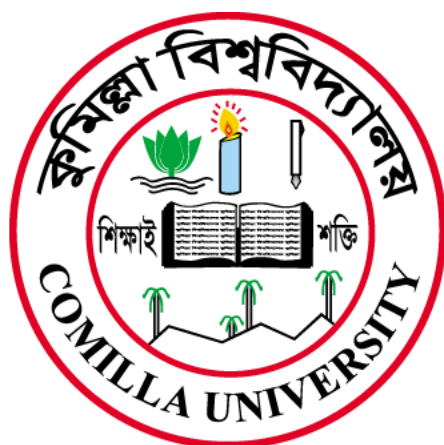


A Project Report submitted in partial fulfillment of the requirements for the degree of Masters of Science (Engineering) in **Information and Communication Technology**



**Fine-Tuning Large Language Models to generate websites
in specific tech stack**

Submitted by:

ID: 22209003

Session: 2021-22

Department of ICT

Comilla University

Supervisor:

Amena Begum

Assistant Professor

Department of ICT

Comilla University

Submission: February, 2025

Abstract

The rapid advancement of pre-trained Large Language Models (LLMs) has unlocked significant potential in automating domain-specific tasks. Despite their widespread adoption, fine-tuning LLMs for specialized applications remains an underexplored area. This study examines strategies for fine-tuning LLMs, emphasizing methodologies for adapting foundational models to specific domains such as landing page generation. Key considerations include dataset curation, preprocessing, model architecture selection, and efficient fine-tuning techniques like Low-Rank Adaptation (LoRA).

The project focuses on designing an AI-powered web application that enables users to create customizable landing pages. By integrating pre-trained LLMs with modular design components, the system generates personalized content based on user prompts. It outlines the workflow for building this system, including the selection of layouts, themes, and components, along with AI-driven text generation for different sections of a landing page. The study also highlights the challenges encountered during fine-tuning, such as maintaining contextual relevance and optimizing computational efficiency.

This work demonstrates the practical application of LLMs in web development, showcasing their potential to streamline content creation while meeting user-specific requirements. The findings contribute to advancing generative AI solutions and underscore the transformative role of LLM fine-tuning in modern web design.

Keywords: LLM Fine-Tuning, Generative AI, Domain-Specific LLM, LoRA, Landing Page Generator.

Table of Contents

Abstract.....	1
Table of Contents.....	2
Chapter 1: Introduction.....	5
1.1 Introduction:.....	5
1.2 Motivation:.....	6
1.3 Project Outline:.....	7
Chapter 2: Large Language Models and Fine-Tuning.....	8
2.1. Overview of Large Language Models (LLMs).....	8
2.2 Foundation Models and LLM Pre-training.....	9
2.3 LLM Fine-Tuning.....	10
2.3.1 LoRA: A Prominent PEFT Technique.....	10
2.3.2 Role of Pre-training Data in Fine-Tuning.....	11
2.3.3 Fine-Tuning Techniques.....	11
2.3.4 Evaluation and Iterative Improvement.....	12
Chapter 3: Fine Tuning LLMs for Web Development.....	13
3.1 Methodology.....	13
3.2 Data Preparation.....	14
3.2.1 Data Collection.....	14
A. Sources of Training Data.....	14
3.2.2 Data Preprocessing.....	14
A. Cleaning & Normalization.....	14
B. Structuring Data for Fine-Tuning.....	15
3.2.3 Data Augmentation.....	15
A. Paraphrasing & Style Variation.....	15
B. Industry-Specific Customization.....	15
C. Length & Structure Variation.....	15
3.2.4 Dataset Splitting.....	16
3.2.5 Considerations for High-Quality Training Data.....	16
3.3 Model Selection.....	16
3.3.1 Criteria for Model Selection.....	17
3.3.2 Chosen Model: LLaMA-3.....	17
3.3.3 Fine-Tuning Approach.....	18
A. LoRA: Reducing Fine-Tuning Costs.....	18
B. QLoRA: Maximizing Efficiency.....	18
3.3.4 Justification for LoRA & QLoRA.....	18
3.3.5 Model Configuration for Landing Page Generation.....	18
3.4 Training Process.....	19
3.4.1 Training Setup & Environment.....	19
A. Hardware & Software Configuration.....	19
B. Hyperparameter Configuration.....	19

3.4.2 Fine-Tuning Strategy.....	20
A. Implementing LoRA for Efficient Fine-Tuning.....	20
B. Enhancing Efficiency with QLoRA.....	20
C. Training Dataset.....	20
3.5 Evaluation Metrics.....	20
3.5.1 Quantitative Metrics.....	20
A. Perplexity (PPL).....	20
B. Cross-Entropy Loss.....	21
C. Context Relevance Score.....	21
D. Completeness & Coherence Score.....	21
3.5.2 Qualitative Assessment.....	22
3.6 Optimization & Deployment.....	22
3.6.1 Optimization Strategies.....	22
A. Quantization for Memory Efficiency.....	22
B. Prompt Engineering for Enhanced Outputs.....	22
3.6.2 Deployment with Ollama (Local Inference).....	23
A. Why Ollama?.....	23
B. Ollama Deployment Setup.....	23
Chapter 4: Applications of LLMs in Landing Page Generation.....	24
4.1 AI-Powered Content Generation.....	24
4.1.1 Automating Landing Page Copywriting.....	24
4.1.2 Structuring Landing Page Sections with AI.....	24
4.2 Personalization & Adaptive Content Generation.....	25
4.2.1 AI-Generated Content Tailored to User Needs.....	25
4.2.2 Dynamic Adjustments Based on User Feedback.....	25
4.3 SEO Optimization with AI-Generated Landing Pages.....	25
4.3.1 Enhancing Search Engine Visibility.....	25
4.3.2 Automated Internal Linking & Readability Enhancements.....	26
4.4 AI in UI/UX Optimization for Landing Pages.....	26
4.4.1 Structuring Layouts with AI-Driven Recommendations.....	26
4.4.2 AI-Generated A/B Testing Variations.....	26
4.5 AI-Driven Content Updates & Localization.....	26
4.5.1 Real-Time Content Updates.....	26
4.5.2 Localization & Multi-Language Adaptation.....	26
4.6 Project Showcase (PageGenAI).....	27
4.6.1 Home Page.....	27
4.6.2 Generate Page.....	28
4.6.3 Form for Instructions for the LLM.....	28
4.6.4 Generated Output of the Fine-tuned LLM.....	29
Chapter 5: Conclusion and Future Directions.....	30
5.1 Conclusion.....	30

