

生成AIパスポート 対策テキスト

# 聞き流し版

公式シラバスのキーワード194語を、一問一答で。動画の副読本です。

---

全5章20節・194語をシラバスの並びのまま収録

節の頭に自己テスト用のチェックリスト／そのあとに答え

巻末に194語の索引

## AI資格ラボ

YouTubeチャンネル登録者への配布テキストです。

出題範囲は一般社団法人 生成AI活用普及協会（GUGA）の公式シラバスに準拠しています。

掲載している問題・解説はすべて自作であり、試験の過去問ではありません。

聞き流し

## はじめに この本の使い方

この本は、聞き流し版の動画の**副読本**です。動画は「言葉 → 4秒の間 → 答え」を194語ぶん繰り返します。耳だけで進むので、**画面を見る必要はありません**。

紙のほうに向いているのは、**思い出せなかった語だけを、あとから拾い直すとき**です。節の頭にその節の語だけを並べたチェックリストを置きました。**先に自分で説明してみて**、言えなかった語だけ、下の答えを読んでください。

### この本の読み方

- ① チェックリストの語を見て、**声に出して説明してみる**
- ② 言えたら次へ。**言えなかった語だけ**、下の答えを読む
- ③ 一周したら、印を付けた語だけをもう一度

⇒ ⚠ **語の並びは公式シラバスのままです**。覚えやすい順に並べ替えていません。試験範囲のどこにいるかが分かるほうが、あとで探しやすいためです。

⚠ ⚠ **この本だけで合格できるとは書きません**。言葉の意味を思い出せるようにするための本です。解き方は「難問対策編」、章ごとの詳しい説明は各章のテキストにあります。

## 第1章 AI（人工知能）

### 1-1 AI（人工知能）の定義（2語）

まず、この語を自分で説明できるか試してください。

- ☐ AI（人工知能） ☐ ダートマス会議

答え

#### AI（人工知能）

人間の知能を人工的に実現しようとする技術・研究分野。専門家の間で合意された定義は存在しない。

#### ダートマス会議

1956年にアメリカのダートマス大学で開かれた研究会です。ジョン・マッカーシーが、ここで「人工知能」という言葉を初めて使いました。

### 1-2 AIに知能をもたらす仕組み（21語）

まず、この語を自分で説明できるか試してください。

- |                                       |                                  |                                     |
|---------------------------------------|----------------------------------|-------------------------------------|
| <input type="checkbox"/> ルールベース       | <input type="checkbox"/> 教師なし学習  | <input type="checkbox"/> 半教師あり学習    |
| <input type="checkbox"/> 機械学習         | <input type="checkbox"/> クラスタリング | <input type="checkbox"/> ノーフリーランチ定理 |
| <input type="checkbox"/> 学習済みモデル      | <input type="checkbox"/> 次元削減    | <input type="checkbox"/> ニューロン      |
| <input type="checkbox"/> 教師あり学習       | <input type="checkbox"/> 強化学習    | <input type="checkbox"/> シナプス       |
| <input type="checkbox"/> 人工ニューロン（ノード） | <input type="checkbox"/> 重み      | <input type="checkbox"/> 正則化        |
| <input type="checkbox"/> ニューラルネットワーク  | <input type="checkbox"/> 情報の重みづけ | <input type="checkbox"/> ドロップアウト    |
| <input type="checkbox"/> ディープラーニング    | <input type="checkbox"/> 過学習     | <input type="checkbox"/> 転移学習       |

答え

#### ルールベース

人間があらかじめルール（条件と処理）を記述し、そのとおりに動かす方式。判断の理由を説明できる一方、書けるルールの量に限界がある。

#### 機械学習

人間がルールを書かず、大量のデータからコンピュータ自身に規則性を見つけさせる方式。

#### 学習済みモデル

機械学習によって学習を終えた成果物。新しいデータを入力すると答えを返す。

## 教師あり学習

正解付きのデータで学習させるやり方です。「この写真は猫」と答えを添えて大量に見せます。

## 教師なし学習

正解を与えずに、データの構造を勝手に見つけさせるやり方です。クラスタリングと次元削減が代表です。

## クラスタリング

似たものどうしをまとまりに分ける処理です。教師なし学習のひとつ。分けたグループが何を意味するかの解釈は、人の仕事です。

## 次元削減

多数の変数を、情報をできるだけ保ったまま、より少ない数の軸にまとめ直す処理です。教師なし学習のひとつ。

## 強化学習

正解は教えず、かわりに報酬を与えるやり方です。うまくいったら点、失敗したら減点。試行錯誤しながら点の高い動きを覚えます。

## 半教師あり学習

正解付きが少しだけ、正解なしが大量、という状況で両方を使うやり方です。現実のデータはたいていこれです。

## ノーフリーランチ定理

あらゆる問題に対して万能に優れたアルゴリズムは存在しないという定理。手法には必ず得手不得手がある。

## ニューロン

人間の脳にある神経細胞のことです。ニューラルネットワークが手本にした、生き物のほうの仕組みです。

## シナプス

ニューロンとニューロンのつなぎ目のことです。ここで信号の伝わりやすさが決まります。

## 人工ニューロン（ノード）

ニューロンを真似て作った、ネットワークの一つひとつの点です。ノードとも呼びます。

## ニューラルネットワーク

人工ニューロンを層状につないだ仕組みです。人間の脳の神経回路を手本にしています。

## ディープラーニング

ニューラルネットワークの層を深く積み重ねた機械学習の手法。層を追うごとに単純な特徴から複雑な特徴へと組み上げていく。

## 重み

つながりの強さを表す数値です。学習とは、この重みを調整していくことです。

## 情報の重みづけ

どの入力をどれだけ重視するかを、重みで決めることです。学習が進むほど、大事な入力のリ重みが大きくなります。

## 過学習

学習データに合わせすぎて、未知のデータに対する性能が落ちてしまう状態。

## 正則化

学習のときに罰則を加えて、モデルが複雑になりすぎるのを抑える方法です。過学習を避ける手法のひとつ。

## ドロップアウト

学習のたびにノードをわざと一部お休みさせる方法です。特定のノードに頼りきりにならず、丸暗記しにくくなります。

## 転移学習

ある目的で学習済みのモデルを別の目的に再利用する手法。少ないデータ・短い時間で学習できる。

## 1-3 AIの種類（3語）

まず、この語を自分で説明できるか試してください。

☐ 特徴量

☐ 弱いAI（ANI）

☐ 強いAI（AGI）

## 答え

### 特徴量

データのうち「どこに注目すれば見分けられるか」という着目点。従来は人間が指定していたが、ディープラーニングではAIが自分で見つける。

### 弱いAI (ANI)

特定の課題に特化したAIです。現在実用化されているAIは、すべて弱いAIにあたります。

### 強いAI (AGI)

汎用的に多様な課題へ対応し、人間のような意識や自我を持つとされるAI。まだ実現していない。

## 1-4 AIの歴史 (8語)

まず、この語を自分で説明できるか試してください。

☐ 第一次AIブーム

☐ 第二次AIブーム

☐ 第三次AIブーム

☐ 探索

☐ エキスパートシステム

☐ ビッグデータ

☐ 推論

☐ AIの冬

## 答え

### 第一次AIブーム

1950年代から60年代のブームです。探索と推論が中心で、迷路やパズルは解けても現実の問題は解けませんでした。

### 探索

考えられる手を順に調べて、答えにたどり着く道を探すやり方です。第一次AIブームの中心でした。

### 推論

既にある知識と規則から、新しい結論を導くことです。探索とならぶ第一次AIブームの柱です。

### 第二次AIブーム

1980年代のブームです。専門家の知識をルールにしたエキスパートシステムが中心でしたが、知識を書き切れず終わりました。

## エキスパートシステム

専門家（エキスパート）の知識をルールとして大量に登録し、専門家のように判断させるシステム。第二次AIブームの主役。

## AIの冬

AIへの期待がしぼみ、研究費や関心が落ち込んだ時期。1970年代と1990年代の2回あった。

## 第三次AIブーム

2010年代からのブームです。扱えるデータの量が増え、計算する力が高まり、ディープラーニングが実用的な成果を出しました。

## ビッグデータ

インターネットの普及などにより蓄積された大量かつ多様なデータ。第三次AIブームを支えた要因の1つ。

## 1-5 シンギュラリティ（技術的特異点）（5語）

まず、この語を自分で説明できるか試してください。

- |                                           |                                     |                                  |
|-------------------------------------------|-------------------------------------|----------------------------------|
| <input type="checkbox"/> シンギュラリティ（技術的特異点） | <input type="checkbox"/> ヴァーナー・ヴィンジ | <input type="checkbox"/> 2045年問題 |
|                                           | <input type="checkbox"/> レイ・カーツワイル  | <input type="checkbox"/> AI効果    |

