

# TINGZHANG HUANG

07536 471901 • [tingzhang.huang05@gmail.com](mailto:tingzhang.huang05@gmail.com) • [github.com/binaryyuki](https://github.com/binaryyuki)

## EDUCATION

King's College London | London, UK

Bachelor of Science in Computer Science with a Year in Industry | September 2024 – June 2028

## SKILLS

**Programming Languages:** Python, Go, Java, JavaScript, C++

**Frameworks & Libraries:** FastAPI, Django, Flask, Spring Boot, Gin, Next.js, React, Nuxt, PyTorch, Pandas, SciPy

**Databases & Cloud:** RESTful API, gRPC, MySQL, PostgreSQL, Oracle Cloud, Kubernetes

**Tools & Technologies:** Linux, Docker, Docker Swarm, Helm, Terraform, Prometheus, CI/CD, Git

## EXPERIENCE

Hangzhou Qimeng Technology Co., Ltd | Hangzhou, Zhejiang, China

Software Engineering Intern | June 2025 – Present

- Maintained and updated company's e-commerce website using **Java** and **Spring Boot** ensuring reliable performance.
- Assisted backend development for product order management systems using **MySQL** and **REST APIs** effectively.
- Improved **UI** design and user experience through layout optimization reducing user complaints by **30%**.
- Worked closely with cross-functional teams implementing **10+ new features** fixing critical bugs improving stability.
- Analyzed business requirements translating into technical solutions meeting customer expectations using **Agile**.

## PROJECTS

**Anime Streaming Platform** | FastAPI, Next.js, Docker, Kubernetes, Terraform, Prometheus, Oracle Cloud, CI/CD

- Founded and developed anime video streaming platform serving **3K+ monthly active users** successfully.
- Built with **FastAPI**, **Next.js**, **Docker** deployed on self-hosted **Kubernetes** cluster on **Oracle Cloud**.
- Implemented autoscaling load balancing and monitoring ensuring high availability and seamless user experience.
- Optimized video streaming pipeline for low-latency playback delivering enhanced **UX** for users globally.
- Automated **CI/CD** pipelines using **Helm**, **Terraform**, **Prometheus** reducing operations workload by **40%**.

**Kubernetes Infrastructure & Cloud-Native Migration** | Kubernetes, Docker, Microservices, Helm, Load Balancing

- Designed multi-node **Kubernetes** clusters for containerized microservices architecture ensuring high availability.
- Migrated legacy deployments to cloud-native microservices improving fault tolerance and scalability significantly.
- Ensured **99.9% uptime** through real-time health checks and failover configurations monitoring system performance.
- Streamlined infrastructure provisioning with **Helm** charts reducing manual deployment time by **50%**.

**LLM Optimization & OpenAI Ecosystem** | Python, OpenAI API, Anthropic API, PyTorch, Prompt Engineering, LLM

- Contributed to **gpt\_academic** open-source project reducing batch processing time by **30%** through optimization.
- Optimized prompt engineering and **Python** pipelines for large-scale academic tasks improving efficiency.
- Experimented with **LLM** finetuning and parameter tuning for summarization and citation extraction tasks.
- Integrated **APIs** from **OpenAI**, **Anthropic** delivering real-world **AI-driven** applications solving complex problems.

**Project STAY** | Python, Machine Learning, ADHD Analysis, Data Visualization, AI Personalization

- Developed predictive model for **ADHD** behavior analysis achieving high accuracy with **90%+ precision**.
- Combined questionnaire data online activity and fragmented learning records creating **AI-driven** plans.
- Created **AI-driven** learning plans enhancing engagement and outcomes for students educators psychologists.
- Built dashboards presenting insights for psychologists and educators improving decision-making processes.

**Distributed Task Queue System** | Go, Redis, RabbitMQ, PostgreSQL, Docker, Kubernetes, gRPC, REST API

- Built distributed task queue system using **Go** and **Redis** handling **20K+ daily tasks** reliably.
- Implemented message broker with **RabbitMQ** supporting priority scheduling and retry mechanisms for task failures.
- Designed task metadata storage with **PostgreSQL** and **Redis** caching reducing query latency by **50%**.
- Deployed with **Docker** and **Kubernetes** achieving horizontal scaling and **99.8% uptime** with automated failover.