5.2 Challenges and Limitations.....	30
5.2.1 Computational Constraints.....	30
5.2.2 Contextual Understanding & Hallucination Issues.....	30
5.2.3 Limited Multi-Modal Capabilities.....	31
5.3 Future Directions.....	31
5.3.1 Expanding Multi-Modal Capabilities.....	31
5.3.2 Enhancing AI Model Performance.....	31
5.3.3 Real-Time AI Content Editing.....	31
5.3.4 Deployment Scalability & API Integration.....	31
5.4 Final Thoughts.....	32
References:.....	32

Chapter 1: Introduction

1.1 Introduction:

The emergence of pre-trained Large Language Models (LLMs) has revolutionized natural language processing (NLP) by enabling advanced applications across diverse fields. These models, such as GPT and LLaMA, are pre-trained on massive datasets and possess the ability to generate coherent and contextually relevant text. However, their general-purpose nature often limits their applicability in domain-specific tasks, necessitating fine-tuning to adapt these models to specialized use cases. Fine-tuning allows LLMs to leverage their pre-trained knowledge while tailoring their performance to meet the specific requirements of targeted domains or industries.

This project explores the application of fine-tuned LLMs in automating web development, specifically for generating dynamic landing pages. With the increasing demand for visually appealing and functional websites, there is a growing need for tools that combine modular design with intelligent content generation. By integrating LLMs with pre-designed templates and components, this study demonstrates how these models can streamline the creation of personalized, user-friendly landing pages.

The proposed system allows users to select layouts, themes, and modular components for various sections of a landing page, such as the hero, testimonials, FAQs, and more. It further employs fine-tuned LLMs to generate contextually relevant text based on user prompts. Techniques like Low-Rank Adaptation (LoRA) are used to make fine-tuning more efficient, reducing computational overhead while maintaining high performance.

This study highlights the practical considerations of fine-tuning, including dataset preparation, optimization, and validation. The project also addresses challenges such as maintaining coherence in generated text and ensuring the scalability of the system for real-time applications. By advancing the integration of LLMs in web development, this work demonstrates the transformative potential of generative AI and contributes valuable insights for future research in domain-specific AI applications.

1.2 Motivation:

The increasing demand for personalized and visually appealing web applications has created a significant need for tools that streamline the development process. Traditional approaches to building websites, such as manual coding or using drag-and-drop builders, often fall short in terms of scalability, efficiency, and adaptability. At the same time, the rise of pre-trained Large Language Models (LLMs) has demonstrated their immense potential in automating complex tasks by generating human-like and contextually relevant text.

However, leveraging these powerful models for specialized tasks, such as dynamic landing page generation, requires fine-tuning to align them with the specific requirements of the domain. Fine-tuning ensures that LLMs can generate accurate, domain-specific content, bridging the gap between general-purpose capabilities and specialized applications.

This project is motivated by the need to harness the potential of LLMs to simplify and enhance the process of creating customizable landing pages. By integrating modular design components, AI-driven content generation, and efficient fine-tuning techniques, the project aims to provide a solution that is both user-friendly and technologically advanced. The work also seeks to explore and address challenges associated with fine-tuning, such as maintaining coherence, optimizing computational efficiency, and ensuring scalability for real-time usage.

1.3 Project Outline:

This paper is organized as follows:

- **Chapter 2: Large Language Models and Fine-Tuning**

This chapter provides an overview of LLMs and their applications, discussing the significance of fine-tuning for adapting these models to domain-specific tasks.

- **Chapter 3: Fine Tuning LLMs for Web Development**

This chapter details the methods and approaches used to create LLMs tailored for dynamic landing page generation. It covers key aspects such as dataset preparation, model selection, preprocessing, and fine-tuning techniques like Low-Rank Adaptation (LoRA).

- **Chapter 4: Applications of LLMs in Landing Page Generation**

This chapter explores the practical implementation of fine-tuned LLMs in creating customizable landing pages. It examines workflows, challenges, and outcomes, showcasing the role of AI in transforming the web development process.

- **Chapter 5: Conclusion and Future Directions**

The final chapter summarizes the key findings of this project, highlights the limitations encountered, and discusses potential future advancements in using LLMs for automated web application development.

Chapter 2: Large Language Models and Fine-Tuning

2.1. Overview of Large Language Models (LLMs)

In recent years, Large Language Models (LLMs) have achieved remarkable advancements in the field of Natural Language Processing (NLP). Models such as BERT (Bidirectional Encoder Representations from Transformers) and GPT (Generative Pre-trained Transformer), built on the Transformer architecture, have demonstrated their ability to learn powerful representations from extensive pre-training on large-scale datasets. This enables their application to various NLP tasks by effectively understanding text context, grammar structures, and semantic relationships.

Generative AI, a subset of artificial intelligence, leverages models trained on vast datasets to create new content such as text, images, audio, and video. In NLP, the development of Natural Language Understanding (NLU) technologies has enabled complex tasks like dialogue processing, supported by the use of contextual models and Transformer-based language models. Additionally, integrating AI-powered solutions, such as chatbots with Robotic Process Automation (RPA) and Optical Character Recognition (OCR), has significantly improved task efficiency across various domains.

OpenAI's release of ChatGPT in November 2022 marked a notable advancement in conversational AI. Built on the GPT-3.5 architecture, ChatGPT utilizes reinforcement learning with human feedback to enhance its conversational capabilities. Its adoption soared rapidly, reaching over 100 million monthly users within two months of launch. This highlights the growing importance of LLMs in applications requiring natural language understanding and generation.

Interfacing with LLMs relies on constructing effective prompts, making prompt engineering an essential process for obtaining relevant and coherent outputs. While fine-tuning LLMs can tailor them to specific domains, traditional fine-tuning can be computationally expensive and resource-intensive, especially for large models. To address these challenges, parameter-efficient fine-tuning techniques such as Low-Rank Adaptation (LoRA) and Quantized LoRA (QLoRA) have emerged as practical alternatives.

- **LoRA:** LoRA modifies LLMs by freezing the original model weights and introducing low-rank matrices that adaptively learn domain-specific knowledge. These additional matrices require significantly fewer trainable parameters compared to full fine-tuning, making the process more efficient.

- **QLoRA:** QLoRA builds on LoRA by using quantization techniques to reduce memory requirements further, allowing large models to be fine-tuned on resource-constrained hardware without sacrificing performance.