### 答え

#### シンギュラリティ（技術的特異点）

AIが人間の知能を超え、その先の変化を人間が予測できなくなるとされる時点のことです。到来するかどうかは、研究者の間でも見方が分かれています。

#### ヴァーナー・ヴィンジ

SF作家であり数学者。技術的特異点という概念を広めた人物です。年を示したのはカーツワイルのほうです。

#### レイ・カーツワイル

技術の進歩の速さから、2045年にAIが人間の知能を超えると予測した人物です。これが2045年問題です。

#### 2045年問題

レイ・カーツワイルが2045年にシンギュラリティが到来すると予測したことに由来する呼び名。

## AI効果

AIが何かを実現した途端、人がそれを「知能ではない、ただの処理だ」と見なすようになる現象。

## 第2章 生成AI

### 2-1 生成AIの誕生まで（32語）

まず、この語を自分で説明できるか試してください。

- |                                              |                                                     |                                                        |
|----------------------------------------------|-----------------------------------------------------|--------------------------------------------------------|
| <input type="checkbox"/> 生成AI（ジェネレーティブAI）    | <input type="checkbox"/> CNN（畳み込みニューラルネットワーク）       | <input type="checkbox"/> エンコーダ                         |
| <input type="checkbox"/> ボルツマンマシン            | <input type="checkbox"/> 畳み込み                       | <input type="checkbox"/> デコーダ                          |
| <input type="checkbox"/> 制約付きボルツマンマシン        | <input type="checkbox"/> VAE（変分自己符号化器）              | <input type="checkbox"/> 潜在ベクトル                        |
| <input type="checkbox"/> 自己回帰モデル             | <input type="checkbox"/> ノイズ                        | <input type="checkbox"/> GAN（敵対的生成ネットワーク）              |
| <input type="checkbox"/> 生成器                 | <input type="checkbox"/> リカレント層                     | <input type="checkbox"/> Attention層                    |
| <input type="checkbox"/> 識別器                 | <input type="checkbox"/> シーケンスデータ                   | <input type="checkbox"/> 自己注意力（Self-Attention）         |
| <input type="checkbox"/> RNN（回帰型ニューラルネットワーク） | <input type="checkbox"/> LSTM（長・短期記憶）               | <input type="checkbox"/> Attention Mechanism           |
| <input type="checkbox"/> 隠れ層                 | <input type="checkbox"/> Transformerモデル             | <input type="checkbox"/> 位置エンコーディング                    |
| <input type="checkbox"/> アーキテクチャ             | <input type="checkbox"/> BERTモデル                    | <input type="checkbox"/> NSP（Next Sentence Prediction） |
| <input type="checkbox"/> GPTモデル              | <input type="checkbox"/> MLM（Masked Language Model） | <input type="checkbox"/> RoBERTa                       |
| <input type="checkbox"/> Open AI             |                                                     | <input type="checkbox"/> ALBERT（a Lite BERT）           |

#### 答え

#### 生成AI（ジェネレーティブAI）

文章・画像・音声・動画など、無かったものを作り出すAI。従来のAIは識別、生成AIは生成。

#### ボルツマンマシン

1985年提案。データのでき方を確率で覚えようとした生成モデルの祖先。全部つながっていて計算量が大きく実用にならなかった。

#### 制約付きボルツマンマシン

同じ層の中のつながりを切り、層と層の間だけをつないだもの。計算が現実的になり、深い層の学習を支えた。

#### 自己回帰モデル

自分が出した出力を次の入力に使い、1つずつ順番に生成していくやり方。GPTがこれ。

#### CNN（畳み込みニューラルネットワーク）

畳み込みを使って画像の特徴を取り出すネットワーク。画像を得意とする。

## 畳み込み

小さな窓をずらしながら当て、局所的な特徴を取り出す操作。

## VAE（変分自己符号化器）

入力を縮めて、揺らして、戻すことで新しいデータを作る生成モデル。

## ノイズ

潜在ベクトルにわざと加える揺らぎ。これがあるので元と違うものが出てくる。

## エンコーダ

入力を圧縮する側（入口）。

## デコーダ

圧縮されたものから復元する側（出口）。

## 潜在ベクトル

エンコーダが作る短い数字の並び。特徴を圧縮した表現。

## GAN（敵対的生成ネットワーク）

生成器と識別器を競い合わせて質を上げる生成モデル。

## 生成器

GANの中で偽物を作る側。

## 識別器

GANの中で本物か偽物かを見分ける側。

## RNN（回帰型ニューラルネットワーク）

隠れ層の出力を次の入力に混ぜ、順番のあるデータを扱えるようにしたモデル。

## 隠れ層

入力層と出力層のあいだの層。RNNでは記憶係の役目をする。

## リカレント層

前の状態を次へ戻す構造を持った層。

## シーケンスデータ

文章・音声・株価など、順番に意味があるデータ。系列データ。

## LSTM（長・短期記憶）

RNNに忘れるかどうかを決めるスイッチを足したもの。長い系列でも忘れにくい。

## Transformerモデル

2017年登場。順番に読むのをやめ、文の全部の語を一度に見るモデル。いまの生成AIのほとんどがこの子孫。

## Attention層

どの語に注目するかを重みとして計算する層。

## 自己注意力（Self-Attention）

同じ文の中の語どうして注目し合う仕組み。

## Attention Mechanism

Attentionの仕組みそのものを指す言い方。上の3つは同じものを指す。

## 位置エンコーディング

それぞれの語に「何番目か」を数値として足す仕組み。Transformerが語順を扱うために必要。

## アーキテクチャ

モデルの構造・設計そのもの。

## GPTモデル

Transformerの作る側を取り出したモデル。自己回帰で前から後ろへ一方向に書く。

## Open AI

GPTモデルを開発した組織。

## BERTモデル

Transformerの読む側を取り出したモデル。Googleが開発。前も後ろも同時に見る。

## MLM (Masked Language Model)

文の中の単語を隠して当てさせる学習方法。穴埋め。

## NSP (Next Sentence Prediction)

2つの文が続いているかを当てさせる学習方法。続き当て。

## RoBERTa

BERTをもっと鍛えた版。データを増やし長く学習させ、NSPをやめた。

## ALBERT (a Lite BERT)

BERTの軽い版。パラメータを大幅に減らし、性能を保ったまま小さくした。

## 2-2 ChatGPT (27語)

まず、この語を自分で説明できるか試してください。

- |                                                   |                                      |                                                                            |
|---------------------------------------------------|--------------------------------------|----------------------------------------------------------------------------|
| <input type="checkbox"/> ChatGPT                  | <input type="checkbox"/> パラメータ       | <input type="checkbox"/> GPT-4                                             |
| <input type="checkbox"/> GPT-1                    | <input type="checkbox"/> GPT-3       | <input type="checkbox"/> データセット                                            |
| <input type="checkbox"/> 自然言語処理 (NLP)             | <input type="checkbox"/> InstructGPT | <input type="checkbox"/> RLHF (Reinforcement Learning from Human Feedback) |
| <input type="checkbox"/> GPT-2                    | <input type="checkbox"/> GPT-3.5     | <input type="checkbox"/> アライメント (Alignment)                                |
| <input type="checkbox"/> ファインチューニング               | <input type="checkbox"/> GPTs        | <input type="checkbox"/> GPT-o4                                            |
| <input type="checkbox"/> ハルシネーション (Hallucination) | <input type="checkbox"/> GPT-4o      | <input type="checkbox"/> GPT-4.1                                           |
| <input type="checkbox"/> マルチモーダル                  | <input type="checkbox"/> GPT-o1      | <input type="checkbox"/> GPT-5                                             |
| <input type="checkbox"/> Code Interpreter         | <input type="checkbox"/> GPT-o3      | <input type="checkbox"/> Sora                                              |
| <input type="checkbox"/> Operator                 | <input type="checkbox"/> Codex       | <input type="checkbox"/> Image Generation                                  |

