

CBSE CLASS X

Artificial Intelligence (417)

ANSWER KEY

From previous CBSE Board Exam questions

Code: 8Z6EEI**Questions: 53****Maximum Marks: 90****Generated: 2026-06-21 02:53**

SELECTIONS USED

Source	Previous-year board
Subject	Artificial Intelligence
Lessons	Natural Language Processing
Year	2022-2026
Questions selected	53

Composition — Types: 24 MCQ · 11 Short · 5 Long · 4 Very short · 2 competency · 2 Case-based · 2 fill_in_blank · 1 Assertion–reason

Q1. § 1.2: Introduction to AI Domains § Test Yourself § Reflection Time

[4]

Categorise the following examples under the given three AI domains — Data Science, NLP and Computer Vision with justification :

- (a) Recommendation Websites
- (b) Voice-based Virtual Assistants
- (c) Spam Filters
- (d) Airline Route Planning

Previously asked in: 2025 104/S Q20

♦ Revisiting AI Project Cycle & Ethical Frameworks for AI

1.2 Introduction to AI Domains

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

Example	Domain	Justification
(a) Recommendation Websites	Data Science (Statistical Data)	They collect and analyse large amounts of user data to suggest products/content, deriving meaningful insights from datasets.
(b) Voice-based Virtual Assistants	NLP	They interact with humans using natural spoken language; NLP algorithms decode and respond to human speech.
(c) Spam Filters	NLP	They detect certain words/phrases in emails to identify spam — one of the earliest NLP applications.
(d) Airline Route Planning	Data Science (Statistical Data)	They analyse large statistical datasets (weather, fuel, traffic) to extract insights and optimise flight routes.

Source: Chapter 1, Section 1.2 — Introduction to AI Domains

Explanation

- **Key rule:** AI domains are categorised by the *type of data* fed into the model — Statistical Data, NLP, or Computer Vision.
- Recommendation/route-planning → statistical/numerical data → **Data Science**.
- Voice assistants and spam filters both deal with **natural language** (spoken or written) → **NLP**.
- Examiners award 1 mark per correct categorisation with valid justification. Simply naming the domain without justification risks losing the mark.
- Note: The passage explicitly lists *email/spam filters* as an NLP example — quote it directly for full marks.

Q2. § Unit Overview § Test Yourself § Reflection Time

[4]

Consider the following scenarios and identify which AI domain would be most appropriate for each, with justification:

(A) An AI based education platform needs to translate to English language and analyze thousands of student essays to provide instant feedback on grammar, content quality and writing style. [2]

(B) An AI based application installed on a busy crossing in a metropolitan city scans all vehicles driving through that crossing during peak traffic hours and categorizes them into four wheelers and two wheelers. [2]

Previously asked in: 2026 104 Q18

♦ Revisiting AI Project Cycle & Ethical Frameworks for AI

1.2 Introduction to AI Domains

6.2 Applications of NLP

Computer Vision Applications – Object Detection and Classification

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

(A) Natural Language Processing (NLP)

The most appropriate AI domain is **NLP**. The platform needs to translate essays into English and then analyze grammar, content quality, and writing style — all of which involve understanding and processing human natural language. NLP algorithms extract meaning from written text, making it ideal for grammar checking, feedback generation, and language translation tasks.

(B) Computer Vision (CV)

The most appropriate AI domain is **Computer Vision**. The application scans vehicles (visual data from a busy crossing) and categorizes them into four-wheelers and two-wheelers. CV enables machines to acquire, analyse, and interpret visual information from images or video feeds, making it perfect for vehicle detection and classification tasks.

Source: Chapter 1, Section 1.2 — Introduction to AI Domains

Explanation

- **Key rule:** Match the *type of data* to the domain — text/language → NLP; images/video → Computer Vision.
- For (A), two tasks are mentioned (translation + analysis of written text) — both fall under NLP. Mention both to get full marks.
- For (B), "scans vehicles" = visual input; "categorizes" = CV's core function. Reference surveillance/detection as a CV example.
- Examiners look for: correct domain name + valid justification linked to the scenario. One without the other loses a mark.

Q3. § 5.2 Computer Vision Tasks

[2]

Differentiate between Computer Vision (CV) and Natural Language Processing (NLP).

Previously asked in: 2024 104 Q11

♦ Computer Vision

5.2 Computer Vision Tasks

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

| Basis | Computer Vision (CV) | Natural Language Processing (NLP) |

|-----|-----|-----|

| **Definition** | Enables machines to interpret and understand visual information (images/videos). | Enables machines to understand and process human language (text/speech). |

| **Input** | Images, videos | Text, audio/speech |

| **Example** | Facial recognition, object detection | Chatbots, language translation |

Source: Unit 5: Computer Vision, Chapter 5

Explanation

- Examiners expect a **tabular format** for "differentiate" questions — it's cleaner and earns full marks faster.
- You need at least **2 clear points of difference** for 2 marks.
- CV focuses on **visual data**; NLP focuses on **language data** — this is the core distinction to highlight.
- Even though NLP is not covered in the source passages, CBSE expects you to know this distinction from general AI knowledge introduced earlier in the course.

Q4. § Unit Overview and Learning Objectives

[2]

With reference to data processing, expand the term TFIDF. Also give any two applications of TFIDF.

Previously asked in: 2023 104 Q16

♦ Natural Language Processing

6.2 Applications of NLP

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

TFIDF stands for **Term Frequency-Inverse Document Frequency**.

Two applications of TFIDF:

1. **Keyword Extraction** – It helps identify the most important/valuable words in a document by assigning higher scores to rare but significant terms.
2. **Text Classification** – It is used to classify documents into categories by determining the relevance of words within them.

Explanation

- The full form alone can fetch 1 mark; each application is worth ½ mark each (totalling 1 mark) — so never skip the expansion.
- TFIDF measures how important a word is to a document in a corpus: high TF but low IDF = common word (stop word); low TF but high IDF = rare, valuable word.
- Other valid applications include: search engines, sentiment analysis, document summarisation, and information retrieval — any two are acceptable.

Q5. § 6.2 Applications of Natural Language Processing § 6.6 Natural Language Processing: Use Case Walkthrough

[1]

This real life application of NLP is used to provide an overview of a news item or blog post, while avoiding redundancy from multiple sources and maximising the diversity of content obtained. Which is this application ?

- (a) Chatbot
- (b) Virtual Assistant
- (c) Sentiment Analysis
- (d) Automatic Summarisation

Previously asked in: 2024 104 Q2 (ii)

♦ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(d) Automatic Summarisation**

This application provides an overview of a news item or blog post, avoids redundancy from multiple sources, and maximises diversity of content.

Explanation

The question describes the key features of **Automatic Summarisation** — summarising content, removing redundancy, and maximising diversity. Chatbots simulate conversation, Virtual Assistants execute voice tasks, and Sentiment Analysis detects opinions — none match this description. Choose the option that best fits all three conditions given in the question.

Q6. § 6.2 Applications of Natural Language Processing § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection

[1]

Time

Which of the following words represent an example of a lemma resulting from lemmatisation for "caring" in context to Natural Language Processing (NLP) ?

- (a) Care
- (b) Cared
- (c) Cares
- (d) Car

Previously asked in: 2024 104 Q2 (iv)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(a) Care**

Lemmatisation reduces a word to its base dictionary form (lemma). The lemma of "caring" is "**care**", which is a meaningful root word.

Explanation

Lemmatisation always produces a valid dictionary word, unlike stemming which may produce incomplete forms. "Car" is unrelated; "cared" and "cares" are inflected forms, not the base lemma. Examiners expect students to know that lemmatisation = meaningful root word (lemma).

Q7. § Unit Overview and Learning Objectives § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time **[1]**

Which of the following applications is not associated with Natural Language Processing (NLP) ?

- (a) Sentiment Analysis
- (b) Speech Recognition
- (c) Spam Filtering in emails
- (d) Stock Market Analysis

Previously asked in: 2024 104 Q3 (vi)

♦ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

(d) Stock Market Analysis

Stock Market Analysis is not an NLP application. Sentiment Analysis, Speech Recognition, and Spam Filtering all involve processing natural language text or speech.

Explanation

The textbook (Chapter 6) lists NLP applications as voice assistants, sentiment analysis, text classification, keyword extraction, language translation, and autogenerated captions. Spam filtering uses text classification (an NLP task), and speech recognition processes natural language. Stock Market Analysis is primarily a financial/statistical task, not an NLP application. Examiners expect direct identification of the odd one out with a brief reason.

Q8. § Unit Overview and Learning Objectives **[1]**

Bag of Words is a _____ model which helps in extracting features out of the text which can be helpful in machine learning algorithms.

- (a) Data Science (DS)
- (b) Virtual Reality (VR)
- (c) Natural Language Processing (NLP)
- (d) Computer Vision (CV)

Previously asked in: 2024 104 Q4 (iii)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

(c) Natural Language Processing (NLP)

Bag of Words is a **Natural Language Processing (NLP)** model which helps in extracting features out of the text which can be helpful in machine learning algorithms.

Explanation

The source passage (Chapter 6) lists "Bag of Words" as a key concept under the **NLP** chapter, and the Test Yourself section (Q6) confirms its purpose is "to extract features from text for machine learning algorithms." Students must remember BoW is an NLP concept, not DS, VR, or CV.

Q9. § 6.2 Applications of Natural Language Processing

[1]

Which NLP application helps in converting natural speech into text in real time ?

- (A) Keyword Extraction tool
- (B) Translation of books from English to Hindi language
- (C) Auto generated captions on YouTube
- (D) Classifying raw text into pre-defined groups

Previously asked in: 2026 104 Q3 (v)

♦ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding rag

Model Answer**(C) Auto generated captions on YouTube**

Auto-generated captions convert natural speech into text in real time, making video content more accessible.

Explanation

The passage explicitly states: "*Captions are generated by turning natural speech into text in real-time*" and gives YouTube as an example. Other options describe different NLP applications — keyword extraction, language translation, and text classification respectively.

Q10. § Test Yourself and Reflection Time

[1]

Assertion (A) : Converting text to lowercase is preferable in text preprocessing.

Reason (R) : It ensures that "Hello" and "hello" are treated as the same word by the machine.

- (A) Both (A) and (R) are true and (R) is the correct explanation of (A).
- (B) Both (A) and (R) are true, but (R) is not the correct explanation of (A).
- (C) (A) is true, but (R) is false.
- (D) (A) is false, but (R) is true.

Previously asked in: 2026 104 Q5 (ii)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding rag

Model Answer**(A) Both (A) and (R) are true and (R) is the correct explanation of (A).**

Converting text to lowercase is a key preprocessing step, and the reason correctly explains that it prevents the machine from treating "Hello" and "hello" as different words.

Source: Chapter 6, Section 6.5 – Text Processing (Converting Text to a Common Case)

Explanation

The textbook explicitly states: "*we convert the whole text into a similar case, preferably lowercase. This ensures that the case sensitivity of the machine does not consider the same words as different just because of different cases.*" The Reason directly and correctly explains the Assertion, so option **(A)** is the right choice. In Assertion-Reason questions, always check if the Reason is a **direct cause** of the Assertion — here it is.

Q11. § 6.2 Applications of Natural Language Processing § 6.4 Chatbots § 6.6 Natural Language Processing: Use Case Walkthrough

[1]

A company wants to analyze customer reviews to understand satisfaction levels. Which NLP application would be most suitable ?

- (A) Text classification
- (B) Sentiment analysis
- (C) Keyword extraction
- (D) Language translation

Previously asked in: 2026 104 Q5 (v)

♦ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding rag

Model Answer

(B) Sentiment Analysis — It analyzes customer reviews to determine whether opinions are positive, negative, or neutral, making it the most suitable for understanding customer satisfaction levels.

Source: Applications of Natural Language Processing, Chapter 6

Explanation

Examiners expect you to directly choose **(B) Sentiment Analysis** and justify it briefly in one line. Remember: Sentiment Analysis specifically detects emotions/opinions in text (positive/negative/neutral), which directly maps to "satisfaction levels." Text classification categorizes documents, keyword extraction finds key terms, and language translation converts languages — none specifically measure satisfaction.

Q12. § Test Yourself and Reflection Time

[2]

How is Stemming different from Lemmatization ? Explain how the word "Wolves" would be processed by stemming and lemmatization.

Previously asked in: 2026 104 Q16

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding rag

Model Answer

Stemming removes affixes from words to get the root/base form, but the result may not be a meaningful word.

Lemmatization also removes affixes, but always produces a meaningful word (called a lemma).