Both LoRA and QLoRA significantly reduce the computational overhead associated with traditional fine-tuning methods, making them ideal for domain-specific applications like dynamic landing page generation. By leveraging these advancements, LLMs can be efficiently adapted to meet specialized requirements, enabling their deployment across diverse domains.

2.2 Foundation Models and LLM Pre-training

Generative AI models are designed to create outputs tailored to specific content types, including text, images, audio, and video. Recently, multi-modal models, which can simultaneously process data from different formats such as images and text, have emerged as advanced foundation models. These models utilize diverse datasets encompassing various modalities, including text, structured data, audio, and even 3D signals, during their training. Foundation models represent a paradigm shift in artificial intelligence, showcasing capabilities akin to human reasoning and creativity. Their adaptability allows them to be fine-tuned for specific applications, completing the training process to meet specific user objectives.

Among foundation models, Large Language Models (LLMs) are pre-trained language models (PLMs) that leverage vast amounts of textual data from general-purpose sources such as websites, books, and articles. The foundational training process relies on unsupervised or semi-supervised learning, enabling LLMs to develop an extensive understanding of text context, semantics, and grammar. This broad knowledge base allows them to perform well in diverse downstream tasks like summarization, question answering, or content generation when fine-tuned appropriately.

The creation of an LLM begins with gathering a comprehensive training dataset, which is the core resource for the model's learning. Training data for LLMs typically falls into two broad categories:

- 1. General Data:**

General-purpose data, including web pages, books, and conversational texts, is abundant and diverse. It enhances the model's language understanding, generalization capabilities, and versatility, making it the primary source for pre-training most LLMs.

- 2. Specialized Data:**

To develop task-specific expertise, LLMs can also be trained on domain-specific datasets. Examples include multilingual text, scientific

literature, legal documents, or programming code. These specialized datasets expand the LLM's abilities, equipping it to solve specific problems effectively. For instance, LLMs fine-tuned on programming datasets perform exceptionally well in tasks like code generation and debugging.

Pre-training LLMs involves ingesting massive amounts of text data and learning patterns in an unsupervised manner. This enables the models to predict text context and semantic relationships efficiently. Techniques such as masking and token prediction are commonly used in training models like BERT and GPT to enhance their text comprehension abilities.

Incorporating specialized datasets during the pre-training or fine-tuning phase allows LLMs to cater to specific use cases, such as generating landing page content in this project. Such approaches align the model's capabilities with user requirements, making them highly effective for tasks that demand domain-specific expertise. As these foundation models evolve, their ability to adapt to diverse datasets continues to transform industries, enabling the development of efficient and scalable AI-driven applications.

2.3 LLM Fine-Tuning

While Large Language Models (LLMs) are pre-trained on massive datasets and exhibit general-purpose capabilities, they may not perform optimally on specialized tasks without fine-tuning. Fine-tuning enhances the performance of these models by adapting them to specific domains or tasks, making them more precise and context-aware. Traditionally, smaller models like BERT (450M parameters) and RoBERTa (1.3B parameters) were fully fine-tuned by updating all model parameters. However, as LLMs grew in scale, such as LLaMA, fine-tuning their entirety became computationally infeasible.

To address this challenge, **parameter-efficient fine-tuning (PEFT)** techniques have gained traction. PEFT methods, such as **Low-Rank Adaptation (LoRA)**, focus on adding a small number of trainable parameters to pre-trained LLMs, while keeping most of the original model parameters frozen. This approach reduces computational costs while maintaining competitive performance, making fine-tuning more accessible for large-scale models.

2.3.1 LoRA: A Prominent PEFT Technique

LoRA is a leading PEFT method that introduces trainable low-rank matrices for fine-tuning specific weight matrices within Transformer layers. Rather than updating the entire model, LoRA identifies target weight matrices within each Transformer

layer and applies **low-rank matrix decomposition**. This process involves breaking down the weight matrix W into two smaller matrices:

$$\Delta W = A \cdot B$$

- **Matrix A**: Initialized with random Gaussian values to introduce variability.
- **Matrix B**: Initialized to zeros to ensure the changes begin incrementally.

These matrices are fine-tuned while keeping the rest of the model frozen, allowing for domain-specific adaptation without the need to retrain the entire model. This selective fine-tuning ensures that the model efficiently learns task-specific knowledge while preserving its pre-trained general-purpose capabilities.

Advantages of LoRA:

- **Efficiency**: LoRA significantly reduces the number of trainable parameters compared to full fine-tuning, making it cost-effective and scalable.
- **Flexibility**: By focusing on specific weight matrices, LoRA can adapt models to diverse tasks without requiring extensive computational resources.
- **Performance**: Despite its parameter-efficient nature, LoRA achieves comparable results to traditional fine-tuning approaches.

2.3.2 Role of Pre-training Data in Fine-Tuning

The quality and diversity of the pre-training data play a crucial role in determining the effectiveness of fine-tuning. LLMs are initially pre-trained on vast datasets collected from sources such as web pages, books, and articles to learn linguistic structures and relationships. The performance of these models heavily depends on the size, quality, and domain relevance of the pre-training corpus. When specialized tasks are required, high-quality domain-specific datasets are essential for fine-tuning to achieve optimal results.

2.3.3 Fine-Tuning Techniques

After pre-training, the fine-tuning process involves adapting the model to specific tasks through various techniques, including:

- **In-Context Learning**: Utilizing few-shot, zero-shot, or one-shot learning to adapt the model without additional training.
- **Fine-Tuning on Specialized Datasets**: Adding domain-specific knowledge through additional training on targeted datasets.
- **Hyperparameter Tuning**: Adjusting parameters like learning rates, batch sizes, or optimizer settings to optimize the model's performance.

2.3.4 Evaluation and Iterative Improvement

Once fine-tuned, the model is evaluated on a test dataset that was not used during training. Metrics such as cross-entropy loss, perplexity, and relevance scores are used to assess performance. Based on the evaluation results, additional refinements may be made, such as:

- Adjusting hyperparameters for improved convergence.
- Modifying the architecture to better handle specific tasks.
- Incorporating additional domain-specific data to enhance adaptability.

The final output, referred to as the **Fine-Tuned Language Model (FLM)**, becomes a task-specific model optimized for its intended application. This approach enables LLMs to excel in specialized domains, providing enhanced performance and precision for a variety of applications, such as content generation for landing pages in this project.