### 答え

## ChatGPT

OpenAIが2022年11月に公開した対話型の生成AI。話し言葉で指示できるようにしたことで、専門知識がなくても使えるようになった。

## GPT-1

2018年。先に大量の文章を読ませ、あとで目的ごとに追加学習させるという2段構えの型を作った最初のGPT。

## 自然言語処理（NLP）

私たちが普段しゃべっている言葉を、コンピュータに扱わせる技術分野。プログラミング言語のような形式的な言葉の反対側。

## GPT-2

2019年。パラメータ約15億。大きくすると賢くなることが見えはじめた。悪用のおそれから段階的に公開された。

## パラメータ

ニューラルネットワークの中で調整される数字の総数。第1章の「重み」が全部で何個あるか、という数え方。

## GPT-3

2020年。パラメータ約1750億。例をいくつか見せるだけで、教えていない仕事もこなせるようになった。

## InstructGPT

2022年。人間の好みを学ばせて、指示に答えるようGPT-3を調整したモデル。ChatGPTの直接の下地。

## GPT-3.5

InstructGPTの流れをくむGPT-3の改良版。これを対話用にして公開したのがChatGPT。

## GPT-4

2023年。ハルシネーションが減り、マルチモーダルになった（文章だけでなく画像も扱える）。

## データセット

学習に使うデータのまとめ。GPT-3ではインターネット上の文章・本・Wikipediaなどが大量に使われた。

## RLHF (Reinforcement Learning from Human Feedback)

人間のフィードバックによる強化学習。人が良し悪しを選び、その選び方を報酬にしてモデルを調整する。

## アライメント (Alignment)

AIの振る舞いを、人間の意図や価値観に沿わせることです。RLHFは、そのための手段のひとつです。

## ファインチューニング

学習済みモデルに追加学習をさせて、特定の用途に合わせること。第1章の転移学習の仲間。

## ハルシネーション (Hallucination)

事実に基づかない、もっともらしい情報を生成してしまう現象。次に来そうな言葉を選んでいて、事実を確かめていないため起きる。

## マルチモーダル

文章・画像・音声など、複数の種類の情報を扱えること。モーダルは様式、マルチは複数。

## Code Interpreter

ChatGPTがプログラムを書いて、その場で実行する機能。集計やグラフ作成まで行う。

## GPTs

目的別のChatGPTを自分で作れる仕組み。プログラムを書かずに用途を絞った相手を用意できる。

## GPT-4o

2024年。oはOmni（すべて）。文章・音声・画像を1つのモデルでまとめて扱い、応答が速くなった。ゼロではない。

## GPT-o1

答える前に内側で手順を組み立ててから答える推論モデルの第一世代。数学やプログラムに強い。

## GPT-o3

o1の系譜を進めた推論モデル。考える力を上げたぶん、返答は遅い。

## GPT-o4

推論モデルの系譜の後の世代。位置づけはo1・o3と同じ「考えてから答える」側。

## GPT-4.1

2025年。開発者向けの色が濃い。長い文章を読める、指示に正確に従う、といった実務寄りの改良。

## GPT-5

2025年。速く答える・必要なら考える・文章も画像も扱う、という別々に育った性質が1つにまとまった。

## Sora

OpenAIの動画生成AI。文章で指示すると映像が出てくる。

## Operator

AIがブラウザを自分で操作するもの。検索・ページ操作・入力を代わりに行う。第3章のAIエージェントにつながる。

## Codex

OpenAIのプログラムを書く側の道具。コードの生成や補完に使う。

## Image Generation

OpenAIの画像生成の機能。文章で指示して絵を作る。

## 2-3 その他の主要生成AI（3語）

まず、この語を自分で説明できるか試してください。

☐ Gemini

☐ Claude

☐ Copilot

### 答え

## Gemini

Googleの生成AI。最初から文章・画像・音声・動画をまとめて扱う設計。

## Claude

Anthropicの生成AI。長い文章を読ませるのが得意で、安全性と受け答えの丁寧さが評価されている。

## Copilot

Microsoftの生成AI。Word・Excel・Teamsの中に組み込まれ、仕事の場所でそのまま使える。

## 第3章 現在の生成AIの動向

### 3-1 生成AIが出来ることと主なサービス（6語）

まず、この語を自分で説明できるか試してください。

- |                                  |                                                |                                  |
|----------------------------------|------------------------------------------------|----------------------------------|
| <input type="checkbox"/> 画像のリサイズ | <input type="checkbox"/> データの水増し（augmentation） | <input type="checkbox"/> リマスタリング |
| <input type="checkbox"/> 正規化     | <input type="checkbox"/> データ拡張技術               | <input type="checkbox"/> Veo3    |

#### 答え

#### 画像のリサイズ

画像の大きさ（縦横のピクセル数）を揃えること。AIは決まった大きさの入力しか受け取れないため、学習の前に揃える。

#### 正規化

データの値の範囲を揃えること。画素の0～255を0～1に直すのが代表例。桁が大きいままだと学習が不安定になるため。

#### データの水増し（augmentation）

手持ちのデータを加工して枚数を増やすこと。反転・回転・明るさ変更・切り取りなど。中身は変わらず、見え方だけ変える。

#### データ拡張技術

データの水増しをするための手口をまとめた呼び名です。水増しが「やること」、データ拡張技術が「その技術の呼び名」です。

#### リマスタリング

古い、あるいは粗いデータを作り直して質を上げることです。解像度を上げる、ノイズを取る、色を戻すなど。水増しが「数」を増やすのに対し、リマスタリングは「質」を上げます。

#### Veo3

Googleの動画生成AI。ヴィオ・スリー。SoraのGoogle版と考えるとよい。第3章で新しく出る語。

### 3-2 ディープフェイク（深層偽造）技術（2語）

まず、この語を自分で説明できるか試してください。

- |                                           |                                            |
|-------------------------------------------|--------------------------------------------|
| <input type="checkbox"/> ディープフェイク（深層偽造）技術 | <input type="checkbox"/> 偽情報（ディスインフォメーション） |
|-------------------------------------------|--------------------------------------------|

## 答え

### ディープフェイク（深層偽造）技術

ディープラーニングを使って本物そっくりの偽物を作る技術。顔の貼り替えや、実在する人の声で喋らせることができる。

### 偽情報（ディスインフォメーション）

意図的に流される嘘の情報。うっかり広まった誤情報とは、悪意があるかどうかで分かれる。

## 3-3 RAG（3語）

まず、この語を自分で説明できるか試してください。

- ☐ RAG (Retrieval-Augmented Generation)      ☐ チャンク      ☐ ベクトルデータベース

## 答え

### RAG (Retrieval-Augmented Generation)

検索拡張生成。答える前に資料を検索し、読ませてから答えさせる仕組み。ハルシネーション対策。モデルは触らない。

### チャンク

資料をあらかじめ細かく切っておいた、その一つひとつ。かたまりの意味。

### ベクトルデータベース

チャンクを意味の数字の並びに変換してしまっておく保管庫。数字なので意味の近さを計算で比べられる。

## 3-4 AIエージェント（5語）

まず、この語を自分で説明できるか試してください。

- ☐ AIエージェント      ☐ Manus      ☐ MCP  
☐ GenSpark      ☐ Skywork AI

## 答え

### AIエージェント

目的を渡すと、手順を自分で考えて、道具を使って、最後までやるAI。第2章のOperatorの正体。

## GenSpark

ジェンスパーク。AIエージェントを使った検索エンジン。自分でいくつも検索して、まとめた答えを作る。

## Manus

マナス。汎用のAIエージェント。調べ物から資料作りまで通しでやる。

## Skywork AI

スカイワーク・エーアイ。汎用型のAIエージェント。資料やスライドを作るのが得意とされる。

## MCP

ModelContextProtocol。Anthropicが出した、AIが外の道具やデータにつながるための共通の決まりごと。

## 第4章 情報リテラシーとAI社会原則

### 4-1 インターネットリテラシー（5語）

まず、この語を自分で説明できるか試してください。

- |                                       |                                        |                                  |
|---------------------------------------|----------------------------------------|----------------------------------|
| <input type="checkbox"/> インターネットリテラシー | <input type="checkbox"/> 情報リテラシー       | <input type="checkbox"/> デジタル市民権 |
| <input type="checkbox"/> テクノロジーの理解    | <input type="checkbox"/> セキュリティとプライバシー |                                  |