Processing "Wolves":

- **Stemming** → *wolv* (not a meaningful word)
- **Lemmatization** → *wolf* (a meaningful word)

Source: Chapter 6, Section 6.5 – Text Processing (Stemming & Lemmatization)

Explanation

- The key distinction examiners look for: stemming is **faster but may give non-meaningful results**; lemmatization is **slower but always gives a valid word**.
- Use a clear example showing both outputs — "wolves" is a standard example. Writing just definitions without an example will cost marks on a 2-mark question.
- The term **lemma** for the output of lemmatization is good to include if space allows.

Q13. § Unit Overview and Learning Objectives § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time [1]

A corpus contains 4 documents in which the words such as 'an, is, the' were appearing frequently. Identify the term that is used for such words.

- (A) Stop word
- (B) Rare word
- (C) Missing word
- (D) Removable word

Previously asked in: 2024 104 Q2 (vi)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

(A) Stop word

Words like 'an', 'is', 'the' that appear frequently in a corpus but add negligible value to the text are called **Stop words**.

Explanation

As per the textbook (Chapter 6), stop words are defined as words with frequent occurrence in the corpus that have negligible value and are often removed during preprocessing. The MCQ option (A) directly matches this definition. Students must not confuse stop words with "removable words" — the correct technical term is **stop word**.

Q14. § 6.2 Applications of Natural Language Processing § 6.3 Stages of Natural Language Processing (NLP) § 6.6 Natural Language

Processing: Use Case Walkthrough

[1]

Which application of NLP helps to provide an overview of a news item or blog post ? It also avoids redundancy from multiple sources and maximises the diversity of content obtained.

- (A) Virtual Assistants
- (B) Sentiment Analysis
- (C) Text Classification
- (D) Automatic Summarization

Previously asked in: 2024 104 Q3 (v)

♦ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

(D) Automatic Summarization

Explanation

Automatic Summarization provides an overview of news items or blog posts, avoids redundancy from multiple sources, and maximises content diversity — all key features that distinguish it from the other options. The source passages list it as a distinct NLP application separate from Text Classification (which categorises documents) and Sentiment Analysis (which detects opinion/emotion).

Q15. § Unit Overview and Learning Objectives § Test Yourself and Reflection Time

[1]

Mention any two commonly used applications of NLP.

Previously asked in: 2022 104 Q8

◆ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

Two commonly used applications of NLP are: **(i) Voice Assistants** (e.g., Google Assistant, Siri) that understand and process natural speech, and **(ii) Language Translation** (e.g., Google Translate) that converts text from one language to another.

Source: Applications of Natural Language Processing, chapter 6, section 6.2

Explanation

For a 1-mark question, name any **two** applications clearly — one line is enough. The source lists six applications: Voice Assistants, Autogenerated Captions, Language Translation, Sentiment Analysis, Text Classification, and Keyword Extraction. Pick any two and briefly state what they do. Avoid writing long descriptions; the examiner awards the mark for correct identification.

Q16. § 6.3 Stages of Natural Language Processing (NLP)

[1]

Name the process of dividing whole corpus into sentences.

Previously asked in: 2022 104 Q10

◆ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

The process of dividing the whole corpus into sentences is called **Lexical Analysis**.

Source: Stages of Natural Language Processing (NLP), chapter 6

Explanation

Lexical Analysis is the first stage of NLP where large chunks of text are broken down into paragraphs, sentences, and words. Examiners expect the exact term "Lexical Analysis" for full marks.

Q17. § 6.3 Stages of Natural Language Processing (NLP) § 6.6 Natural Language Processing: Use Case Walkthrough

[1]

Which of the following applications of NLP (Natural Language Processing) is associated with spam filtering in e-mails ?

- (A) Virtual Assistants
- (B) Sentiment Analysis
- (C) Text Classification
- (D) Automatic Summarization

Previously asked in: 2024 104 Q4 (v)

♦ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(C) Text Classification**

Text Classification categorizes documents into predefined groups, which is how spam filters sort emails into "spam" or "not spam" categories.

Explanation

The textbook defines Text Classification as a tool that "classifies a sentence or document category-wise" into predefined groups. Spam filtering works by classifying incoming emails into categories (spam/not spam), making Text Classification the correct answer. Virtual Assistants relate to voice processing, Sentiment Analysis detects emotions/opinions, and Automatic Summarization condenses text — none of these match spam filtering.

Q18. § Unit Overview and Learning Objectives § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time

[1]

In the context of NLP, which of the following words represents a stem resulting from stemming for "Studies" ?

- (A) Study
- (B) Stud
- (C) Studi
- (D) Studied

Previously asked in: 2024 104 Q5 (iv)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(C) Studi**

Stemming chops off suffixes mechanically without checking meaning. For "Studies," the stemmer removes the suffix, producing "**Studi**" — not a real word, unlike lemmatization which gives "Study."

Explanation

Stemming is a crude, rule-based process that strips affixes regardless of whether the result is a valid word. "Studi" is the classic example of stemming output for "Studies." This is also what distinguishes stemming from lemmatization — lemmatization would correctly return "Study." Examiners specifically test this distinction, so remember: stemming → "Studi" (may be meaningless); lemmatization → "Study" (meaningful root word).

Q19. § Test Yourself and Reflection Time**[2]**

What is Tokenization? Count how many tokens are present in the following statement:

I find that the harder I work, the more luck I seem to have.

Previously asked in: 2022 104 Q15

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

Tokenization is a step in Text Normalisation where each sentence is divided into smaller units called **tokens**. A token can be any word, number, or special character occurring in a sentence.

Counting tokens in: *I find that the harder I work, the more luck I seem to have.*

Tokens: I / find / that / the / harder / I / work / , / the / more / luck / I / seem / to / have / .

Total number of tokens = 16

Source: Text Normalisation – Tokenization, Chapter 6

Explanation

- Define tokenization clearly (1 mark) — mention that tokens include words AND special characters/punctuation.
- Count carefully (1 mark) — the comma (,) and the full stop (.) are also separate tokens. Students often miss punctuation marks and lose the mark. Total = 16 tokens.

Q20. § Test Yourself and Reflection Time**[2]**

Kaira, a beginner in the field of NLP is trying to understand the process of Stemming. Help her in filling up the following table by suggesting appropriate affixes and stem of the words mentioned there:

Previously asked in: 2022 104 Q16

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 · 2026-06-21 03:18 · grounding rag

Model Answer

S. No.	Word	Affixes	Stem
i.	Tries	-es	Try
ii.	Learning	-ing	Learn

Stemming removes affixes (prefixes/suffixes) from a word to obtain its root/stem form.

- *Tries* → suffix **-es** is removed → Stem: **Try**
- *Learning* → suffix **-ing** is removed → Stem: **Learn**

Source: Chapter 6, Text Processing / NLP Stages

Explanation

- Stemming is the process of stripping affixes (usually suffixes) from a word to get its base/stem. The stem may not always be a meaningful word (unlike lemmatization).
- For board exams, clearly identify the **affix** (e.g., -es, -ing) and the resulting **stem**.
- The examiner looks for the correct identification of suffix and the correct stem — 1 mark per row (1/2 for affix + 1/2 for stem, or 1 mark each row depending on the scheme).

Q21. § Unit Overview and Learning Objectives § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time **[4]**

With reference to NLP, explain the following terms in detail with the help of suitable example:

- Term frequency
- Inverse Document Frequency

Previously asked in: 2022 104 Q20

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

Term Frequency (TF):

Term Frequency measures how often a word appears in a document. It helps identify how important a word is within a single document.

Formula: $TF = (\text{Number of times word appears in document}) / (\text{Total words in document})$

Example: In Document: "Johnny Johnny Yes Papa" — TF of "Johnny" = $2/4 = 0.5$

Inverse Document Frequency (IDF):

IDF measures how rare or valuable a word is across all documents in the corpus. Common words (like "the") get a low IDF; rare words get a high IDF.

Formula: $IDF = \log(\text{Total documents} / \text{Documents containing the word})$

Example: If "Papa" appears in 3 out of 4 documents — IDF of "Papa" = $\log(4/3) \approx 0.125$

TFIDF = TF × IDF — words that are frequent in one document but rare across the corpus are most valuable.

Source: Chapter 6, Section 6 (TFIDF concept)

Explanation

- **TF** focuses on a single document; **IDF** focuses on the entire corpus — examiners expect both distinctions to be clear.
- Always include the formula and a numeric example for a competency-based question — that's what earns full marks.
- **TFIDF = TF × IDF** is a bonus line that shows you understand the combined purpose; include it if word count allows.
- Keep examples consistent (use the same corpus for both TF and IDF if possible) to show connected understanding.

Q22. § 6.5 Text Processing § Test Yourself and Reflection Time

[2]

What is the primary difference between Human Language and Computer Language ?

Previously asked in: 2024 104 Q15

◆ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 · 2026-06-21 03:19 · grounding rag

Model Answer**Human Language vs Computer Language:**

Human languages (natural languages) are complex, redundant, context-dependent, and change over time. They use syntax, semantics, and lexicon to convey meaning. Computer languages (e.g., Python, C++, Java) are numerical/binary in nature — structured, precise, and follow strict instructions. Computers cannot directly understand natural language; NLP is needed to bridge this gap.

Source: Chapter 6, Section 6.1 – Introduction

Explanation

- Examiners expect **two clear points of contrast** for 2 marks — one about human language features, one about computer language features.
- Key terms to include: **natural language, complex/context-dependent, numerical/binary, NLP.**
- Avoid writing a long paragraph; two compact sentences covering both sides is ideal for 2 marks.
- The mention of NLP as the bridge adds value and shows understanding of the chapter's theme.

Q23. § Unit Overview and Learning Objectives**[4]**

Document 1 : NLP is a domain of AI.

Document 2 : NLP stands for Natural Language Processing.

Consider the following documents : Implement all the four steps of Bag of Words (BoW) model to create a document vector table.

Previously asked in: 2024 104 Q20

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding stimulus

Model Answer**Step 1 – Collect Data (Documents):**

- Doc 1: "NLP is a domain of AI."
- Doc 2: "NLP stands for Natural Language Processing."

Step 2 – Tokenisation (create unique word list/vocabulary):

NLP, is, a, domain, of, AI, stands, for, Natural, Language, Processing

(Total 11 unique words)

Step 3 – Create Document Vectors (word frequency count):

Word	Doc 1	Doc 2
NLP	1	1
is	1	0
a	1	0
domain	1	0
of	1	0
AI	1	0
stands	0	1
for	0	1
Natural	0	1
Language	0	1
Processing	0	1

Step 4 – Apply BoW Model:

Each document is represented as a vector of word frequencies:

- Doc 1 → [1,1,1,1,1,0,0,0,0,0]
- Doc 2 → [1,0,0,0,0,0,1,1,1,1]

Source: AI Chapter, Natural Language Processing – Bag of Words

Explanation

Examiners award 1 mark per step. Ensure all four steps are clearly labelled. The vocabulary must list **unique** words only. The document vector table (frequency count) is the core deliverable — write it neatly. Final vectors in bracket notation reinforce Step 4 and show you understand the BoW output format.

Q24. § 6.2 Applications of Natural Language Processing § 6.4 Chatbots § 6.6 Natural Language Processing: Use Case Walkthrough

[1]

Email filters, spam filters, smart assistants are the examples of:

- (a) Pocket Assistants
- (b) CV
- (c) NLP
- (d) Evaluation

Previously asked in: 2023 104 Q3 (i)

♦ Natural Language Processing 6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

(c) NLP

Email filters, spam filters, and smart assistants are real-world applications of **Natural Language Processing (NLP)**.

Explanation

The source passages (Chapter 6, Section 6.2) list voice assistants (smart assistants like Siri, Alexa, Google) as key NLP applications. Email/spam filters also use NLP to classify and process text. Examiners expect you to directly identify the correct option and briefly justify it in one line.

Q25. § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time

[1]

For _____ the whole corpus is divided into sentences. Each sentence is taken as a different data so now the whole corpus gets reduced to sentences.

- (a) Text Regulation
- (b) Sentence Segmentation
- (c) Tokenisation
- (d) Stemming

Previously asked in: 2023 104 Q3 (iii)

♦ Natural Language Processing 6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

(b) Sentence Segmentation

Under Sentence Segmentation, the whole corpus is divided into sentences, and each sentence is taken as different data, reducing the corpus to sentences.

Source: Text Processing, Section 6.5

Explanation

The passage in Section 6.5 directly defines Sentence Segmentation as the step where "the whole corpus is divided into sentences" — the question uses this exact language. Tokenisation comes *after* segmentation (it breaks sentences into individual tokens). Students must not confuse the two steps.

Q26. § 6.1 Introduction § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time

[1]

With reference to NLP, consider the following plot of occurrence of words versus their value: In the given graph, X represents:

- (a) Rare / valuable words
- (b) Punctuation words
- (c) Popular words
- (d) Pronoun

Previously asked in: 2023 104 Q3 (vi)

♦ Natural Language Processing 6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

(a) Rare / valuable words

X represents **rare / valuable words** — words that occur the least in the corpus but add the most value to the text analysis.