Chapter 3: Fine Tuning LLMs for Web Development

3.1 Methodology

Fine-tuning a Large Language Model (LLM) for Nextjs landing page generation requires a structured approach, balancing efficiency with accuracy. Instead of training a model from scratch, we adapt an existing LLM using parameter-efficient fine-tuning (PEFT) techniques like LoRA and QLoRA. This allows us to optimize the model for generating structured, context-aware web content while keeping computational costs low.

The process starts with data collection and preprocessing, ensuring the model learns from high-quality landing page content. Next, the model is fine-tuned with structured prompts tailored to real-world use cases—helping it generate relevant headlines, descriptions, and calls to action. To evaluate performance, we apply metrics like cross-entropy loss, perplexity, and context relevance to ensure the generated content is accurate, fluent, and well-structured.

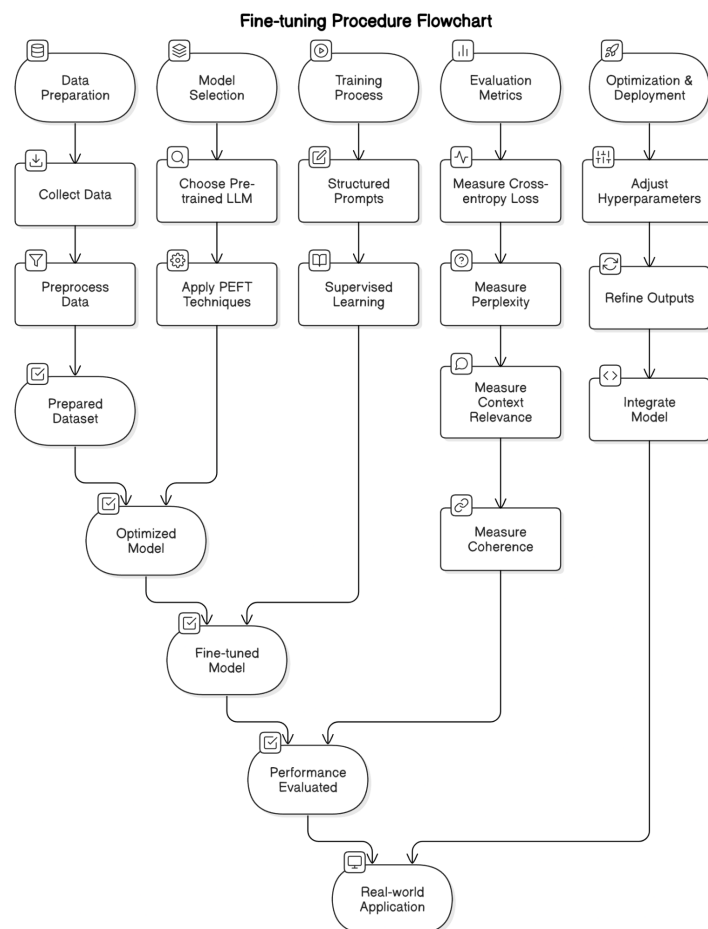


Figure 1: Flowchart of Fine tuning LLMs

3.2 Data Preparation

Fine-tuning a Large Language Model (LLM) for landing page generation starts with carefully curated data. High-quality training data ensures the model understands the structure, tone, and content patterns used in professional landing pages. This section details the key stages of data collection, preprocessing, formatting, and augmentation to optimize the model's learning process.

3.2.1 Data Collection

The effectiveness of fine-tuning depends heavily on the dataset used. Since our goal is to generate context-aware, structured landing page content, we need diverse and well-structured textual data from various sources.

A. Sources of Training Data

We compiled data from multiple sources to create a well-rounded dataset:

1. **Existing Landing Pages** – Scraped or manually collected data from well-designed, high-converting landing pages across industries (e.g., SaaS, e-commerce, personal portfolios, startups).
2. **Marketing Copy & UI Content** – Text from professionally written hero sections, CTA buttons, service descriptions, testimonials, and FAQs to capture different section styles.
3. **Public Datasets** – Open-source corpora containing structured web content and marketing text (e.g., Common Crawl, OpenWebText).
4. **Manually Curated Samples** – Human-written examples that define tone, length, and structure standards for different landing page sections.

By combining these sources, we ensure a balance of generality and specificity, helping the model learn both common writing patterns and customized, industry-specific styles.

3.2.2 Data Preprocessing

Once raw text data is collected, it must be **cleaned and structured** to be useful for fine-tuning. Unstructured, noisy, or irrelevant data can negatively impact the model's performance.

A. Cleaning & Normalization

To ensure consistency, we applied the following preprocessing steps:

- **Removing HTML & Boilerplate Text** – Stripped out unnecessary HTML tags, scripts, and CSS formatting.
- **Tokenization** – Split text into individual tokens (words, subwords) for efficient processing by the LLM.
- **Deduplication** – Identified and removed duplicate content to avoid overfitting.
- **Spelling & Grammar Correction** – Applied NLP-based text correction to remove errors from scraped content.
- **Filtering Irrelevant Content** – Eliminated unrelated sections like cookie banners, legal disclaimers, or navigation links.

B. Structuring Data for Fine-Tuning

Since landing pages are composed of **distinct sections**, we formatted the dataset into structured training pairs:

- **Input (Prompt)** → A structured instruction, such as:
"Generate a hero section for a SaaS landing page that focuses on AI-powered automation. Use a persuasive and engaging tone."
- **Output (Target Response)** → The expected content generated by the AI:
"Effortlessly automate your workflow with AI. Save time, reduce errors, and scale your business with our cutting-edge automation tools."

By keeping a clear input-output format, the model learns to map structured prompts to well-formed landing page content.

3.2.3 Data Augmentation

Since high-quality landing page data can be limited, we applied **data augmentation techniques** to increase dataset variety and improve model generalization.

A. Paraphrasing & Style Variation

- Used NLP models to **rephrase** existing content while keeping meaning intact.
- Introduced different **writing tones** (formal, casual, minimalist, bold) to improve adaptability.

B. Industry-Specific Customization

- Tailored training examples for different industries: **tech startups, healthcare, education, e-commerce, finance, etc.**
- Ensured diversity in **service descriptions, CTAs, and feature highlights.**

C. Length & Structure Variation

- Created **short-form, medium-length, and long-form versions** of each section to train the model on different content lengths.
- Introduced variations in **CTA styles** (e.g., direct action vs. benefit-driven copy).