#### 答え

##### インターネットリテラシー

インターネットを適切に利用するために必要な力。テクノロジーの理解・情報リテラシー・セキュリティとプライバシー・デジタル市民権の4つを含む親の言葉。

##### テクノロジーの理解

インターネットや情報技術の仕組みを知っていること。知らないものは疑えないため、身を守る土台になる。

##### 情報リテラシー

情報を見極める力。その情報が本当か、誰が発信しているのかを確かめられること。

##### セキュリティとプライバシー

身を守る力。攻撃の手口を知り、設定によって自分の情報を守れること。

##### デジタル市民権

ネット上でもひとりの市民として振る舞うという考え方。権利と責任がともにあるとし、匿名を理由にした無責任な振る舞いを否定する。

### 4-2 セキュリティとプライバシー（11語）

まず、この語を自分で説明できるか試してください。

- |                                   |                                          |                                  |
|-----------------------------------|------------------------------------------|----------------------------------|
| <input type="checkbox"/> フィッシング詐欺 | <input type="checkbox"/> アンチウイルスソフトウェア   | <input type="checkbox"/> ベイト攻撃   |
| <input type="checkbox"/> スミッシング   | <input type="checkbox"/> ランサムウェア         | <input type="checkbox"/> ブラックメール |
| <input type="checkbox"/> ヴィッシング   | <input type="checkbox"/> ソーシャルエンジニアリング攻撃 | <input type="checkbox"/> プレテキスト  |
| <input type="checkbox"/> マルウェア    | <input type="checkbox"/> スピアフィッシング       |                                  |

## 答え

### フィッシング詐欺

偽のメールやWebサイトで本物を装い、ID・パスワード・カード番号などを入力させて盗む詐欺。

### スミッシング

SMS（ショートメッセージ）を使ったフィッシング。SMS+Phishing。宅配便の不在通知を装う手口が代表例。

### ヴィッシング

音声通話を使ったフィッシング。Voice+Phishing。電話で本人に喋らせて聞き出す。

### マルウェア

悪意のあるソフトウェアの総称。malicious+software。ウイルス・ワーム・トロイの木馬・ランサムウェアなどをすべて含む。

### アンチウイルスソフトウェア

マルウェアを検出・遮断・駆除するソフトウェア。名前に反して、ウイルスだけでなくマルウェア全般が対象。

### ランサムウェア

ファイルを暗号化して使えなくし、復旧と引き換えに金銭を要求するマルウェア。ransom=身代金。払っても戻る保証はない。

### ソーシャルエンジニアリング攻撃

技術的な解析ではなく、人間の心理や行動につけこんで情報を盗む攻撃の総称。スパイフィッシング・ベイト攻撃・ブラックメール・プレテキストを含む。

### スパイフィッシング

特定の個人や組織を狙い撃ちにするフィッシング。spear=槍。相手の所属や事情を調べたうえで送るため信じやすい。

### ベイト攻撃

えさで相手を釣る攻撃。bait=えさ。無料・限定といった欲を刺激する誘いで、リンクを踏ませたり機器を使わせたりする。

## ブラックメール

脅迫。弱みを握ったと告げて金銭や情報を要求する。黒いメールという意味ではない。

## プレテキスト

作り話（口実）で身分をいつわり、信用させてから情報を聞き出す手口。上司・取引先・システム管理者などになりきる。

## 4-3 個人情報保護の観点（9語）

まず、この語を自分で説明できるか試してください。

- |                                    |                                    |                                       |
|------------------------------------|------------------------------------|---------------------------------------|
| <input type="checkbox"/> 個人情報保護法   | <input type="checkbox"/> 個人情報取扱事業者 | <input type="checkbox"/> 機微（センシティブ）情報 |
| <input type="checkbox"/> 改正個人情報保護法 | <input type="checkbox"/> 個人識別符号    | <input type="checkbox"/> 匿名加工情報       |
| <input type="checkbox"/> 個人情報保護委員会 | <input type="checkbox"/> 要配慮個人情報   | <input type="checkbox"/> マスキング        |

### 答え

#### 個人情報保護法

個人情報の取得・利用・保管・提供のルールを定めた法律。正式名称は個人情報の保護に関する法律。

#### 改正個人情報保護法

改正を重ねている個人情報保護法を指す言い方。3年ごとの見直しが定められており、社会や技術の変化に合わせて更新される。

#### 個人情報保護委員会

個人情報保護法にもとづき事業者を監督・指導する国の機関。名前に反して民間団体ではない。

#### 個人情報取扱事業者

個人情報をデータベース化して事業に使っている者。法律上の義務を負う側で、規模の下限はなく1件でも該当する。

#### 個人識別符号

それ単体で特定の個人を識別できる符号。身体の特徴をデータ化したもの（指紋・顔・虹彩など）と、公的な番号（マイナンバー・旅券番号・運転免許証番号など）。

#### 要配慮個人情報

本人が不当な差別や偏見を受けないよう特に配慮を要する個人情報。人種・信条・社会的身分・病歴・犯罪の経歴・犯罪被害の事実など。原則、本人の同意なしに取得できない。

### 機微（センシティブ）情報

取り扱いに配慮が要する情報の一般的な呼び方。要配慮個人情報とほぼ同義だが、法律用語ではなくより広い言い方。

### 匿名加工情報

特定の個人を識別できないように加工し、元に戻せないようにした情報。条件を満たせば本人の同意なしに第三者提供ができる。

### マスキング

情報の一部を隠す処理。氏名の一部を伏せる、番号の一部だけ残す、顔にぼかしを入れるなど。

## 4-4 制作物に関わる権利（13語）

まず、この語を自分で説明できるか試してください。

- |                                |                                  |                                  |
|--------------------------------|----------------------------------|----------------------------------|
| <input type="checkbox"/> 知的財産権 | <input type="checkbox"/> 意匠権     | <input type="checkbox"/> 営業秘密    |
| <input type="checkbox"/> 著作権   | <input type="checkbox"/> 肖像権     | <input type="checkbox"/> 限定提供データ |
| <input type="checkbox"/> 特許権   | <input type="checkbox"/> パブリシティ権 | <input type="checkbox"/> 著作権侵害   |
| <input type="checkbox"/> 商標権   | <input type="checkbox"/> 不正競争防止法 | <input type="checkbox"/> 名誉毀損    |
| <input type="checkbox"/> 生成物   |                                  |                                  |

### 答え

#### 知的財産権

人が頭で作り出したものを守る権利の総称。著作権・特許権・商標権・意匠権などを含む。

#### 著作権

表現したもの（文章・絵・音楽・プログラムなど）を守る権利。作った瞬間に発生し、登録は要らない。

#### 特許権

発明（技術的なアイデア）を守る権利。出願・審査・登録を経て発生する。

#### 商標権

商品やサービスの名前・マークを守る権利。ロゴやブランド名が対象。登録が必要。

#### 意匠権

物品の見た目のデザイン（形・模様・色）を守る権利。中身ではなく姿かたちが対象。登録が必要。

## 肖像権

自分の顔や姿を勝手に撮られたり公開されたりしない権利。誰にでもある人格の権利。

## パブリシティ権

有名人の名前や顔がもつ顧客を引きつける力（経済的価値）を守る権利。

## 不正競争防止法

事業者どうしの不正な競争を防ぐ法律。他社商品の模倣、営業秘密の不正取得などを規制する。

## 営業秘密

不正競争防止法で守られる情報。①秘密として管理②事業に有用③公然と知られていない、の3つを全部満たすもの。

## 限定提供データ

特定の相手に限って提供しているデータ。営業秘密と公開データの中間を守る区分。

## 著作権侵害

他人の著作物を許諾なく利用すること。AI生成物が既存作品と類似した場合も対象で、責任は公開した人が負う。

## 名誉毀損

人の社会的評価を下げる行為。事実かどうかにかかわらず成立しうる。

## 生成物

生成AIが作り出したもの。文章・画像・音楽・プログラムなど。

## 4-5 AIを取り巻く理念と原則・ガイドライン（21語）

まず、この語を自分で説明できるか試してください。

- |                                                                           |                                                   |                                       |
|---------------------------------------------------------------------------|---------------------------------------------------|---------------------------------------|
| <input type="checkbox"/> 人間の尊厳が尊重される社会 (Dignity)                          | <input type="checkbox"/> 持続可能な社会 (Sustainability) | <input type="checkbox"/> 透明性          |
| <input type="checkbox"/> 多様な背景を持つ人々が多様な幸せを追求できる社会 (Diversity & Inclusion) | <input type="checkbox"/> 人間中心の考え方                 | <input type="checkbox"/> アカウンタビリティ    |
| <input type="checkbox"/> 教育・リテラシー                                         | <input type="checkbox"/> 安全性・公平性                  | <input type="checkbox"/> 人間中心         |
| <input type="checkbox"/> 公正競争確保                                           | <input type="checkbox"/> プライバシー保護                 | <input type="checkbox"/> 安全性          |
|                                                                           | <input type="checkbox"/> セキュリティ確保                 | <input type="checkbox"/> 公平性          |
|                                                                           | <input type="checkbox"/> イノベーション                  | <input type="checkbox"/> AIガバナンス・ゴール  |
|                                                                           | <input type="checkbox"/> 環境・リスク分析                 | <input type="checkbox"/> AIマネジメントシステム |

☐ AI開発者 (AI Developer)

☐ AI提供者 (AI Provider)

☐ AI利用者 (AI Business User)

## 答え

### 人間の尊厳が尊重される社会 (Dignity)

AIに使われる人間ではなく、AIを使う人間でいるという基本理念。

### 多様な背景を持つ人々が多様な幸せを追求できる社会 (Diversity & Inclusion)

いろいろな背景の人が、それぞれの幸せを目指せる社会という基本理念。

### 持続可能な社会 (Sustainability)

続けられる社会という基本理念。環境も格差もここに入る。

### 人間中心の考え方

AI社会原則の1つ。AIは人間を助けるもので、人間が判断の主であること。

### 安全性・公平性

AI社会原則の1つ。安全に動き、偏った差別をしないこと。原則ではこの2つでひとつの項目。

### プライバシー保護

個人の情報を守ること。AI社会原則と共通の指針の両方にある。

### セキュリティ確保

攻撃や事故に耐えること。AI社会原則と共通の指針の両方にある。

### 透明性

どう動いているかを説明できること。AI社会原則と共通の指針の両方にある。

### アカウンタビリティ

説明責任。結果に責任を持つこと。透明性（見える）とは別。

### 人間中心

共通の指針の1つ。原則では「人間中心の考え方」と表記される。

## 安全性

共通の指針の1つ。生命・身体・財産や環境に危害を及ぼさないようにすることです。なお、AI社会原則のほうでは「安全性・公平性」でひとつの項目です。

## 公平性

共通の指針の1つ。学習データのかたよりなどによって、特定の人や集団を不当に差別しないようにすることです。なお、AI社会原則のほうでは「安全性・公平性」でひとつの項目です。

## 教育・リテラシー

共通の指針にだけある項目。AIを使う人を育てること。

## 公正競争確保

共通の指針にだけある項目。特定の企業が独占しないようにすること。

## イノベーション

共通の指針にだけある項目。新しいものが生まれる余地を残すこと。

## 環境・リスク分析

AIガバナンス構築の1段目。自分たちの環境とリスクを調べる。

## AIガバナンス・ゴール

AIガバナンス構築の2段目。どういう状態を目指すのかを決める。

## AIマネジメントシステム

AIガバナンス構築の3段目。回し続ける仕組みを作る。

## AI開発者 (AI Developer)

AIそのものを作る側。モデルを学習させ、システムを開発する。

## AI提供者 (AI Provider)

作られたAIをサービスとして提供する側。

## AI利用者 (AI Business User)

そのAIを事業に使う側。英語は単なるUserではなくBusinessUser。

## 4-6 AI新法（2語）

まず、この語を自分で説明できるか試してください。

☐ AI新法

☐ 人工知能関連技術の研究開発及び  
活用の推進に関する法律

### 答え

#### AI新法

人工知能関連技術の研究開発及び活用の推進に関する法律の通称。2025年6月4日公布。

#### 人工知能関連技術の研究開発及び活用の推進に関する法律

AI新法の正式名称。名前のとおり推進のための法律で、規制のための法律ではない。

## 第5章 テキスト生成AIのプロンプト制作と実例

### 5-1 LMとLLM（10語）

まず、この語を自分で説明できるか試してください。

- |                                     |                                      |                                        |
|-------------------------------------|--------------------------------------|----------------------------------------|
| <input type="checkbox"/> LM         | <input type="checkbox"/> プレトレーニング    | <input type="checkbox"/> プロンプト         |
| <input type="checkbox"/> n-gramモデル  | <input type="checkbox"/> ハイパーパラメータ   | <input type="checkbox"/> プロンプトエンジニアリング |
| <input type="checkbox"/> ニューラル言語モデル | <input type="checkbox"/> Temperature |                                        |
| <input type="checkbox"/> LLM        | <input type="checkbox"/> Top-p       |                                        |

#### 答え

##### LM

言語モデル（LanguageModel）。次に来る言葉を予測する仕組みの総称。

##### n-gramモデル

直前のn個の単語だけを見て次の語を予測する統計的な言語モデル。古いやり方で、離れた言葉のつながりは扱えない。

##### ニューラル言語モデル

ニューラルネットワークを使って次の語を予測する言語モデル。離れた語のつながりも扱える。現在の主流。

##### LLM

大規模言語モデル（LargeLanguageModel）。とても大きな言語モデル。大きいのは学習データの量とパラメータの数。

##### プレトレーニング

事前学習。大量のデータであらかじめ広く学習させ、基礎的な言語の能力を身につけさせる段階。あとから特化させるのがファインチューニング。

##### ハイパーパラメータ

学習の前に人間があらかじめ決めておく設定値。学習でモデル内部が調整するパラメータとは別物。

##### Temperature

出力のばらつきの強さを決めるハイパーパラメータ。低いと安定して無難、高いと多様で意外な出力になる。

## Top-p

次の語の候補をどこまで広げるかを決めるハイパーパラメータ。確率の高い順に足して、合計がpになるまでを候補にする。

## プロンプト

生成AIへ与える指示文。利用者が打ち込む文章そのもの。

## プロンプトエンジニアリング

求める出力が得られるようにプロンプトを設計・工夫すること。モデル自体には手を加えない。

## 5-2 プロンプティングの基礎（6語）

まず、この語を自分で説明できるか試してください。

- |                                      |                                           |                                             |
|--------------------------------------|-------------------------------------------|---------------------------------------------|
| <input type="checkbox"/> Instruction | <input type="checkbox"/> Input Data       | <input type="checkbox"/> Zero-Shot プロンプティング |
| <input type="checkbox"/> Context     | <input type="checkbox"/> Output Indicator | <input type="checkbox"/> Few-Shot プロンプティング  |