Explanation

The graph follows Zipf's Law distribution in NLP. Stop words (high occurrence, low value) appear at the top-left; frequent words are in the middle; and at the bottom-right (lowest occurrence), point X marks rare/valuable words — they appear least but carry the most meaning. The textbook (Test Yourself Q8) confirms: *"They occur the least but add the most value to the corpus."* Examiners expect students to read the graph direction correctly and link low occurrence with high value.

Q27. § Unit Overview and Learning Objectives § 6.6 Natural Language Processing: Use Case Walkthrough § Test Yourself and Reflection Time

[1]

Which of the following is a feature of document classification?

- (a) Helps in classifying the type and genre of a document.
- (b) Helps in creating a document.
- (c) Helps to display important information of a corpus.
- (d) Helps in including the necessary words in the text body.

Previously asked in: 2023 104 Q4 (i)

♦ Natural Language Processing 6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

(a) Helps in classifying the type and genre of a document.

Text Classification (document classification) classifies a sentence or document category-wise into predefined groups or categories.

Explanation

The source passage under **Applications of NLP** describes *Text Classification* as a tool that "classifies a sentence or document category-wise" into predefined groups — i.e., it classifies the type and genre of a document. The other options describe unrelated NLP tasks (corpus display, document creation, word inclusion), so (a) is the only correct answer.

Q28.

[1]

Which of the following words represents an example of stemming for the word 'Sharing'?

- (A) Share
- (B) Shared
- (C) Shares
- (D) Shar

Previously asked in: 2025 104/S Q2 (iv)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(D) Shar**

Stemming cuts off affixes to get a base/root form, which may not be a real word. For 'Sharing', stemming removes '-ing' to give **Shar**, not a valid word.

Explanation

Stemming is a crude, rule-based process that simply chops off suffixes. It does **not** guarantee a meaningful word (unlike lemmatization, which gives 'Share'). 'Shar' is the classic stemmed output of 'Sharing'. Examiners expect you to distinguish stemming (may produce non-words) from lemmatization (always produces valid base words).

Q29. § Test Yourself and Reflection Time

[1]

_____ is a term used for any word or number or special character occurring in a sentence. (Token / Punctuator)

Previously asked in: 2023 104 Q5 (i)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

Token is a term used for any word or number or special character occurring in a sentence.

Source: Chapter 6, Section 6.3 / Text Normalisation – Tokenisation

Explanation

The examiner expects the single correct term "Token." In tokenisation (a step of Text Normalisation in NLP), the text is broken into smaller units called **tokens** — these can be words, numbers, or special characters. Do not write "Punctuator," as that refers only to punctuation marks, not all types of units in a sentence.

Q30. § Unit Overview and Learning Objectives § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time [1]

In the sentence 'She reads the book', which of the following is a stop-word that should be removed during text preprocessing?

- (A) She
- (B) reads
- (C) the
- (D) book

Previously asked in: 2025 104/S Q3 (vi)

♦ Natural Language Processing 6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

(C) the

"The" is a stop word — a common word with negligible meaning that is removed during text preprocessing, as it adds no significant value to text analysis.

Explanation

Stop words are high-frequency words (like *the*, *is*, *a*, *an*, *of*) that carry little meaningful information and are removed during text normalisation. The source (Q.7, Test Yourself) defines stop words as "words with negligible value that are often removed during preprocessing." In the sentence, *she* and *reads* and *book* carry meaning; only *the* is a typical stop word.

Q31. § Unit Overview and Learning Objectives § 6.3 Stages of Natural Language Processing (NLP) § 6.6 Natural Language Processing: Use [1]

Case Walkthrough

In Natural Language Processing (NLP), _____ occur/s very frequently in the corpus but do/does not add any value to it.

- (A) Text Normalisation
- (B) Stop words
- (C) Start words
- (D) Tokenisation

Previously asked in: 2025 104/S Q4 (i)

♦ Natural Language Processing 6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

(B) Stop words

Stop words occur very frequently in the corpus (e.g., "is", "the", "and") but do not add any meaningful value to it.

Explanation

The source passage (Test Yourself, Q.7) defines stop words as "words with negligible value that are often removed during preprocessing." They appear frequently but carry no significant meaning, making **(B)** the correct answer. Text Normalisation and Tokenisation are processes, not word types; "Start words" is not an NLP term.

Q32. § Test Yourself and Reflection Time

[1]

The first step of Bag of Words algorithm is Text Normalisation. Which of the following task is done in this step?

- (A) Creating document vectors
- (B) Collecting and pre-processing data
- (C) Adding the words to a dictionary
- (D) Creating a vector of words

Previously asked in: 2025 104/S Q4 (iv)

♦ Natural Language Processing

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(B) Collecting and pre-processing data**

In the Bag of Words algorithm, the first step is **Text Processing**, which involves collecting data and pre-processing it (text normalisation).

Explanation

The source passage explicitly lists the steps of Bag of Words: Step 1 is "Collecting data and pre-processing it," Step 2 is creating a dictionary, and Step 3–4 involve creating document vectors. Examiners expect you to directly identify option (B) as the answer. Do not confuse "Text Normalisation" (the broader NLP concept) with the first *step of the BoW algorithm*, which is data collection and pre-processing.

Q33. § 6.2 Applications of Natural Language Processing § 6.6 Natural Language Processing: Use Case Walkthrough

[1]

Sentiment analysis of customer reviews on various online stores is an example of _____.

- (A) Machine Learning
- (B) Computer Vision
- (C) Natural Language Processing (NLP)
- (D) Speech Recognition

Previously asked in: 2025 104/S Q5 (i)

♦ Natural Language Processing

6.2 Applications of NLP

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(C) Natural Language Processing (NLP)**

Sentiment analysis of customer reviews involves analysing textual data to detect positive, negative, or neutral opinions — which is an application of NLP.

Source: Chapter 6, Section 6.2 Applications of Natural Language Processing

Explanation

- Examiners expect you to directly state the correct option with a one-line justification.
- Sentiment Analysis is explicitly listed as an NLP application in the textbook (Section 6.2), making **(C)** the only correct answer.
- Do **not** confuse it with Machine Learning (which is broader) or Speech Recognition (which deals with audio, not text reviews).

Q34. § 6.1 Introduction**[1]**

Consider the following sentence:

On seeing her son's result, Pooja's face turned red with anger.

