Ook komen bedrijfsintegratie, modelvergelijking en onderzoekstrends aan bod.

## Modules Large Language Models

| Module 1: Intro to LLMs                                                                                                                                                                                                                                                                   | Module 2: Model Architectures                                                                                                                                                                                                                                                                                               | Module 3: Training LLMs                                                                                                                                                                                                                                                    |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| What are LLMs?<br>Transformer architecture<br>Training Objectives (causal, masked)<br>Evolution of LLMs (GPT, BERT, T5)<br>Open Source vs Proprietary LLMs<br>Tokenization and Vocabulary<br>Attention Mechanism<br>Model Scaling Laws<br>Transfer Learning<br>Pretraining vs Fine-Tuning | Decoder vs Encoder-Decoder Models<br>GPT, LLaMA, T5, and PaLM<br>Training Pipeline Overview<br>Optimizers (Adam, Adafactor)<br>Precision (FP32, FP16, quantization)<br>Transformers (HF), Megatron, Deepspeed<br>Parameter vs Instruction Suning<br>LoRA and QLoRA<br>In-context Learning<br>Reinforcement Learning with HF | Dataset Creation and Curation<br>Tokenizer Customization<br>Data Preprocessing<br>Fine-Tuning with Hugging Face<br>SFT (Supervised Fine-Tuning)<br>Adapters and LoRA<br>Evaluation Metrics<br>Avoiding Overfitting<br>Model Alignment<br>Model Evaluation and Benchmarking |
| Module 4: LLM Deployment                                                                                                                                                                                                                                                                  | Module 5: Safety and Bias                                                                                                                                                                                                                                                                                                   | Module 6: LLM Use Cases                                                                                                                                                                                                                                                    |
| Inference Optimization<br>Model Distillation<br>Quantization Techniques<br>Hosting on AWS, GCP, Azure<br>Using Model Gateways<br>LangChain and Semantic Search<br>Vector Stores and Embeddings<br>Caching Responses<br>Load Balancing<br>Cost Optimization Strategies                     | Understanding Model Biases<br>Mitigation Strategies<br>Model Auditing<br>Adversarial Prompts<br>User Privacy<br>Filtering and Moderation<br>Red Teaming<br>Explainability in LLMs<br>Interpreting Outputs<br>Regulatory and Legal Issues                                                                                    | Coding Assistants<br>AI for Legal and Finance<br>Education and Learning<br>Health Care and Biotech<br>Chatbots and Agents<br>RAG Systems<br>Tool Use and Plugins<br>Enterprise Use of LLMs<br>Evaluating New Models<br>Future Directions LLM Research                      |