

Q1. What is Pinecone and what machine learning or AI problems is it designed to solve?

Answer:

Q2. Explain the core abstractions or APIs in Pinecone (models, tensors, pipelines, agents).

Answer:

Q3. How does Pinecone handle training: defining a model, specifying a loss, and running optimization?

Answer:

Q4. What is Pinecone's approach to data ingestion, preprocessing, and batching?

Answer:

Q5. How do you evaluate model performance in Pinecone (metrics, validation sets, cross-validation)?

Answer:

Q6. How does Pinecone support model deployment and serving in a production environment?

Answer:

Q7. What hardware acceleration does Pinecone support (GPU, TPU) and how do you configure it?

Answer:

Q8. How does Pinecone handle distributed or multi-GPU training?

Answer:

Q9. What are common debugging techniques for model training issues in Pinecone?

Answer:

Q10. How do you version and experiment-track models in a Pinecone-based ML workflow?

Answer: