

NCA GENERATIVE AI LLM (NCA-GENL) PRACTICE EXAM (SAMPLE)

STUDY GUIDE



**SAMPLE STUDY GUIDE
WITH LIMITED QUESTIONS**

**VISIT US ONLINE FOR
THE FULL VERSION STUDY GUIDE**

<https://ncagenl.examzify.com>



Copyright © 2026 by Examzify - A Kaluba Technologies Inc. product.

ALL RIGHTS RESERVED.

No part of this book may be reproduced or transferred in any form or by any means, graphic, electronic, or mechanical, including photocopying, recording, web distribution, taping, or by any information storage retrieval system, without the written permission of the author.

Notice: Examzify makes every reasonable effort to obtain accurate, complete, and timely information about this product from reliable sources.

SAMPLE

Table of Contents

Copyright	1
Table of Contents	2
Introduction	3
How to Use This Guide	4
Questions	5
Answers	8
Explanations	10
Next Steps	16

SAMPLE

Introduction

Preparing for a certification exam can feel overwhelming, but with the right tools, it becomes an opportunity to build confidence, sharpen your skills, and move one step closer to your goals. At Examzify, we believe that effective exam preparation isn't just about memorization, it's about understanding the material, identifying knowledge gaps, and building the test-taking strategies that lead to success.

This guide was designed to help you do exactly that.

Whether you're preparing for a licensing exam, professional certification, or entry-level qualification, this book offers structured practice to reinforce key concepts. You'll find a wide range of multiple-choice questions, each followed by clear explanations to help you understand not just the right answer, but why it's correct.

The content in this guide is based on real-world exam objectives and aligned with the types of questions and topics commonly found on official tests. It's ideal for learners who want to:

- Practice answering questions under realistic conditions,
- Improve accuracy and speed,
- Review explanations to strengthen weak areas, and
- Approach the exam with greater confidence.

We recommend using this book not as a stand-alone study tool, but alongside other resources like flashcards, textbooks, or hands-on training. For best results, we recommend working through each question, reflecting on the explanation provided, and revisiting the topics that challenge you most.

Remember: successful test preparation isn't about getting every question right the first time, it's about learning from your mistakes and improving over time. Stay focused, trust the process, and know that every page you turn brings you closer to success.

Let's begin.

How to Use This Guide

This guide is designed to help you study more effectively and approach your exam with confidence. Whether you're reviewing for the first time or doing a final refresh, here's how to get the most out of your Examzify study guide:

1. Start with a Diagnostic Review

Skim through the questions to get a sense of what you know and what you need to focus on. Your goal is to identify knowledge gaps early.

2. Study in Short, Focused Sessions

Break your study time into manageable blocks (e.g. 30 – 45 minutes). Review a handful of questions, reflect on the explanations.

3. Learn from the Explanations

After answering a question, always read the explanation, even if you got it right. It reinforces key points, corrects misunderstandings, and teaches subtle distinctions between similar answers.

4. Track Your Progress

Use bookmarks or notes (if reading digitally) to mark difficult questions. Revisit these regularly and track improvements over time.

5. Simulate the Real Exam

Once you're comfortable, try taking a full set of questions without pausing. Set a timer and simulate test-day conditions to build confidence and time management skills.

6. Repeat and Review

Don't just study once, repetition builds retention. Re-attempt questions after a few days and revisit explanations to reinforce learning. Pair this guide with other Examzify tools like flashcards, and digital practice tests to strengthen your preparation across formats.

There's no single right way to study, but consistent, thoughtful effort always wins. Use this guide flexibly, adapt the tips above to fit your pace and learning style. You've got this!

Questions

SAMPLE

- 1. Which technique focuses on improving speed while reducing memory use during training cycles?**
 - A. Gradient Clipping
 - B. Mixed Precision Training
 - C. Data Augmentation
 - D. Transfer Learning
- 2. Which mechanism optimizes memory usage and computation managing attention in models?**
 - A. Adaptive Attention
 - B. Paged Attention
 - C. Dynamic Attention
 - D. Contextual Attention
- 3. Which concept represents a comparison between model output distributions and true output forms?**
 - A. Objective Function
 - B. Cross-Entropy Loss
 - C. Gradient Checkpointing
 - D. Penalization Mechanism
- 4. What technique enables an LLM to dynamically update and maintain relevant information from various input sources?**
 - A. Contextual Learning
 - B. Dynamic Memory Learning
 - C. In-Context Learning with Dynamic Memory
 - D. Adaptive Input Learning
- 5. Which of these tools helps in ensemble learning methods?**
 - A. Tri Library
 - B. NVIDIA Jarvis
 - C. NVIDIA Transfer Learning Toolkit (TLT)
 - D. None of the above