The word "red" demonstrates which characteristic of natural language ?

- (A) Redundancy
- (B) Context-dependent meaning
- (C) Grammatical structure
- (D) Temporal change

Previously asked in: 2026 104 Q2 (ii)

♦ **Natural Language Processing**

Characteristics of Natural Language – Ambiguity (Polysemy/Context-dependent meaning)

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding rag**Model Answer****(B) Context-dependent meaning**

The word "red" here means anger, not a colour. Its meaning changes based on context, which is a key characteristic of natural language.

Explanation

The source passage directly uses a similar example — "*His face turned red after he found out...*" — to illustrate that the word "red" can mean shame, anger, or a physical reaction depending on context. Examiners expect students to identify this as **context-dependent meaning**. The other options (redundancy, grammatical structure, temporal change) are also features of natural language but do not apply here.

Q35.**[1]**

Two popular examples of pocket assistants are _____ and _____.

Previously asked in: 2023 104 Q2 (i)

♦ **Natural Language Processing**

Pocket Assistants / Virtual Assistants (e.g., Siri, Google Assistant)

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag**Model Answer**

Two popular examples of pocket assistants (voice assistants) are **Google Assistant (Hey Google)** and **Siri (Hey Siri)**.

Source: Chapter 6, Section 6.2 – Applications of Natural Language Processing

Explanation

The source passage lists voice assistants such as **Google, Alexa, and Siri** as examples. "Pocket assistants" refers to voice/virtual assistants on personal devices. Acceptable answers include any two from: Google Assistant, Siri, or Alexa. Examiners expect recognisable names — just naming two correctly earns the mark.

Q36. § Unit Overview and Learning Objectives**[4]**

Create a document vector table from the following documents by implementing all the four steps of Bag of words model. Also depict the outcome of each step.

Document 1: Sameera and Sanya are classmates.

Document 2: Sameera likes dancing but Sanya loves to study mathematics.

Previously asked in: 2023 104 Q19

♦ Natural Language Processing

6.3 Stages of NLP

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**Step 1: Text Processing (Pre-processing)**

Remove stop words (and, are, but, to), convert to lowercase.

- Doc 1: [sameera, sanya, classmates]
- Doc 2: [sameera, likes, dancing, sanya, loves, study, mathematics]

Step 2: Create Dictionary (Vocabulary)

List all unique words from both documents:

sameera	sanya	classmates	likes	dancing	loves	study	mathematics
-----	-----	-----	-----	-----	-----	-----	-----

Step 3: Create Document Vector for Doc 1

sameera	sanya	classmates	likes	dancing	loves	study	mathematics
-----	-----	-----	-----	-----	-----	-----	-----
1	1	1	0	0	0	0	0

Step 4: Document Vector Table for All Documents

sameera	sanya	classmates	likes	dancing	loves	study	mathematics
-----	-----	-----	-----	-----	-----	-----	-----
1	1	1	0	0	0	0	0
1	1	0	1	1	1	1	1

Source: Chapter 6, Section 6.5 – Bag of Words

Explanation

- Examiners look for **all four steps clearly labelled** with the outcome of each step shown.
- Stop words removed in Step 1: *and, are, but, to* — these are grammar words with no informational value.
- The dictionary must contain **only unique words** (no repetitions).
- In the final table, each row = one document; each cell = frequency of that word in that document (0 if absent, 1 if present once).
- Keep the table neat — marks are awarded for correct 0/1 entries.

Q37. § Unit Overview and Learning Objectives

[4]

Document 1 : CV is an upcoming field.

Document 2 : Image Feature is an important part of CV.

You have two documents :

Document 1 : CV is an upcoming field.

Document 2 : Image Feature is an important part of CV.

Implement all four steps of the Bag of Words (BoW) model to create a document vector table. Depict the outcome of each step.

Previously asked in: 2025 104/S Q19

◆ Natural Language Processing

6.3 Stages of NLP

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding stimulus

Model Answer**Step 1 – Collect Training Data (Documents):**

- Doc 1: "CV is an upcoming field."
- Doc 2: "Image Feature is an important part of CV."

Step 2 – Design the Vocabulary (unique words):

{ CV, is, an, upcoming, field, Image, Feature, important, part, of }

(Stop words retained; total = 10 unique words)

Step 3 – Create Document Vectors (word frequency count):

| Word | CV | is | an | upcoming | field | Image | Feature | important | part | of |

| --- | --- | --- | --- | --- | --- | --- | --- | --- | --- | --- |

| Doc 1 | 1 | 1 | 1 | 1 | 1 | 0 | 0 | 0 | 0 | 0 |

| Doc 2 | 1 | 1 | 1 | 0 | 0 | 1 | 1 | 1 | 1 | 1 |

Step 4 – Use Vectors for ML Model:

Doc 1 → [1,1,1,1,1,0,0,0,0,0]

Doc 2 → [1,1,1,0,0,1,1,1,1,1]

These numerical vectors are fed into a Machine Learning model for text classification or analysis.

Explanation

- **Examiners look for:** All four steps clearly labelled with correct output at each step.
- The vocabulary must list only **unique** words across both documents.
- Frequency counts must be accurate — CV appears once in each document.
- The final vector representation ties the answer together; don't skip it.
- Roughly 1 mark per step — label them clearly to ensure full credit.

Q38. § Unit Overview and Learning Objectives § Test Yourself and Reflection Time

[4]

Consider the following two documents :

Document 1 : ML and DL are part of AI.

Document 2 : DL is a subset of ML.

Implement all four steps of the Bag of Words (BoW) model to create a document vector table. Depict the outcome of each step.

Previously asked in: 2024 104 Q19

♦ Natural Language Processing

6.3 Stages of NLP

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**Step 1: Text Processing (Pre-processing)**

Remove stop words (is, a, of) and convert to lowercase.