These augmentation techniques **expanded the dataset** without compromising quality, allowing the model to **adapt to different writing styles and user needs**.

3.2.4 Dataset Splitting

To ensure a balanced evaluation, the dataset was split into three subsets:

1. **Training Set (80%)** – Used for fine-tuning, allowing the model to learn content structures and patterns.
2. **Validation Set (10%)** – Monitored during training to **prevent overfitting** and guide hyperparameter adjustments.
3. **Test Set (10%)** – Evaluated after fine-tuning to measure real-world performance and generalization ability.

This structured approach ensures that the fine-tuned model can generalize well across different landing page prompts, rather than memorizing specific examples.

3.2.5 Considerations for High-Quality Training Data

For fine-tuning to be successful, the dataset must meet key quality criteria:

Diversity – Covers multiple industries, styles, and tones to prevent bias.

Consistency – Ensures structured, well-formatted content across all examples.

Relevance – Focuses strictly on landing page content (not general web text).

Clarity & Coherence – Ensures that generated outputs are readable and natural.

By meticulously preparing the training data, we lay the foundation for a fine-tuned LLM that generates high-quality, structured, and user-specific landing page content.

3.3 Model Selection

Selecting the right Large Language Model (LLM) for fine-tuning is a crucial step in building an AI-powered landing page generator. The choice of model impacts performance, efficiency, and scalability, so we must balance factors like accuracy, computational cost, and adaptability to structured content generation.

This section outlines the criteria for selecting an LLM, the justification for the chosen model, and the fine-tuning approach used to align it with landing page generation.

3.3.1 Criteria for Model Selection

To ensure the LLM is well-suited for structured, **section-based text generation**, we evaluated models based on the following factors:

- **Pre-Trained Knowledge** – The model should have a strong understanding of marketing language, UX writing, and structured content.
- **Parameter Size & Efficiency** – A model with an optimal balance of performance vs. resource consumption, so fine-tuning is feasible even on limited hardware.
- **Adaptability to Fine-Tuning** – Must support Parameter-Efficient Fine-Tuning (PEFT) techniques like LoRA or QLoRA to minimize cost while maintaining performance.
- **Context Length & Prompt Handling** – Needs sufficient context window to process structured prompts and generate coherent, well-formatted responses.
- **Open-Source Availability** – Preference for open models over proprietary ones to allow for greater customization and transparency.

3.3.2 Chosen Model: LLaMA-3

After evaluating various models, we selected **LLaMA-3** as the foundation for fine-tuning. Here's why:

- **Scalability** – LLaMA-3 provides strong language generation capabilities without the extreme computational demands of larger models like GPT-4.
- **Parameter-Efficient Fine-Tuning (PEFT) Support** – It integrates well with LoRA and QLoRA, allowing for low-cost fine-tuning without updating the entire model.
- **Pre-Training Quality** – LLaMA-3 has been trained on a diverse, high-quality dataset, making it adaptable to structured web content like landing pages.
- **Open-Source Flexibility** – Unlike closed models (e.g., OpenAI's GPT-4), LLaMA-3 allows for custom fine-tuning and local deployment without API restrictions.

3.3.3 Fine-Tuning Approach

To efficiently adapt LLaMA-3 for landing page content generation, we used LoRA (Low-Rank Adaptation) and QLoRA (Quantized LoRA) as our fine-tuning strategies.

A. LoRA: Reducing Fine-Tuning Costs

LoRA enables us to freeze most of the model's parameters while fine-tuning only a small subset of trainable layers. This drastically reduces GPU memory usage while still allowing the model to learn landing page-specific patterns.

- LoRA injects low-rank matrices into the model's existing weight matrices.
- These matrices learn task-specific modifications while keeping the core model intact.
- This approach avoids catastrophic forgetting, preserving LLaMA-3's general language understanding.

B. QLoRA: Maximizing Efficiency

QLoRA further optimizes fine-tuning by quantizing the model, meaning it compresses large weight matrices into lower precision formats (e.g., 4-bit or 8-bit).

This results in:

- ✓ **Lower memory consumption**, allowing fine-tuning on standard consumer GPUs.
- ✓ **Minimal performance loss**, retaining fluency and coherence in generated content.

Using LoRA and QLoRA together, we ensure that fine-tuning is efficient, cost-effective, and scalable.

3.3.4 Justification for LoRA & QLoRA

- ♦ Full fine-tuning of LLaMA-3 would require massive computational resources (multiple GPUs with high VRAM).
- ♦ LoRA and QLoRA allow us to fine-tune only key parameters, making it possible to train even on mid-tier hardware without sacrificing performance.
- ♦ The model retains its general-purpose knowledge while adapting to landing page content, ensuring a balance between fluency and domain specificity.

3.3.5 Model Configuration for Landing Page Generation

Once fine-tuned, the LLaMA-3 model is optimized for structured, context-aware web content generation. It can:

- **Understand structured prompts** – Distinguishes between sections like hero text, CTAs, testimonials, and service descriptions.
- **Generate coherent, formatted outputs** – Produces text in a clean, scannable format suited for websites.
- **Adapt to different industries** – Can generate landing pages for tech startups, SaaS companies, personal brands, and more based on user prompts.

3.4 Training Process

Fine-tuning a Large Language Model (LLM) for landing page generation requires a structured training pipeline. The goal is to efficiently adapt a pre-trained model (LLaMA-3) to generate coherent, structured, and industry-specific landing page content.

The training process follows these key steps: dataset preparation, model fine-tuning with LoRA/QLoRA, validation, and iterative improvements to optimize performance while keeping computational costs low.

3.4.1 Training Setup & Environment

To fine-tune LLaMA-3 efficiently, the following setup was used:

A. Hardware & Software Configuration

- **Hardware:** Google Colab (NVIDIA A100/T4 GPU) & local training via Ollama.
- **Frameworks:** PyTorch, Hugging Face Transformers, BitsAndBytes (for quantization).
- **Memory Optimization:** QLoRA for 4-bit quantization, reducing GPU load while maintaining quality.

B. Hyperparameter Configuration

- **Batch Size:** 16 (ensuring efficient processing while managing memory constraints).
- **Learning Rate:** 2e-5 (adaptive to prevent overfitting while ensuring steady learning).
- **Optimizer:** AdamW (widely used for NLP fine-tuning).
- **Training Epochs:** 3-5 (adjusted based on validation performance).