- 6. What aspect of generative AI does steerable AI primarily control?**
- A. Data preprocessing and cleaning
 - B. Model training speed and efficiency
 - C. Model output characteristics such as style and tone
 - D. Hardware resource allocation
- 7. What does the Recall-Oriented Understudy for Gisting Evaluation (ROUGE) metric evaluate?**
- A. Word error rates in speech recognition
 - B. The performance of summarization and translation systems
 - C. The speed of GPU operations
 - D. Similarity scores between input texts
- 8. What method can help mitigate bias in an LLM after its deployment?**
- A. Regularly retraining the model
 - B. Adding more layers to the model
 - C. Reducing the model size
 - D. Using simpler algorithms
- 9. Which term refers broadly to the function that the model aims to minimize during training?**
- A. Synchronous Updates
 - B. Objective Function
 - C. Asynchronous Updates
 - D. Gradient Checkpointing
- 10. Which product refers to NVIDIA's microservices architecture for inference?**
- A. NVIDIA TensorRT
 - B. NVIDIA NIM
 - C. NVIDIA Docker
 - D. NVIDIA GPU Cloud

Answers

SAMPLE

1. B
2. B
3. B
4. C
5. D
6. C
7. B
8. A
9. B
10. B

SAMPLE

Explanations

SAMPLE

1. Which technique focuses on improving speed while reducing memory use during training cycles?

A. Gradient Clipping

B. Mixed Precision Training

C. Data Augmentation

D. Transfer Learning

Mixed Precision Training is a technique that optimizes the use of computational resources during the training of machine learning models. By leveraging lower-precision data types, such as half-precision floating-point (FP16) rather than full-precision (FP32), Mixed Precision Training can significantly enhance training speed due to reduced memory bandwidth requirements and faster arithmetic operations on compatible hardware (like GPUs). When using lower precision, the model utilizes less memory, allowing larger batch sizes or more complex models to fit within the same memory constraints. This results in improved performance during the training cycles, as the model can operate more efficiently while still maintaining accuracy. Moreover, hardware that supports mixed precision can execute multiple operations in parallel, further accelerating the training process. Other techniques like Gradient Clipping, Data Augmentation, and Transfer Learning serve different purposes. Gradient Clipping addresses exploding gradients by constraining their values, Data Augmentation increases the diversity of the training dataset to help the model generalize better, and Transfer Learning utilizes pre-trained models on new tasks to save resources. None of these techniques primarily focus on balancing speed and memory use in the way that Mixed Precision Training does.

2. Which mechanism optimizes memory usage and computation managing attention in models?

A. Adaptive Attention

B. Paged Attention

C. Dynamic Attention

D. Contextual Attention

Paged attention is a mechanism designed to optimize memory usage and computation for managing attention in language models. It effectively reduces the amount of memory required by only keeping relevant parts of the input data in the attention mechanism at any given time. This is particularly important in deep learning models, where the attention mechanism can become a bottleneck due to the quadratic scaling of memory usage with the length of the input sequence. In paged attention, segments or "pages" of the input are used instead of trying to attend to the entire input sequence all at once. This approach allows the model to selectively focus on specific portions of the input, which not only minimizes memory consumption but also enables faster computation. By leveraging this segmented approach, models can handle longer sequences without overwhelming computational resources, making it a practical solution for scaling attention mechanisms effectively. Other mechanisms, while they may also aim to improve attention management, do not specifically target optimizations related to memory and computation in the way that paged attention does.

3. Which concept represents a comparison between model output distributions and true output forms?

- A. Objective Function
- B. Cross-Entropy Loss**
- C. Gradient Checkpointing
- D. Penalization Mechanism

The concept of cross-entropy loss is fundamentally tied to measuring the difference between two probability distributions: the predicted output from a model and the actual distribution of the true labels. This metric is particularly important in classification tasks, where we often want to compare how close the predicted probabilities are to the one-hot encoded vectors representing the true classes. Cross-entropy loss quantifies the dissimilarity between the predicted distribution (what the model thinks is the right answer) and the true distribution (what the actual answers are). It effectively provides a single value that reflects not just whether a prediction is right or wrong, but how confident it is in its predictions. A lower cross-entropy loss indicates a closer match between the model's predictions and the actual labels, promoting better learning during the training process. The objective function is a broader term that encompasses various types of loss functions used during training to measure performance, but it does not specifically represent the comparison between predicted and true distributions. Gradient checkpointing is a memory optimization technique used during training to save computational resources, not a measure of output distribution comparison. A penalization mechanism might involve adding constraints or penalties to the learning process, but it does not define the specific comparison between output distributions.