### 答え

#### Instruction

プロンプトの4要素のひとつ。指示——何をしてほしいのか。

#### Context

プロンプトの4要素のひとつ。文脈・背景——どういう状況で、誰に向けたものか。

#### Input Data

プロンプトの4要素のひとつ。入力データ——AIに処理してほしい中身そのもの。

#### Output Indicator

プロンプトの4要素のひとつ。出力指示——どんな形式で返してほしいか。

#### Zero-Shot プロンプティング

例を1つも見せずにいきなり指示するやり方。ふだんの使い方はほとんどこれ。

#### Few-Shot プロンプティング

いくつか例を見せてから指示するやり方。出力の形式や語調をそろえたいときに効く。

## 索引 全194語

シラバスの並び順です。節の番号を添えてあるので、動画の該当箇所も探せます。

### 第1章 AI（人工知能）（39語）

#### AI（人工知能） p.2

人間の知能を人工的に実現しようとする技術・研究分野。専門家の間で合意された定義は存在しない。

#### ダートマス会議 p.2

1956年にアメリカのダートマス大学で開かれた研究会です。ジョン・マッカーシーが、ここで「人工知能」という言葉を初めて使いました。

#### ルールベース p.2

人間があらかじめルール（条件と処理）を記述し、そのとおりに動かす方式。判断の理由を説明できる一方、書けるルールの量に限界がある。

#### 機械学習 p.2

人間がルールを書かず、大量のデータからコンピュータ自身に規則性を見つけさせる方式。

#### 学習済みモデル p.2

機械学習によって学習を終えた成果物。新しいデータを入力すると答えを返す。

#### 教師あり学習 p.3

正解つきのデータで学習させるやり方です。「この写真は猫」と答えを添えて大量に見せます。

#### 教師なし学習 p.3

正解を与えずに、データの構造を勝手に見つけさせるやり方です。クラスタリングと次元削減が代表です。

#### クラスタリング p.3

似たものどうしをまとまりに分ける処理です。教師なし学習のひとつ。分けたグループが何を意味するかの解釈は、人の仕事です。

#### 次元削減 p.3

多数の変数を、情報をできるだけ保ったまま、より少ない数の軸にまとめ直す処理です。教師なし学習のひとつ。

#### 強化学習 p.3

正解は教えず、かわりに報酬を与えるやり方です。うまくいったら点、失敗したら減点。試行錯誤しながら点の高い動きを覚えます。

#### 半教師あり学習 p.3

正解つきが少しだけ、正解なしが大量、という状況で両方を使うやり方です。現実のデータはたいていこれです。

#### ノーフリーランチ定理 p.3

あらゆる問題に対して万能に優れたアルゴリズムは存在しないという定理。手法には必ず得手不得手がある。

#### ニューロン p.3

人間の脳にある神経細胞のことです。ニューラルネットワークが手本にした、生き物のほうの仕組みです。

#### シナプス p.3

ニューロンとニューロンのつなぎ目のことです。ここで信号の伝わりやすさが決まります。

#### 人工ニューロン（ノード） p.4

ニューロンを真似て作った、ネットワークの一つひとつの点です。ノードとも呼びます。

#### ニューラルネットワーク p.4

人工ニューロンを層状につないだ仕組みです。人間の脳の神経回路を手本にしています。

#### ディープラーニング p.4

ニューラルネットワークの層を深く積み重ねた機械学習の手法。層を追うごとに単純な特徴から複雑な特徴へと組み上げていく。

#### 重み p.4

つながりの強さを表す数値です。学習とは、この重みを調整していくことです。

#### 情報の重みづけ p.4

どの入力をどれだけ重視するかを、重みで決めることです。学習が進むほど、大事な入力の重みが大きくなります。

#### 過学習 p.4

学習データに合わせすぎて、未知のデータに対する性能が落ちてしまう状態。

## 正則化 p.4

学習のときに罰則を加えて、モデルが複雑になりすぎるのを抑える方法です。過学習を避ける手法のひとつ。

## ドロップアウト p.4

学習のたびにノードをわざと一部お休みさせる方法です。特定のノードに頼りきりにならず、丸暗記しにくくなります。

## 転移学習 p.4

ある目的で学習済みのモデルを別の目的に再利用する手法。少ないデータ・短い時間で学習できる。

## 特徴量 p.5

データのうち「どこに注目すれば見分けられるか」という着目点。従来は人間が指定していたが、ディープラーニングではAIが自分で見つける。

## 弱いAI (ANI) p.5

特定の課題に特化したAIです。現在実用化されているAIは、すべて弱いAIにあたります。

## 強いAI (AGI) p.5

汎用的に多様な課題へ対応し、人間のような意識や自我を持つとされるAI。まだ実現していない。

## 第一次AIブーム p.5

1950年代から60年代のブームです。探索と推論が中心で、迷路やパズルは解けても現実の問題は解けませんでした。

## 探索 p.5

考えられる手を順に調べて、答えにたどり着く道を探すやり方です。第一次AIブームの中心でした。

## 推論 p.5

既にある知識と規則から、新しい結論を導くことです。探索とならぶ第一次AIブームの柱です。

## 第二次AIブーム p.5

1980年代のブームです。専門家の知識をルールにしたエキスパートシステムが中心でしたが、知識を書き切れず終わりました。

## エキスパートシステム p.6

専門家（エキスパート）の知識をルールとして大量に登録し、専門家のように判断させるシステム。第二次AIブームの主役。

## AIの冬 p.6

AIへの期待がしばみ、研究費や関心が落ち込んだ時期。1970年代と1990年代の2回あった。

## 第三次AIブーム p.6

2010年代からのブームです。扱えるデータの量が増え、計算する力が高まり、ディープラーニングが実用的な成果を出しました。

## ビッグデータ p.6

インターネットの普及などにより蓄積された大量かつ多様なデータ。第三次AIブームを支えた要因の1つ。

## シンギュラリティ（技術的特異点） p.6

AIが人間の知能を超え、その先の変化を人間が予測できなくなるとされる時点のことです。到来するかどうかは、研究者の間でも見方が分かれています。

## ヴァーナー・ヴィンジ p.6

SF作家であり数学者。技術的特異点という概念を広めた人物です。年を示したのはカーツワイルのほうです。

## レイ・カーツワイル p.6

技術の進歩の速さから、2045年にAIが人間の知能を超えると予測した人物です。これが2045年問題です。

## 2045年問題 p.6

レイ・カーツワイルが2045年にシンギュラリティが到来すると予測したことに由来する呼び名。

## AI効果 p.7

AIが何かを実現した途端、人がそれを「知能ではない、ただの処理だ」と見なすようになる現象。

## 第2章 生成AI (62語)

### 生成AI (ジェネレーティブAI) p.8

文章・画像・音声・動画など、無かったものを作り出すAI。従来のAIは識別、生成AIは生成。

### ボルツマンマシン p.8

1985年提案。データのでき方を確率で覚えようとした生成モデルの祖先。全部つながっていて計算量が大きく実用にならなかった。

### 制約付きボルツマンマシン p.8

同じ層の中のつながりを切り、層と層の間だけをつないだもの。計算が現実的になり、深い層の学習を支えた。

### 自己回帰モデル p.8

自分が出した出力を次の入力に使い、1つずつ順番に生成していくやり方。GPTがこれ。

### CNN (畳み込みニューラルネットワーク) p.8

畳み込みを使って画像の特徴を取り出すネットワーク。画像を得意とする。

### 畳み込み p.9

小さな窓をずらしながら当て、局所的な特徴を取り出す操作。

### VAE (変分自己符号化器) p.9

入力を縮めて、揺らして、戻すことで新しいデータを作る生成モデル。

### ノイズ p.9

潜在ベクトルにわざと加える揺らぎ。これがあるので元と違うものが出てくる。

### エンコーダ p.9

入力を圧縮する側 (入口)。

### デコーダ p.9

圧縮されたものから復元する側 (出口)。