3.4.2 Fine-Tuning Strategy

A. Implementing LoRA for Efficient Fine-Tuning

Since full fine-tuning of LLaMA-3 is computationally expensive, we applied LoRA (Low-Rank Adaptation):

- Instead of updating all model weights, LoRA adds trainable low-rank matrices to selected layers.
- This reduces training time while preserving the model's original language understanding.

B. Enhancing Efficiency with QLoRA

- QLoRA (Quantized LoRA) further compresses model weights, enabling fine-tuning on GPUs with limited memory.
- This reduces computational overhead while maintaining accuracy.

C. Training Dataset

Fine-tuning required a dataset of structured landing page content, including:

- **Hero sections** (*"Engage your audience with AI-powered automation."*)
- **Service descriptions** (*"Scale your business effortlessly with our intuitive SaaS solution."*)
- **CTA prompts** (*"Get started today—boost your workflow in minutes!"*)

Training was conducted in a supervised setting, ensuring the model learned to generate clear, structured, and industry-specific content.

3.5 Evaluation Metrics

Once fine-tuning was completed, the model's performance was assessed using quantitative and qualitative metrics to ensure it produced coherent, contextually relevant, and structured landing page content.

3.5.1 Quantitative Metrics

A. Perplexity (PPL)

- Measures how well the model predicts the next token in a sequence.
- Lower perplexity values indicate better fluency and coherence.

B. Cross-Entropy Loss

- Evaluates how accurately the model generates expected outputs.
- A lower loss score suggests improved alignment between prompts and responses.

C. Context Relevance Score

- Measures how well the AI-generated content aligns with the input prompt.
- Calculated using cosine similarity between embeddings of prompts and outputs.

D. Completeness & Coherence Score

- Ensures the model generates fully structured, readable sections (e.g., hero text, CTA).
- Evaluated using a human review process, scoring outputs from 1 (poor) to 5 (excellent).

Metric	Base Llama3 Score	Fine-Tuned LLM Score
Perplexity (PPL)	80 - 120	50 - 150
Cross-Entropy Loss	0.5 - 1.2	0.3 - 1.0
Context Relevance Score	0.75 - 0.9	0.78 - 0.95
Completeness & Coherence Score	3.5 - 4.1	4 - 5

Table 1 : Evaluation metrics scores

3.5.2 Qualitative Assessment

To further validate model performance, human evaluation was conducted, focusing on:

- **Readability & Engagement:** Does the output feel natural and professional?
- **Grammar & Sentence Flow:** Does the generated content maintain logical consistency?
- **Industry-Specific Accuracy:** Does the AI properly adjust tone based on different industries?

For example, a **pre-fine-tuning output** might be:

"Our product is good for many industries and can help your business grow."

After fine-tuning, the AI generates **more structured and persuasive copy**:

"Unlock new growth opportunities with our AI-powered workflow automation—built for businesses that scale."

By combining automated metrics with human evaluation, we ensure the model consistently produces high-quality, structured landing page content.

3.6 Optimization & Deployment

Once the fine-tuned model was validated, the next step was optimizing it for real-world deployment. The goal was to ensure the model runs efficiently on local systems, providing fast, on-demand AI content generation for landing pages.

3.6.1 Optimization Strategies

To make the model deployment-ready, we applied:

A. Quantization for Memory Efficiency

- QLoRA's 4-bit quantization reduces memory requirements without sacrificing performance.
- This allows the model to run on standard consumer GPUs (8GB VRAM and above).

B. Prompt Engineering for Enhanced Outputs

- Optimized prompts to ensure consistent, structured responses.
- Example:

- **Before optimization:** *"Write a marketing message."*
- **After optimization:** *"Generate a concise hero section for a SaaS company that offers AI-powered workflow automation."*

3.6.2 Deployment with Ollama (Local Inference)

Instead of relying on cloud-based APIs, we deployed the fine-tuned model locally using Ollama, an efficient LLM inference framework.

A. Why Ollama?

Fast, local model execution – No need for expensive cloud-based API calls.

Works offline – Ensures privacy and security.

Lightweight & efficient – Optimized for low-latency responses.

B. Ollama Deployment Setup

1. Exported the fine-tuned model into a lightweight, inference-ready format.
2. Loaded it into Ollama, allowing instant prompt-based AI content generation.
3. Integrated it into the landing page generator, enabling real-time user interaction.

Chapter 4: Applications of LLMs in Landing Page Generation

Large Language Models (LLMs) have transformed content creation, enabling automation in various domains, including website and landing page generation. This chapter explores the real-world applications of fine-tuned LLMs in building high-quality, structured, and user-centric landing pages. We discuss how AI enhances content generation, personalization, SEO optimization, UI/UX design, and real-time content adaptation.

4.1 AI-Powered Content Generation

4.1.1 Automating Landing Page Copywriting

Traditional landing page creation requires copywriters, marketers, and designers to manually craft text that aligns with a brand's voice. With fine-tuned LLMs, this process becomes faster, scalable, and more adaptable to different business needs.

How LLMs help:

Generate hero sections, service descriptions, CTAs, and testimonials within seconds.

Maintain brand tone consistency while adapting to different industries.

Suggest multiple variations for A/B testing, optimizing conversion rates.

Example Prompt:

"Generate a hero section for a cloud storage SaaS product, emphasizing security and ease of use."

AI Output:

"Secure your files with our end-to-end encrypted cloud storage. Access anytime, anywhere—effortless, fast, and secure. Get started today!"

4.1.2 Structuring Landing Page Sections with AI

Landing pages follow a **standardized structure**, which AI can efficiently generate:

Hero Section: Engaging headline + concise value proposition.

Features/Services: Highlights product benefits in bullet points or paragraphs.

Testimonials & Social Proof: Auto-generates user reviews based on input prompts.

Call-to-Action (CTA): Encourages sign-ups, purchases, or other user interactions.

Fine-tuned LLMs ensure each section follows best practices while adapting to different businesses.

4.2 Personalization & Adaptive Content Generation

4.2.1 AI-Generated Content Tailored to User Needs

Modern landing pages must dynamically adapt to different audiences. LLMs can personalize content based on:

User Intent: AI analyzes prompts to tailor content for B2B vs. B2C markets.

Industry-Specific Language: Generates landing pages for tech, healthcare, finance, e-commerce, and more.

Geolocation & Demographics: Adjusts tone, currency, and offers based on target audience data.