4. What technique enables an LLM to dynamically update and maintain relevant information from various input sources?

- A. Contextual Learning
- B. Dynamic Memory Learning
- C. In-Context Learning with Dynamic Memory**
- D. Adaptive Input Learning

The selected answer, "In-Context Learning with Dynamic Memory," accurately reflects a technique that allows a large language model (LLM) to dynamically update and maintain relevant information from various input sources. In this context, "In-Context Learning" refers to the LLM's ability to utilize previously provided examples or data during the current interaction to inform its responses. This enables the model to adapt its understanding and responses based on real-time input, making it highly responsive to user queries. Dynamic Memory complements this by allowing the model to retain and recall pertinent information from prior interactions or inputs, effectively creating a memory-like system. This combination empowers the LLM to fluidly adjust to new data while retaining important context that has been accumulated over time. Therefore, it becomes adept at providing nuanced and contextual responses that are informed by both static knowledge and dynamic updates. The other techniques mentioned do not encapsulate the same depth of real-time adaptability combined with memory capacity. Contextual Learning generally pertains to understanding the context in which inputs are provided but does not extensively cover the aspect of dynamically maintaining memory. Dynamic Memory Learning focuses primarily on storage and retrieval but may not account for in-context examples effectively. Adaptive Input Learning, while useful for adjusting processing based on input, does not specifically

5. Which of these tools helps in ensemble learning methods?

- A. Tri Library
- B. NVIDIA Jarvis
- C. NVIDIA Transfer Learning Toolkit (TLT)
- D. None of the above**

The correct answer hinges on the understanding of ensemble learning methods, which combine multiple models to improve the overall performance compared to a single model. Ensemble learning techniques such as bagging, boosting, and stacking rely on the integration of various algorithms to enhance accuracy and reduce overfitting. In this context, the tools listed are recognized for their primary applications. The Tri Library is not specifically associated with ensemble learning but rather focuses on other functionalities. NVIDIA Jarvis is a toolkit for building and deploying conversational AI applications, which doesn't directly support the ensemble learning framework. The NVIDIA Transfer Learning Toolkit (TLT) is designed to facilitate transfer learning and model customization primarily for deep learning, rather than specifically enabling ensemble methods. Consequently, considering the purpose of the mentioned tools, none of them are primarily designed to assist with ensemble learning methods, making "None of the above" the accurate conclusion.

6. What aspect of generative AI does steerable AI primarily control?

- A. Data preprocessing and cleaning
- B. Model training speed and efficiency
- C. Model output characteristics such as style and tone**
- D. Hardware resource allocation

Steerable AI primarily focuses on controlling model output characteristics such as style and tone. This aspect allows users to dictate how the generated content should be shaped according to specific preferences or requirements. For example, if an application demands a formal tone or a creative style, steerable AI can adjust the generation process accordingly. This capability enhances the versatility of generative AI applications, enabling developers and users to tailor the responses or outputs of models for different contexts and audiences. The emphasis on output characteristics distinguishes steerable AI from other aspects like data preprocessing, which involves preparing data for training, model training speed, which pertains to the efficiency of the training process, and hardware resource allocation, which deals with managing computational resources. Steerable AI's capability to modify how content is presented is crucial for applications that require targeted communication and engagement with users.

7. What does the Recall-Oriented Understudy for Gisting Evaluation (ROUGE) metric evaluate?

- A. Word error rates in speech recognition
- B. The performance of summarization and translation systems**
- C. The speed of GPU operations
- D. Similarity scores between input texts

The Recall-Oriented Understudy for Gisting Evaluation (ROUGE) metric primarily evaluates the performance of summarization and translation systems. This metric assesses the quality of summaries generated by comparing them to reference summaries, which can be particularly useful in understanding how well a system captures essential information. ROUGE focuses on the recall aspect, measuring the overlap of n-grams (contiguous sequences of n items) between the generated summaries and the reference summaries. By doing so, it provides insights into how comprehensively the system is conveying the key ideas from a larger text. The importance of ROUGE lies in its ability to facilitate objective comparisons between different text generation models, making it a widely used evaluation tool in natural language processing tasks related to summarization and translation. This allows researchers and practitioners to gauge improvements in their models over time or relative to other systems. In contrast, other options do not align with the function of ROUGE. For instance, evaluating word error rates pertains to aspects of speech recognition rather than summarization. Similarly, measuring the speed of GPU operations is unrelated to content evaluation metrics like ROUGE, and while similarity scores between texts might seem relevant, they do not specifically address summarization and translation performance. Thus, the choice focusing on summarization and translation

