

NVIDIA NCA-GENL復習解答例、NCA-GENL試験対応



BONUS!!! PassTest NCA-GENLダンプの一部を無料でダウンロード：https://drive.google.com/open?id=1bZ_ePG0kXmIgvpydwZpzvR3ewpx8qZQd

PassTestはNVIDIAのNCA-GENL認定試験についてすべて資料を提供するの唯一サイトでございます。受験者はPassTestが提供した資料を利用してNCA-GENL認証試験は問題にならないだけでなく、高い点数も合格することができます。

NVIDIA NCA-GENL 認定試験の出題範囲：

トピック	出題範囲
トピック 1	データ分析と可視化：この試験セクションでは、データサイエンティストのスキルを測定します。データの解釈、クレンジング、そして視覚的なストーリーテリングによる提示方法を網羅しています。特に、可視化を用いて洞察を抽出し、モデルの挙動、パフォーマンス、あるいはトレーニングデータのパターンを評価する方法に重点が置かれています。
トピック 2	ソフトウェア開発：この試験セクションでは、機械学習開発者のスキルを測定し、AIアプリケーション向けの効率的でモジュール化されたスケーラブルなコードの作成方法を網羅します。LLMベースの開発に関連するソフトウェアエンジニアリングの原則、バージョン管理、テスト、ドキュメント作成の実践が含まれます。
トピック 3	実験：このセクションでは、MLエンジニアのスキルを測定し、LLMを用いた構造化された実験の実施方法を網羅します。テストケースの設定、パフォーマンス指標の追跡、そして実験結果に基づいた情報に基づいた意思決定などが含まれます。
トピック 4	データ前処理と特徴量エンジニアリング：この試験セクションでは、データエンジニアのスキルを測定し、モデルのトレーニングや微調整に使用可能な形式への生データの準備について学習します。堅牢なLLMパイプラインの構築に不可欠な、クリーニング、正規化、トークン化、特徴抽出手法も網羅しています。
トピック 5	機械学習とニューラルネットワークの基礎：このセクションでは、AI研究者のスキルを測定します。機械学習とニューラルネットワークの基礎原理を網羅し、これらの概念が大規模言語モデル（LLM）の開発にどのように貢献しているかに焦点を当てます。学習者が生成型AIシステムの学習に関わる基本構造と学習メカニズムを理解できるようにします。

トピック 6	• アライメント：このセクションでは、AIポリシーエンジニアのスキルを測定し、LLMの出力を人間の意図や価値観と整合させるための手法を網羅します。これには、モデルから生じる有害、偏った、または不正確な結果を削減するための安全メカニズム、倫理的セーフガード、チューニング戦略が含まれます。
トピック 7	• この試験セクションでは、AI製品開発者のスキルを測定し、仮説の検証、モデルのバリエーションの比較、モデルの応答のテストなどを行う実験を戦略的に計画する方法を扱います。実験における構造、制御、変数に焦点を当てます。

>> NVIDIA NCA-GENL復習解答例 <<

NVIDIA NCA-GENL試験対応、NCA-GENLトレーニング費用

有名なブランドPassTestとして、非常に成功しているにもかかわらず、NVIDIA現状に満足することはなく、常にNCA-GENL試験トレントの内容を更新する用意があります。最も重要なことは、NCA-GENLガイドトレントの新しいバージョンをコンパイルしている限り、購入後1年間無料で最新バージョンのNCA-GENLトレーニング資料をお客様に送信します。NCA-GENL試験に合格するのに役立つ、絶え間なく更新される試験の要求に合わせて、NVIDIA Generative AI LLMsガイド急流を引き続きお届けします。

NVIDIA Generative AI LLMs 認定 NCA-GENL 試験問題 (Q44-Q49):

質問 # 44

Which calculation is most commonly used to measure the semantic closeness of two text passages?

- A. Hamming distance
- **B. Cosine similarity**
- C. Euclidean distance
- D. Jaccard similarity

正解: B

解説:

Cosine similarity is the most commonly used metric to measure the semantic closeness of two text passages in NLP. It calculates the cosine of the angle between two vectors (e.g., word embeddings or sentence embeddings) in a high-dimensional space, focusing on the direction rather than magnitude, which makes it robust for comparing semantic similarity. NVIDIA's documentation on NLP tasks, particularly in NeMo and embedding models, highlights cosine similarity as the standard metric for tasks like semantic search or text similarity, often using embeddings from models like BERT or Sentence-BERT. Option A (Hamming distance) is for binary data, not text embeddings. Option B (Jaccard similarity) is for set-based comparisons, not semantic content. Option D (Euclidean distance) is less common for text due to its sensitivity to vector magnitude.