- Doc 1: [ml, dl, part, ai]
- Doc 2: [dl, subset, ml]

Step 2: Create a Dictionary (Vocabulary)

List all unique words from both documents:

```
| ml | dl | part | ai | subset |
|----|----|-----|----|-----|
```

Step 3: Create Document Vector for Doc 1

```
| ml | dl | part | ai | subset |
|----|----|-----|----|-----|
| 1 | 1 | 1 | 1 | 0 |
```

Step 4: Create Document Vectors for All Documents

```
| Document | ml | dl | part | ai | subset |
|-----|----|----|-----|----|-----|
| Doc 1 | 1 | 1 | 1 | 1 | 0 |
| Doc 2 | 1 | 1 | 0 | 0 | 1 |
```

Source: Chapter 6, Section 6.5 – Bag of Words

Explanation

- Examiners award 1 mark per step, so label each step clearly.
- Step 1 must show stop-word removal ("and", "is", "a", "of") and lowercasing — these are the pre-processing actions visible to the examiner.
- Step 2: the dictionary contains only **unique** words (5 words here).
- Steps 3 & 4 must show the actual frequency table; since no word repeats within a single document here, all values are 0 or 1.
- Do not skip labelling rows as Doc 1 / Doc 2 — that shows you understand "document vectors."

Q39. § Unit Overview and Learning Objectives

[4]

Document 1 : Data Science requires information.

Document 2 : Information analysis requires data.

Implement all the four steps of Bag of Words (BoW) model to create a document vector table.

Previously asked in: 2026 104 Q21

♦ Natural Language Processing

6.3 Stages of NLP

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding stimulus

Model Answer**Step 1 – Collect Documents:**

- Doc 1: "Data Science requires information."
- Doc 2: "Information analysis requires data."

Step 2 – Create Vocabulary (unique words, lowercase):

{data, science, requires, information, analysis}

Step 3 – Create Word Frequency Vectors:

			data		science		requires		information		analysis	
		---		---		---		---		---		---
Doc 1		1		1		1		1		0		
Doc 2		1		0		1		1		1		

Step 4 – Document Vector Table ready.

Each document is represented as a numeric vector based on word frequency.

Source: Bag of Words Model, Document Representation

Explanation

- Examiners look for **all four steps clearly labelled**: collect text → build vocabulary → count frequencies → form vector table.
- The table is the most important part — ensure rows = documents, columns = vocabulary words.
- Treat words as **case-insensitive** and ignore punctuation.
- "requires" appears in both documents → frequency = 1 in each row.
- Do not skip the vocabulary listing step; it earns separate marks.

Q40. § 6.3 Stages of Natural Language Processing (NLP)**[1]**

What is NLP?

Previously asked in: 2022 104 Q7

♦ Natural Language Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

NLP (Natural Language Processing) is a branch of Artificial Intelligence that enables computers to understand, interpret, and process human (natural) language.

Source: Chapter 6, Unit Overview

Explanation

For a 1-mark question, one clear definitional sentence is sufficient. Examiners look for two key ideas: (1) it is an AI technique, and (2) it deals with human/natural language understanding and processing. Avoid listing applications or stages — that exceeds the mark allocation.

Q41. § Unit Overview and Learning Objectives § 6.3 Stages of Natural Language Processing (NLP) § Test Yourself and Reflection Time**[2]**

Explain the following picture which depicts one of the processes on NLP. Also mention the purpose which will be achieved by this process.

Previously asked in: 2022 104 Q18

♦ Natural Language Processing

6.3 Stages of NLP

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

The diagram depicts **Converting Text to a Common Case**, a step in Text Normalisation (Text Processing) in NLP.

In this process, all words — regardless of their capitalisation (e.g., HELLO, HeLlo, HELLo) — are converted to **lowercase** (hello).

Purpose: This ensures uniformity in word representation, so that the machine does not treat the same word written in different cases as different words, thereby improving accuracy of text analysis.

Source: Chapter 6, Section 6.5 – Text Processing (Text Normalisation)

Explanation

- The examiner wants you to **name the step** (Converting Text to Common Case / Lowercase Conversion) and **state its purpose** — both are needed for full 2 marks.
- Quote from the textbook: *"This ensures that the case sensitivity of the machine does not consider the same words as different just because of different cases."*
- Don't confuse this with Tokenisation or Stop Word Removal — those are different steps. This specifically deals with **case uniformity**.

Q42. § Unit Overview and Learning Objectives**[4]**

Consider the text of following documents:

Document 1: Sahil likes to play cricket

Document 2: Sajal likes cricket too

Document 3: Sajal also likes to play basketball

Apply all the four steps of Bag of words model of NLP on the above given documents and generate the output.

Previously asked in: 2022 104 Q19

♦ Natural Language Processing

6.3 Stages of NLP

6.5 Text Processing

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer**Step 1: Collect Data (Documents)**

- Doc 1: Sahil likes to play cricket
- Doc 2: Sajal likes cricket too
- Doc 3: Sajal also likes to play basketball

Step 2: Create a List of Unique Words (Vocabulary)

{sahil, likes, to, play, cricket, sajal, too, also, basketball}

(9 unique words)

Step 3: Remove Stop Words (optional normalisation)

Stop words like "to", "too", "also" may be removed → Vocabulary: {sahil, likes, play, cricket, sajal, basketball}

Step 4: Create Document Vectors (Frequency Table)

Word	Doc1	Doc2	Doc3
sahil	1	0	0
likes	1	1	1
play	1	0	1
cricket	1	1	0
sajal	0	1	1
basketball	0	0	1

Each document is now represented as a numerical vector based on word frequency.

Source: Chapter 6, Bag of Words Model

Explanation

- Examiners expect **all four steps clearly labelled**: data collection → vocabulary creation → stop word removal → document vector/frequency table.
- The frequency table is the key output; missing it will cost marks.
- You don't need to calculate TF-IDF here — just the BoW frequency table.
- Stop word removal is considered one of the steps in this model as taught in the chapter; include it even briefly.

Q43. § 6.4 Chatbots

[2]

Define Chatbot. What are its types?

Previously asked in: 2023 104 Q13

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

A **chatbot** is a computer program designed to simulate human conversation through voice commands, text chats, or both. It can learn over time how to best interact with humans.

Types of Chatbots:

1. **Script-bot** – A traditional, scripted chatbot that follows predefined rules.
2. **Smart-bot** – An AI-powered chatbot with greater knowledge and learning ability.

Source: Chapter 6, Section 6.4 – Chatbots

Explanation

- The definition must mention "simulate human conversation" and "voice/text."
- Two types must be named: **Script-bot** and **Smart-bot** — examiners expect both terms exactly as given in the textbook.
- For 2 marks, one mark is typically for the definition and one for the types — keep it concise, no extra elaboration needed.

Q44. § 6.2 Applications of Natural Language Processing § 6.4 Chatbots

[1]

Which domain of AI is used for interacting with virtual assistants such as Siri and Alexa ?