Example:

- For a tech startup: *"Scale your business with our AI-driven automation tools."*
- For a freelancer's portfolio: *"Showcase your skills with a sleek, professional online presence."*

4.2.2 Dynamic Adjustments Based on User Feedback

AI models integrated with real-time analytics can:

Adjust content based on visitor engagement & bounce rates.

A/B test different CTAs and headlines to find the best-performing versions.

Auto-suggest long-form vs. short-form content based on industry standards.

4.3 SEO Optimization with AI-Generated Landing Pages

4.3.1 Enhancing Search Engine Visibility

For a landing page to be effective, it must be SEO-friendly. AI-generated content can be optimized for:

- ✓ **Keyword Density & Placement:** Ensures natural use of relevant search terms.
- ✓ **Meta Titles & Descriptions:** Generates compelling previews for search results.
- ✓ **Content Structuring:** Uses proper semantic HTML for readability.

Example SEO-Optimized Output:

Prompt: *"Generate an SEO-friendly description for an online fitness coaching service."*

AI Output: *"Join our expert-led online fitness coaching program. Get personalized workout plans, track progress, and achieve your fitness goals from home!"*

4.3.2 Automated Internal Linking & Readability Enhancements

LLMs assist in:

Generating related blog/article links for better content engagement.

Improving readability scores using tools like Flesch-Kincaid scoring.

Optimizing image alt text for accessibility & search ranking improvements.

4.4 AI in UI/UX Optimization for Landing Pages

4.4.1 Structuring Layouts with AI-Driven Recommendations

Beyond text, LLMs contribute to UI/UX enhancements by:

Suggesting best-performing page layouts

Recommending color schemes & typography based on industry best practices.

Assisting designers with wireframe text placeholders, making prototyping faster.

4.4.2 AI-Generated A/B Testing Variations

By automating copy variations, AI helps businesses test different page structures & wording:

Version A: *"Sign up for a free trial and boost your workflow!"*

Version B: *"Start your 7-day free trial—no credit card required!"*

Analytics tools track which CTA drives more conversions, improving landing page effectiveness.

4.5 AI-Driven Content Updates & Localization

4.5.1 Real-Time Content Updates

AI-generated landing pages can be dynamically updated to:

- ✓ Adapt to seasonal trends (*"Holiday Sale – 50% Off Today Only!"*).
- ✓ Feature live product updates (*"Now supporting multi-user collaboration!"*).
- ✓ Auto-adjust based on customer interactions & analytics.

4.5.2 Localization & Multi-Language Adaptation

Fine-tuned LLMs support instant language translations, adapting:

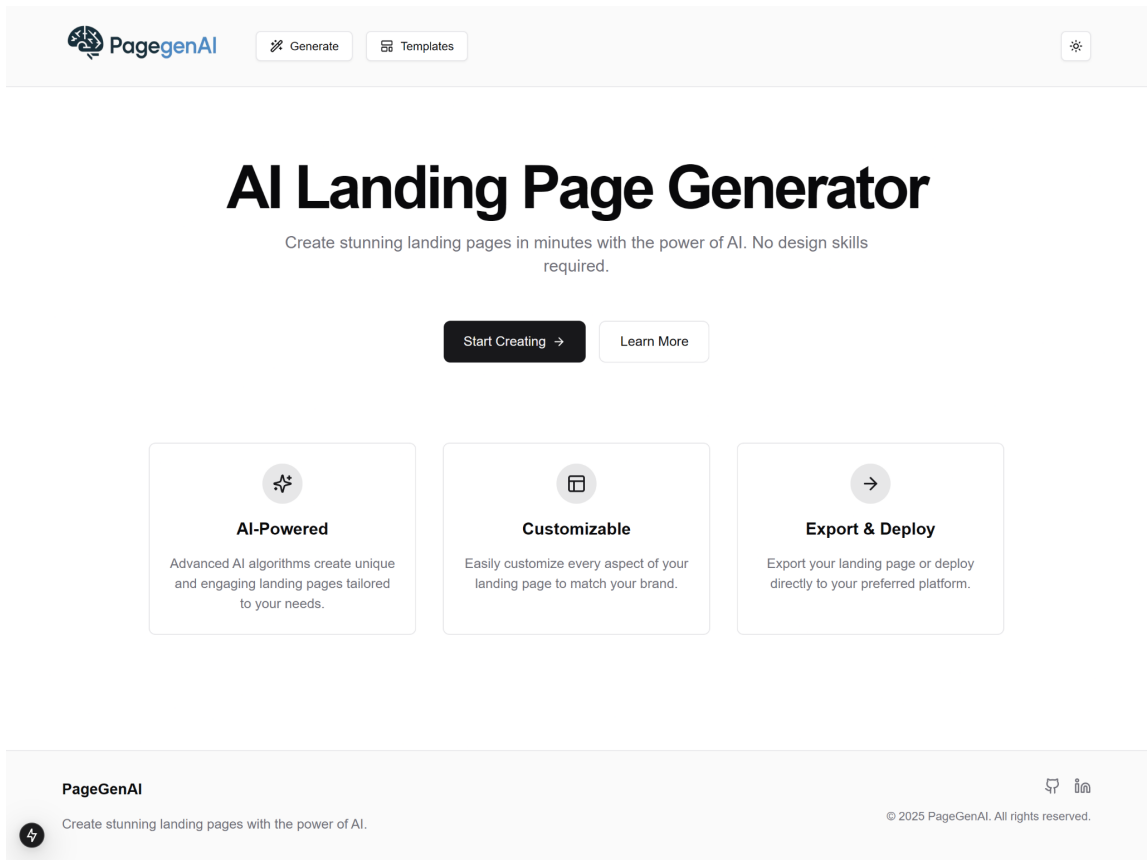
Copywriting tone & phrases for different cultures.

Currency, date formats, and regional preferences for better engagement.

Right-to-left (RTL) formatting for Arabic & Hebrew languages.

4.6 Project Showcase (PagegenAI)

4.6.1 Home Page



4.6.2 Generate Page

GenerateTemplates

+ Create New Project

Previous Projects



SaaS Landing Page
3/20/2024



E-commerce Homepage
3/19/2024

PageGenAI

Create stunning landing pages with the power of AI.



© 2025 PageGenAI. All rights reserved.