References:

NVIDIA NeMo Documentation: <https://docs.nvidia.com/deeplearning/nemo/user-guide/docs/en/stable/nlp/intro.html>

質問 # 45

Which feature of the HuggingFace Transformers library makes it particularly suitable for fine-tuning large language models on NVIDIA GPUs?

- A. Built-in support for CPU-based data preprocessing pipelines.
- B. Automatic conversion of models to ONNX format for cross-platform deployment.
- **C. Seamless integration with PyTorch and TensorRT for GPU-accelerated training and inference.**
- D. Simplified API for classical machine learning algorithms like SVM.

正解: C

解説:

The HuggingFace Transformers library is widely used for fine-tuning large language models (LLMs) due to its seamless integration with PyTorch and NVIDIA's TensorRT, enabling GPU-accelerated training and inference. NVIDIA's NeMo documentation

references HuggingFace Transformers for its compatibility with CUDA and TensorRT, which optimize model performance on NVIDIA GPUs through features like mixed-precision training and dynamic shape inference. This makes it ideal for scaling LLM fine-tuning on GPU clusters. Option A is incorrect, as Transformers focuses on GPU, not CPU, pipelines. Option C is partially true but not the primary feature for fine-tuning. Option D is false, as Transformers is for deep learning, not classical algorithms.

References:

NVIDIA NeMo Documentation: <https://docs.nvidia.com/deeplearning/nemo/user-guide/docs/en/stable/nlp/intro.html>

HuggingFace Transformers Documentation: <https://huggingface.co/docs/transformers/index>

質問 # 46

What is a Tokenizer in Large Language Models (LLM)?

- A. A method to remove stop words and punctuation marks from text data.
- B. A technique used to convert text data into numerical representations called tokens for machine learning.
- C. A machine learning algorithm that predicts the next word/token in a sequence of text.
- **D. A tool used to split text into smaller units called tokens for analysis and processing.**

正解: D

解説:

A tokenizer in the context of large language models (LLMs) is a tool that splits text into smaller units called tokens (e.g., words, subwords, or characters) for processing by the model. NVIDIA's NeMo documentation on NLP preprocessing explains that tokenization is a critical step in preparing text data, with algorithms like WordPiece, Byte-Pair Encoding (BPE), or SentencePiece breaking text into manageable units to handle vocabulary constraints and out-of-vocabulary words. For example, the sentence "I love AI" might be tokenized into ["I", "love", "AI"] or subword units like ["I", "lov", "##e", "AI"]. Option A is incorrect, as removing stop words is a separate preprocessing step. Option B is wrong, as tokenization is not a predictive algorithm. Option D is misleading, as converting text to numerical representations is the role of embeddings, not tokenization.

References:

NVIDIA NeMo Documentation: <https://docs.nvidia.com/deeplearning/nemo/user-guide/docs/en/stable/nlp/intro.html>

質問 # 47

In the development of Trustworthy AI, what is the significance of 'Certification' as a principle?

- **A. It involves verifying that AI models are fit for their intended purpose according to regional or industry-specific standards.**
- B. It ensures that AI systems are transparent in their decision-making processes.
- C. It requires AI systems to be developed with an ethical consideration for societal impacts.
- D. It mandates that AI models comply with relevant laws and regulations specific to their deployment region and industry.

正解: A

解説:

In the development of Trustworthy AI, 'Certification' as a principle involves verifying that AI models are fit for their intended purpose according to regional or industry-specific standards, as discussed in NVIDIA's Generative AI and LLMs course. Certification ensures that models meet performance, safety, and ethical benchmarks, providing assurance to stakeholders about their reliability and appropriateness. Option A is incorrect, as transparency is a separate principle, not certification. Option B is wrong, as ethical considerations are broader and not specific to certification. Option D is inaccurate, as compliance with laws is related but distinct from certification's focus on fitness for purpose. The course states: "Certification in Trustworthy AI verifies that models meet regional or industry-specific standards, ensuring they are fit for their intended purpose and reliable." References: NVIDIA Building Transformer-Based Natural Language Processing Applications course; NVIDIA Introduction to Transformer-Based Natural Language Processing.

質問 # 48

You are working with a data scientist on a project that involves analyzing and processing textual data to extract meaningful insights and patterns. There is not much time for experimentation and you need to choose a Python package for efficient text analysis and manipulation. Which Python package is best suited for the task?

myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw,
myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt,
myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, Disposable vapes

BONUS!!! PassTest NCA-GENLダンプの一部を無料でダウンロード: https://drive.google.com/open?id=1bZ_ePG0kXm1gpydwZpzvR3ewpx8qZQd