- (a) Machine Learning (ML)
- (b) Computer Vision (CV)
- (c) Natural Language Processing (NLP)
- (d) Technical Vision (TV)

Previously asked in: 2024 104 Q5 (vi)

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer**(c) Natural Language Processing (NLP)**

Voice assistants like Siri and Alexa take natural speech, process it, and execute tasks — this is done using Natural Language Processing (NLP).

Explanation

The passage explicitly states: "Hey Alexa, play some music" and "Hey Siri, what's the weather today" as examples of **Voice Assistants**, listed under **Applications of Natural Language Processing**. NLP enables machines to understand and respond to human speech. ML is a broader domain; CV deals with images; "Technical Vision (TV)" does not exist as an AI domain.

Q45. § 6.4 Chatbots**[1]**

Which type of chat bot requires coding and works on bigger databases directly ?

- (A) Script bot
- (B) Smart bot
- (C) Traditional bot
- (D) Rule-based bot

Previously asked in: 2026 104 Q4 (vi)

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:20 • grounding rag

Model Answer**(B) Smart bot**

A Smart bot requires coding and works on bigger databases directly.

Explanation

The source passage distinguishes two types of chatbots: **Script-bot** (scripted/traditional, rule-based, simpler) and **Smart-bot** (AI-powered, works with larger databases, requires coding). The question describes characteristics of a Smart bot. Options C and D are not the two types named in the passage; option A (Script bot) is the simpler, non-coding type.

Q46. § 6.4 Chatbots**[2]**

What are the primary differences between Script-bots and Smart-bots ?

Previously asked in: 2024 104 Q12

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

Script-bot: A scripted or traditional chatbot that follows pre-defined rules and fixed responses. It can only answer questions within its programmed script and cannot handle queries outside it.

Smart-bot: An AI-powered chatbot that uses machine learning and NLP. It has broader knowledge, learns over time, and can handle a wider variety of conversations more naturally.

Source: Chapter 6, Section 6.4 – Chatbots

Explanation

- The passage directly states that some chatbots are "scripted or traditional" while others are "AI-powered and have more knowledge" — use these exact terms.
- For 2 marks, examiners expect one clear point for each type. Avoid over-explaining.
- Key contrast to highlight: **rule-based vs. AI/learning-based**.

Q47. § 6.2 Applications of Natural Language Processing § 6.4 Chatbots

[1]

Name any two currently popular virtual assistants.

Previously asked in: 2022 104 Q9

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

Two currently popular virtual assistants are **Google Assistant** (Hey Google) and **Siri** (Apple). **Alexa** (Amazon) is also a widely used voice assistant.

Source: Applications of Natural Language Processing, chapter 6, section 6.2

Explanation

The source passage lists three voice assistants as examples under "Voice assistants" in section 6.2 — Google Assistant, Alexa, and Siri. For a 1-mark question, naming any two is sufficient. Avoid over-explaining; just name them clearly.

Q48. § 6.4 Chatbots

[2]

Differentiate between Script-bot and Smart-bot.

Previously asked in: 2025 104/S Q12; 2022 104 Q13 — 2×

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer

Basis	Script-bot	Smart-bot
Nature	Traditional/scripted chatbot	AI-powered chatbot
Knowledge	Follows pre-written scripts; limited responses	Has broader knowledge; learns and adapts over time
Interaction	Feels robotic; cannot handle unknown queries	Feels more human-like; handles varied conversations

Script-bots respond only within fixed, programmed rules, while Smart-bots use Artificial Intelligence to understand context and give more flexible, intelligent responses.

Source: Chapter 6, Section 6.4 — Chatbots

Explanation

- The passage directly states: "some of them are scripted or traditional chatbots while others are AI-powered and have more knowledge" — this is the core distinction examiners expect.
- For a 2-mark question, a simple table OR two clear contrasting points is sufficient. Avoid writing long paragraphs.
- Key terms to use: **scripted/traditional, AI-powered, fixed responses, learns over time.**

Q49. § 6.4 Chatbots

[1]

Which type of chat-bot has a wide functionality, is flexible and powerful, and works on bigger databases directly ?

Previously asked in: 2024 104 Q5 (vi)

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:19 • grounding rag

Model Answer

Smart-bot has wide functionality, is flexible and powerful, and works on bigger databases directly.

Explanation

The question tests knowledge of the two types of chatbots: Script-bot (rule-based, limited) and Smart-bot (AI-powered, flexible, works on large databases). Examiners expect the specific term **Smart-bot** as the answer.

Q50. § 6.4 Chatbots

[1]

Select the correct features of Smart Bot:

- (a) Smart-bots are flexible and powerful
- (b) Coding is required to take this up on board
- (c) Smart bots work on bigger databases and other resources directly
- (d) All of the above

Previously asked in: 2023 104 Q3 (ii)

♦ Natural Language Processing

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer**(a) Smart-bots are flexible and powerful**

The correct answer is **(a)**. Smart-bots are flexible and powerful and work on bigger databases, but **no coding is required** — so option (b) is incorrect, making (d) wrong too.

Explanation

The passage states there are two types of chatbots: Script-bots and Smart-bots. Smart-bots are AI-powered, flexible, and powerful. A key feature of Smart-bots is that **no coding is required** (they are no-code tools), which directly eliminates option (b) and therefore option (d). Examiners expect you to identify the false statement in such MCQs to eliminate "All of the above."

Q51. § 6.2 Applications of Natural Language Processing**[1]**

Smart Assistants such as Alexa, Siri are the examples of:

- (a) Natural Language Processing
- (b) Data Science
- (c) Machine Learning
- (d) Computer Vision

Previously asked in: 2023 104 Q5 (iii)

◆ **Natural Language Processing**

6.2 Applications of NLP

6.4 Chatbots

Generated by claude-sonnet-4-6 • 2026-06-21 03:18 • grounding rag

Model Answer**(a) Natural Language Processing**

Alexa and Siri are voice assistants that leverage NLP to understand natural speech and execute tasks efficiently.

Source: Chapter 6, Section 6.2 – Applications of Natural Language Processing

Explanation

The passage explicitly states: "*Hey Alexa, play some music*" and "*Hey Siri, what's the weather today*" as examples under **Voice Assistants**, which is listed as an application of **Natural Language Processing**. In MCQs, always look for the exact term used in the textbook against the given examples.

Available for free from:

<https://cbsegrade10studyguide.com>

<https://github.com/orgs/cbse-free-resources/repositories>

Available for free from:

<https://cbsegrade10studyguide.com>

<https://github.com/orgs/cbse-free-resources/repositories>