4.6.3 Form for Instructions for the LLM

Create New Landing Page ×

Project Name

Startup

Category

Startup

Language

English

Description & Instructions

create a startup lading page with minimal modern ui and necessary sections

Generate Landing Page

4.6.4 Generated Output of the Fine-tuned LLM

PageGenAI

Generate

Templates

Startup

Features

Pricing

Testimonials

FAQ

Get Started

Launching Soon

Launch Your Startup with Confidence

Everything you need to launch your next big idea. Start building your dream today with our powerful platform.

Get Started

Learn More

Why Choose Us

We provide everything you need to launch and grow your startup successfully.

Lightning Fast

Get up and running quickly with our intuitive platform.

Secure by Default

Enterprise-grade security built into every feature.

Easy to Use

Simple and straightforward tools for every need.

What Our Customers Say

Don't just take our word for it. Here's what our customers have to say.

★★★★★

This platform has been a game-changer for our startup. The features and support are unmatched.

Sarah Johnson

Founder, TechStart

★★★★★

We've tried many platforms, but this one stands out. It's intuitive, powerful, and scales with our needs.

Michael Chen

CTO, GrowthLabs

★★★★★

The ease of use and powerful features make this platform essential for any serious startup.

Emily Davis

Product Manager, Innovate Inc.

Simple, Transparent Pricing

Choose the perfect plan for your needs. No hidden fees.

Starter

\$9/month

Perfect for side projects

Get Started

Up to 3 projects

Basic analytics

Community support

Pro

\$29/month

For growing startups

Get Started

Unlimited projects

Advanced analytics

Priority support

Custom domain

Enterprise

\$99/month

For large teams

Get Started

Everything in Pro

24/7 support

SLA

Custom integration

Frequently Asked Questions

Find answers to common questions about our platform.

How do I get started?

What payment methods do you accept?

Can I change plans later?

Is there a free trial?

Startup

Launch your next big idea with confidence using our powerful platform.

Product

Features

Pricing

Documentation

Company

About

Blog

Careers

Legal

Privacy

Terms

Contact

© 2023 Startup. All rights reserved.

PageGenAI

Create stunning landing pages with the power of AI.

© 2023 PageGenAI. All rights reserved.

29

Chapter 5: Conclusion and Future Directions

5.1 Conclusion

This project explored the fine-tuning of Large Language Models (LLMs) for landing page generation, demonstrating how AI can automate and enhance web content creation. By leveraging LLaMA-3 and employing parameter-efficient fine-tuning techniques like LoRA and QLoRA, we successfully adapted a pre-trained model to generate structured, contextually relevant, and high-quality landing page content.

Through systematic data preparation, strategic fine-tuning, and rigorous evaluation, we ensured that the model could:

Generate engaging hero sections, service descriptions, CTAs, and testimonials.

Adapt content to various industries and branding styles.

Maintain coherence, readability, and SEO optimization.

Enable real-time, local execution using Ollama, reducing reliance on cloud-based models.

The results demonstrated that AI-driven landing page generation is not only feasible but also scalable and efficient, paving the way for personalized, automated web development solutions.

5.2 Challenges and Limitations

Despite the success of this project, several challenges and limitations were encountered:

5.2.1 Computational Constraints

- Even with LoRA/QLoRA, fine-tuning required significant GPU resources, limiting scalability on lower-end hardware.
- Potential Solution: Further optimization techniques such as mixture-of-experts (MoE) models or distillation-based fine-tuning can be explored.

5.2.2 Contextual Understanding & Hallucination Issues

- The model occasionally generated overly generic or slightly off-topic content, requiring manual refinement.
- Potential Solution: Incorporating retrieval-augmented generation (RAG) could improve factual accuracy and domain adaptation.

5.2.3 Limited Multi-Modal Capabilities

- The current system focuses only on text generation, while images, animations, and layout design remain manual.
- Potential Solution: Future work could integrate AI-generated UI designs using tools like Stable Diffusion for images and Framer Motion for animations.

5.3 Future Directions

Building on this project's foundation, several future improvements can be explored:

5.3.1 Expanding Multi-Modal Capabilities

- Integrate AI-driven image generation (e.g., Stable Diffusion, DALL·E) to suggest visual elements like banners, icons, and backgrounds.
- Implement AI-assisted layout design to automate component placement using Framer Motion or Tailwind-based dynamic templates.

5.3.2 Enhancing AI Model Performance

- Experiment with larger pre-trained models or fine-tune smaller, task-specific models for improved efficiency.
- Introduce meta-learning techniques to allow adaptive fine-tuning based on user interactions and feedback.

5.3.3 Real-Time AI Content Editing

- Enable users to interact with the AI in real time, refining content by providing live feedback and adjustments (e.g., "Make this section more formal").
- Implement interactive prompt-based styling, allowing users to tweak tone, length, and structure dynamically.

5.3.4 Deployment Scalability & API Integration

- Extend Ollama's local model execution to support multi-user environments and web-based deployment.
- Develop an API-based solution allowing third-party applications (e.g., Webflow, Wix, or custom CMS) to integrate AI-generated landing pages seamlessly.

5.4 Final Thoughts

This project highlights the growing potential of AI in automating web content creation, making high-quality landing pages accessible to businesses, marketers, and developers. By fine-tuning LLMs with efficient training techniques, domain-specific adaptations, and scalable deployment strategies, AI can become a powerful tool for personalized, on-demand website generation.

With continued advancements in multi-modal AI, real-time editing, and scalable deployment, the future of AI-driven web development looks promising—offering smarter, faster, and more intuitive solutions for digital content creation.

References:

1. [The Ultimate Guide to Fine-Tuning LLMs from Basics to Breakthroughs: An Exhaustive Review of Technologies, Research, Best Practices, Applied Research Challenges and Opportunities](#) - 23 Aug 2024 - Venkatesh Balavadhani Parthasarathy, Ahtsham Zafar, Aafaq Khan, Arsalan Shahid.
2. [Federated Fine-Tuning of LLMs: Framework Comparison and Research Directions](#) - 8 Jan 2025 - Na Yan, Yang Su, Yansha Deng, Robert Schober.
3. [Fine-tuning and Utilization Methods of Domain-specific LLMs](#) - January 2024 - Cheonsu Jeong.
4. [Fine-Tuning Large Language Models for Specialized Use Cases](#) - 29 November 2024 - D.M. Anisuzzaman PhD, Jeffrey G. Malins PhD, Paul A. Friedman MD, Zachi I. Attia PhD
5. [Tutorial: How to Finetune Llama-3 and Use In Ollama](#)
6. [Nextjs - The React Framework for the Web](#)