### 潜在ベクトル p.9

エンコーダが作る短い数字の並び。特徴を圧縮した表現。

### GAN (敵対的生成ネットワーク) p.9

生成器と識別器を競い合わせて質を上げる生成モデル。

### 生成器 p.9

GANの中で偽物を作る側。

### 識別器 p.9

GANの中で本物か偽物かを見分ける側。

### RNN (回帰型ニューラルネットワーク) p.9

隠れ層の出力を次の入力に混ぜ、順番のあるデータを扱えるようにしたモデル。

### 隠れ層 p.9

入力層と出力層のあいだの層。RNNでは記憶係の役目をする。

### リカレント層 p.10

前の状態を次へ戻す構造を持った層。

### シーケンスデータ p.10

文章・音声・株価など、順番に意味があるデータ。系列データ。

### LSTM (長・短期記憶) p.10

RNNに忘れるかどうかを決めるスイッチを足したもの。長い系列でも忘れにくい。

### Transformerモデル p.10

2017年登場。順番に読むのをやめ、文の全部の語を一度に見るモデル。いまの生成AIのほとんどがこの子孫。

## Attention層 p.10

どの語に注目するかを重みとして計算する層。

## 自己注意力 (Self-Attention) p.10

同じ文の中の語どうして注目し合う仕組み。

## Attention Mechanism p.10

Attentionの仕組みそのものを指す言い方。上の3つは同じものを指す。

## 位置エンコーディング p.10

それぞれの語に「何番目か」を数値として足す仕組み。Transformerが語順を扱うために必要。

## アーキテクチャ p.10

モデルの構造・設計そのもの。

## GPTモデル p.10

Transformerの作る側を取り出したモデル。自己回帰で前から後ろへ一方向に書く。

## Open AI p.10

GPTモデルを開発した組織。

## BERTモデル p.11

Transformerの読む側を取り出したモデル。Googleが開発。前も後ろも同時に見る。

## MLM (Masked Language Model) p.11

文の中の単語を隠して当てさせる学習方法。穴埋め。

## NSP (Next Sentence Prediction) p.11

2つの文が続いているかを当てさせる学習方法。続き当てる。

## RoBERTa p.11

BERTをもっと鍛えた版。データを増やし長く学習させ、NSPをやめた。

## ALBERT (a Lite BERT) p.11

BERTの軽い版。パラメータを大幅に減らし、性能を保ったまま小さくした。

## ChatGPT p.11

OpenAIが2022年11月に公開した対話型の生成AI。話し言葉で指示できるようにしたことで、専門知識がなくても使えるようになった。

## GPT-1 p.12

2018年。先に大量の文章を読ませ、あとで目的ごとに追加学習させるという2段階の型を作った最初のGPT。

## 自然言語処理 (NLP) p.12

私たちが普段しゃべっている言葉を、コンピュータに扱わせる技術分野。プログラミング言語のような形式的な言葉の反対側。

## GPT-2 p.12

2019年。パラメータ約15億。大きくすると賢くなることが見えはじめた。悪用のおそれから段階的に公開された。

## パラメータ p.12

ニューラルネットワークの中で調整される数字の総数。第1章の「重み」が全部で何個あるか、という数え方。

## GPT-3 p.12

2020年。パラメータ約1750億。例をいくつか見せるだけで、教えていない仕事もこなせるようになった。

## InstructGPT p.12

2022年。人間の好みを学ばせて、指示に答えるようGPT-3を調整したモデル。ChatGPTの直接の下地。

## GPT-3.5 p.12

InstructGPTの流れをくむGPT-3の改良版。これを対話用にして公開したのがChatGPT。

## GPT-4 p.12

2023年。ハルシネーションが減り、マルチモーダルになった（文章だけでなく画像も扱える）。

## データセット p.12

学習に使うデータのまとまり。GPT-3ではインターネット上の文章・本・Wikipediaなどが大量に使われた。

## RLHF (Reinforcement Learning from Human Feedback) p.13

人間のフィードバックによる強化学習。人が良し悪しを選び、その選び方を報酬にしてモデルを調整する。

## アライメント (Alignment) p.13

AIの振る舞いを、人間の意図や価値観に沿わせることです。RLHFは、そのための手段のひとつです。

## ファインチューニング p.13

学習済みモデルに追加学習をさせて、特定の用途に合わせる。第1章の転移学習の仲間。

## ハルシネーション (Hallucination) p.13

事実に基づかない、もっともらしい情報を生成してしまう現象。次に来そうな言葉を選んでいただけで事実を確かめていないため起きる。

## マルチモーダル p.13

文章・画像・音声など、複数の種類の情報を扱えること。モデルは様式、マルチは複数。

## Code Interpreter p.13

ChatGPTがプログラムを書いて、その場で実行する機能。集計やグラフ作成まで行う。

## GPTs p.13

目的別のChatGPTを自分で作れる仕組み。プログラムを書かずに用途を絞った相手を用意できる。

## GPT-4o p.13

2024年。oはOmni（すべて）。文章・音声・画像を1つのモデルでまとめて扱い、応答が速くなった。ゼロではない。

## Claude p.14

Anthropicの生成AI。長い文章を読ませるのが得意で、安全性と受け答えの丁寧さが評価されている。

## GPT-o1 p.13

答える前に内側で手順を組み立ててから答える推論モデルの第一世代。数学やプログラムに強い。

## GPT-o3 p.13

o1の系譜を進めた推論モデル。考える力を上げたぶん、返答は遅い。

## GPT-o4 p.13

推論モデルの系譜の後の世代。位置づけはo1・o3と同じ「考えてから答える」側。

## GPT-4.1 p.14

2025年。開発者向けの色が濃い。長い文章を読める、指示に正確に従う、といった実務寄りの改良。

## GPT-5 p.14

2025年。速く答える・必要なら考える・文章も画像も扱う、という別々に育った性質が1つにまとまった。

## Sora p.14

OpenAIの動画生成AI。文章で指示すると映像が出てくる。

## Operator p.14

AIがブラウザを自分で操作するもの。検索・ページ操作・入力を代わりに行う。第3章のAIエージェントにつながる。

## Codex p.14

OpenAIのプログラムを書く側の道具。コードの生成や補完に使う。

## Image Generation p.14

OpenAIの画像生成の機能。文章で指示して絵を作る。

## Gemini p.14

Googleの生成AI。最初から文章・画像・音声・動画をまとめて扱う設計。

## Copilot p.14

Microsoftの生成AI。Word・Excel・Teamsの中に組み込まれ、仕事の場所でそのまま使える。

## 第3章 現在の生成AIの動向（16語）

### 画像のリサイズ p.15

画像の大きさ（縦横のピクセル数）を揃えること。AIは決まった大きさの入力しか受け取れないため、学習の前に揃える。

### 正規化 p.15

データの値の範囲を揃えること。画素の0～255を0～1に直すのが代表例。桁が大きいままだと学習が不安定になるため。

### データの水増し（augmentation） p.15

手持ちのデータを加工して枚数を増やすこと。反転・回転・明るさ変更・切り取りなど。中身は変えず、見え方だけ変える。

### データ拡張技術 p.15

データの水増しをするための手口をまとめた呼び名です。水増しが「やること」、データ拡張技術が「その技術の呼び名」です。

### リマスタリング p.15

古い、あるいは粗いデータを作り直して質を上げることです。解像度を上げる、ノイズを取る、色を戻すなど。水増しが「数」を増やすのに対し、リマスタリングは「質」を上げます。