8. What method can help mitigate bias in an LLM after its deployment?

- A. Regularly retraining the model**
- B. Adding more layers to the model
- C. Reducing the model size
- D. Using simpler algorithms

Regularly retraining the model is an effective method for mitigating bias in a large language model (LLM) after its deployment. This process allows the model to continuously learn from new data, which can be more representative of diverse populations and viewpoints over time. As societal norms and languages evolve, models need to be updated to reflect these changes, helping to correct biases that may have been present in the original training data. By incorporating more current information, retraining can alleviate the effects of historical biases or gaps in the initial dataset. In contrast, adding more layers to the model does not inherently address bias; it may complicate the model's architecture without ensuring a more equitable representation of information. Similarly, reducing the model size can potentially impact its performance and capability to understand nuanced scenarios, which does not directly target bias mitigation. Using simpler algorithms could lead to generalized outcomes that may overlook the complex interactions required to address biases in language understanding effectively. Therefore, regular retraining stands out as the most practical and direct approach to tackle bias in LLMs post-deployment.

9. Which term refers broadly to the function that the model aims to minimize during training?

- A. Synchronous Updates
- B. Objective Function**
- C. Asynchronous Updates
- D. Gradient Checkpointing

The term that refers broadly to the function that the model aims to minimize during training is the objective function. In the context of machine learning and generative AI, the objective function quantifies how well the model performs in relation to the task at hand. During training, the model seeks to adjust its parameters in order to minimize the value of this function, thereby improving its predictions or outputs. The objective function can take various forms depending on the specific task—such as mean squared error for regression problems or cross-entropy loss for classification tasks. The minimization process typically involves algorithms that compute the gradients of the objective function with respect to the model parameters, guiding the optimization process. The other terms mentioned serve different purposes. Synchronous updates and asynchronous updates refer to methodologies for updating model parameters in distributed training environments, focusing on how and when the model parameters are adjusted based on training data. Gradient checkpointing is a technique employed to manage memory usage during the training of large models, allowing for efficient computation of gradients without holding all intermediate activations in memory. These concepts, while essential to model training and optimization, do not encapsulate the core concept of the function being minimized, which is the essence of the objective function.

10. Which product refers to NVIDIA's microservices architecture for inference?

- A. NVIDIA TensorRT
- B. NVIDIA NIM**
- C. NVIDIA Docker
- D. NVIDIA GPU Cloud

NVIDIA NIM (NVIDIA Inference Management) refers to NVIDIA's microservices architecture specifically designed for deploying and managing inference workloads. It allows developers to create a flexible environment where they can manage multiple machine learning models and services as microservices. This architecture is particularly beneficial for handling scalable and efficient inference operations in production settings, facilitating the integration of diverse models and making it easier to leverage GPUs for processing. NVIDIA TensorRT is primarily a high-performance inference optimizer and runtime for deep learning models, focusing on optimizing neural networks for deployment on NVIDIA GPUs. Although it plays a crucial role in inference, it does not encompass the broader microservices architecture that NIM provides. NVIDIA Docker enables the creation and management of Docker containers tailored for NVIDIA GPUs, but it does not specifically address the microservices architecture for inference workloads. It is more about containerization than the architecture itself. NVIDIA GPU Cloud is a platform that provides a catalog of GPU-optimized software and resources. It facilitates cloud-based access to NVIDIA's technologies, but it does not specifically define a microservices architecture for managing inference tasks. Thus, NVIDIA NIM is the correct answer as it specifically addresses the needs of managing inference as a set of microservices.

Next Steps

Congratulations on reaching the final section of this guide. You've taken a meaningful step toward passing your certification exam and advancing your career.

As you continue preparing, remember that consistent practice, review, and self-reflection are key to success. Make time to revisit difficult topics, simulate exam conditions, and track your progress along the way.

If you need help, have suggestions, or want to share feedback, we'd love to hear from you. Reach out to our team at hello@examzify.com.

Or visit your dedicated course page for more study tools and resources:

<https://ncagenl.examzify.com>

We wish you the very best on your exam journey. You've got this!

SAMPLE