### Veo3 p.15

Googleの動画生成AI。ヴィオ・スリー。SoraのGoogle版と考えるとよい。第3章で新しく出る語。

### ディープフェイク（深層偽造）技術 p.16

ディープラーニングを使って本物そっくりの偽物を作る技術。顔の貼り替えや、実在する人の声で喋らせることができる。

### 偽情報（ディスインフォメーション） p.16

意図的に流される嘘の情報。うっかり広まった誤情報とは、悪意があるかどうかで分かれる。

### RAG（Retrieval-Augmented Generation） p.16

検索拡張生成。答える前に資料を検索し、読ませてから答えさせる仕組み。ハルシネーション対策。モデルは触らない。

### チャンク p.16

資料をあらかじめ細かく切っておいた、その一つひとつ。かたまりの意味。

### ベクトルデータベース p.16

チャンクを意味の数字の並びに変換してしまっておく保管庫。数字なので意味の近さを計算で比べられる。

### AIエージェント p.16

目的を渡すと、手順を自分で考えて、道具を使って、最後までやるAI。第2章のOperatorの正体。

### GenSpark p.17

ジェンスパーク。AIエージェントを使った検索エンジン。自分でいくつも検索して、まとめた答えを作る。

### Manus p.17

マナス。汎用のAIエージェント。調べ物から資料作りまで通しでやる。

### Skywork AI p.17

スカイワーク・エーアイ。汎用型のAIエージェント。資料やスライドを作るのが得意とされる。

### MCP p.17

ModelContextProtocol。Anthropicが出した、AIが外の道具やデータにつながるための共通の決まりごと。

## 第4章 情報リテラシーとAI社会原則 (61語)

### インターネットリテラシー p.18

インターネットを適切に利用するために必要な力。テクノロジーの理解・情報リテラシー・セキュリティとプライバシー・デジタル市民権の4つを含む親の言葉。

### テクノロジーの理解 p.18

インターネットや情報技術の仕組みを知っていること。知らないものは疑えないため、身を守る土台になる。

### 情報リテラシー p.18

情報を見極める力。その情報が本当か、誰が発信しているのかを確かめられること。

### セキュリティとプライバシー p.18

身を守る力。攻撃の手口を知り、設定によって自分の情報を守れること。

### デジタル市民権 p.18

ネット上でもひとりの市民として振る舞うという考え方。権利と責任がともにあるとし、匿名を理由にした無責任な振る舞いを否定する。

### フィッシング詐欺 p.19

偽のメールやWebサイトで本物を装い、ID・パスワード・カード番号などを入力させて盗む詐欺。

### スミッシング p.19

SMS（ショートメッセージ）を使ったフィッシング。SMS + Phishing。宅配便の不在通知を装う手口が代表例。

### ヴィッシング p.19

音声通話を使ったフィッシング。Voice + Phishing。電話で本人に喋らせて聞き出す。

### マルウェア p.19

悪意のあるソフトウェアの総称。malicious + software。ウイルス・ワーム・トロイの木馬・ランサムウェアなどをすべて含む。

### アンチウイルスソフトウェア p.19

マルウェアを検出・遮断・駆除するソフトウェア。名前に反して、ウイルスだけでなくマルウェア全般が対象。

### ランサムウェア p.19

ファイルを暗号化して使えなくし、復旧と引き換えに金銭を要求するマルウェア。ransom = 身代金。払っても戻る保証はない。

### ソーシャルエンジニアリング攻撃 p.19

技術的な解析ではなく、人間の心理や行動につけこんで情報を盗む攻撃の総称。スピアフィッシング・ベイト攻撃・ブラックメール・プレテキストを含む。

### スピアフィッシング p.19

特定の個人や組織を狙い撃ちにするフィッシング。spear = 槍。相手の所属や事情を調べたうえで送るため信じやすい。

### ベイト攻撃 p.19

えさで相手を釣る攻撃。bait = えさ。無料・限定といった欲を刺激する誘いで、リンクを踏ませたり機器を使わせたりする。

### ブラックメール p.20

脅迫。弱みを握ったと告げて金銭や情報を要求する。黒いメールという意味ではない。

### プレテキスト p.20

作り話（口実）で身分をいつわり、信用させてから情報を聞き出す手口。上司・取引先・システム管理者などになりきる。

### 個人情報保護法 p.20

個人情報の取得・利用・保管・提供のルールを定めた法律。正式名称は個人情報の保護に関する法律。

### 改正個人情報保護法 p.20

改正を重ねている個人情報保護法を指す言い方。3年ごとの見直しが定められており、社会や技術の変化に合わせて更新される。

### 個人情報保護委員会 p.20

個人情報保護法にもとづき事業者を監督・指導する国の機関。名前に反して民間団体ではない。

### 個人情報取扱事業者 p.20

個人情報をデータベース化して事業に使っている者。法律上の義務を負う側で、規模の下限はなく1件でも該当する。

## 個人識別符号 p.20

それ単体で特定の個人を識別できる符号。身体の特徴をデータ化したもの（指紋・顔・虹彩など）と、公的な番号（マイナンバー・旅券番号・運転免許証番号など）。

## 要配慮個人情報 p.20

本人が不当な差別や偏見を受けないよう特に配慮を要する個人情報。人種・信条・社会的身分・病歴・犯罪の経歴・犯罪被害の事実など。原則、本人の同意なしに取得できない。

## 機微（センシティブ）情報 p.21

取り扱いに配慮が要する情報の一般的な呼び方。要配慮個人情報とほぼ同義だが、法律用語ではなくより広い言い方。

## 匿名加工情報 p.21

特定の個人を識別できないように加工し、元に戻せないようにした情報。条件を満たせば本人の同意なしに第三者提供ができる。

## マスキング p.21

情報の一部を隠す処理。氏名の一部を伏せる、番号の一部だけ残す、顔にぼかしを入れるなど。

## 知的財産権 p.21

人が頭で作り出したものを守る権利の総称。著作権・特許権・商標権・意匠権などを含む。

## 著作権 p.21

表現したもの（文章・絵・音楽・プログラムなど）を守る権利。作った瞬間に発生し、登録は要らない。

## 特許権 p.21

発明（技術的なアイデア）を守る権利。出願・審査・登録を経て発生する。

## 商標権 p.21

商品やサービスの名前・マークを守る権利。ロゴやブランド名が対象。登録が必要。

## 意匠権 p.21

物品の見た目のデザイン（形・模様・色）を守る権利。中身ではなく姿かたちが対象。登録が必要。

## 肖像権 p.22

自分の顔や姿を勝手に撮られたり公開されたりしない権利。誰にでもある人格の権利。

## パブリシティ権 p.22

有名人の名前や顔がもつ顧客を引きつける力（経済的価値）を守る権利。

## 不正競争防止法 p.22

事業者どうしの不正な競争を防ぐ法律。他社商品の模倣、営業秘密の不正取得などを規制する。

## 営業秘密 p.22

不正競争防止法で守られる情報。①秘密として管理②事業に有用③公然と知られていない、の3つを全部満たすものの。

## 限定提供データ p.22

特定の相手に限って提供しているデータ。営業秘密と公開データの中間を守る区分。

## 著作権侵害 p.22

他人の著作物を許諾なく利用すること。AI生成物が既存作品と類似した場合も対象で、責任は公開した人が負う。

## 名誉毀損 p.22

人の社会的評価を下げる行為。事実かどうかにかかわらず成立しうる。

## 生成物 p.22

生成AIが作り出したもの。文章・画像・音楽・プログラムなど。

## 人間の尊厳が尊重される社会（Dignity） p.23

AIに使われる人間ではなく、AIを使う人間でいるという基本理念。

## 多様な背景を持つ人々が多様な幸せを追求できる社会（Diversity & Inclusion） p.23

いろいろな背景の人が、それぞれの幸せを目指せる社会という基本理念。

## 持続可能な社会（Sustainability） p.23

続けられる社会という基本理念。環境も格差もここに入る。

## 人間中心の考え方 p.23

AI社会原則の1つ。AIは人間を助けるもので、人間が判断の主であること。

## 安全性・公平性 p.23

AI社会原則の1つ。安全に動き、偏った差別をしないこと。原則ではこの2つでひとつの項目。

## プライバシー保護 p.23

個人の情報を守ること。AI社会原則と共通の指針の両方にある。

## セキュリティ確保 p.23

攻撃や事故に耐えること。AI社会原則と共通の指針の両方にある。

## 透明性 p.23

どう動いているかを説明できること。AI社会原則と共通の指針の両方にある。

## アカウントビリティ p.23

説明責任。結果に責任を持つこと。透明性（見える）とは別。

## 人間中心 p.23

共通の指針の1つ。原則では「人間中心の考え方」と表記される。

## 安全性 p.24

共通の指針の1つ。生命・身体・財産や環境に危害を及ぼさないようにすることです。なお、AI社会原則のほうでは「安全性・公平性」でひとつの項目です。

## 人工知能関連技術の研究開発及び活用の推進に関する法律 p.25

AI新法の正式名称。名前のとおり推進のための法律で、規制のための法律ではない。

## 公平性 p.24

共通の指針の1つ。学習データのかたよりなどによって、特定の人や集団を不当に差別しないようにすることです。なお、AI社会原則のほうでは「安全性・公平性」でひとつの項目です。

## 教育・リテラシー p.24

共通の指針にだけある項目。AIを使う人を育てること。

## 公正競争確保 p.24

共通の指針にだけある項目。特定の企業が独占しないようにすること。

## イノベーション p.24

共通の指針にだけある項目。新しいものが生まれる余地を残すこと。

## 環境・リスク分析 p.24

AIガバナンス構築の1段目。自分たちの環境とリスクを調べる。

## AIガバナンス・ゴール p.24

AIガバナンス構築の2段目。どういう状態を目指すのかを決める。

## AIマネジメントシステム p.24

AIガバナンス構築の3段目。回し続ける仕組みを作る。

## AI開発者（AI Developer） p.24

AIそのものを作る側。モデルを学習させ、システムを開発する。

## AI提供者（AI Provider） p.24

作られたAIをサービスとして提供する側。

## AI利用者（AI Business User） p.24

そのAIを事業に使う側。英語は単なるUserではなくBusinessUser。

## AI新法 p.25

人工知能関連技術の研究開発及び活用の推進に関する法律の通称。2025年6月4日公布。

## 第5章 テキスト生成AIのプロンプト制作と実例（16語）

### LM p.26

言語モデル（LanguageModel）。次に来る言葉を予測する仕組みの総称。

### n-gramモデル p.26

直前のn個の単語だけを見て次の語を予測する統計的な言語モデル。古いやり方で、離れた言葉のつながりは扱えない。

### ニューラル言語モデル p.26

ニューラルネットワークを使って次の語を予測する言語モデル。離れた語のつながりも扱える。現在の主流。

### LLM p.26

大規模言語モデル（LargeLanguageModel）。とても大きな言語モデル。大きいのは学習データの量とパラメータの数。

### プレトレーニング p.26

事前学習。大量のデータであらかじめ広く学習させ、基礎的な言語の能力を身につけさせる段階。あとから特化させるのがファインチューニング。

### ハイパーパラメータ p.26

学習の前に人間があらかじめ決めておく設定値。学習でモデル内部が調整するパラメータとは別物。

### Temperature p.26

出力のばらつきの強さを決めるハイパーパラメータ。低いと安定して無難、高いと多様で意外な出力になる。

### Top-p p.27

次の語の候補をどこまで広げるかを決めるハイパーパラメータ。確率の高い順に足して、合計がpになるまでを候補にする。

### プロンプト p.27

生成AIへ与える指示文。利用者が打ち込む文章そのもの。

### プロンプトエンジニアリング p.27

求める出力が得られるようにプロンプトを設計・工夫すること。モデル自体には手を加えない。

### Instruction p.27

プロンプトの4要素のひとつ。指示——何をしてほしいのか。

### Context p.27

プロンプトの4要素のひとつ。文脈・背景——どういう状況で、誰に向けたものか。

### Input Data p.27

プロンプトの4要素のひとつ。入力データ——AIに処理してほしい中身そのもの。

### Output Indicator p.27

プロンプトの4要素のひとつ。出力指示——どんな形式で返してほしいか。

### Zero-Shot プロンプティング p.27

例を1つも見せずにいきなり指示するやり方。ふだんの使い方はほとんどこれ。

### Few-Shot プロンプティング p.27

いくつか例を見せてから指示するやり方。出力の形式や語調をそろえたいときに効く。