

生成AI活用普及協会（GUGA）公式シラバス準拠

生成AIパスポート 完全攻略テキスト

全5章の講義＋難問対策＋一問一答を、この1冊に。
2026年2月試験より適用のシラバスに対応。

- 1 読む 第1部で全範囲をひととおり
- 2 確かめる 第3部の一問一答194語で穴を洗い出す
- 3 詰める 第2部の4つの道具で、取りこぼしを減らす

この1冊で、全範囲・全キーワード194語

79

図解

194

キーワード

256

ページ

537

練習問題（Web）



解説動画（YouTube）

全5章の講義＋
難問対策編。無料



練習問題537問・模試

登録不要・広告なし
ai-shikaku-lab.pages.dev



AI資格ラボ YouTubeチャンネル

掲載の問題・解説はすべて自作であり、試験の過去問ではありません。

はじめに この本の使い方

本書は、YouTubeチャンネル「AI資格ラボ」で公開している講義動画の教科書を、**1冊にまとめたもの**です。動画を見ていなくても、これだけで読み切れます。

生成AIパスポート試験の出題範囲は、生成AI活用普及協会（GUGA）が公開している**公式シラバス**で決まっています。本書はその**詳細キーワードをひとつ残らず**収録しました。

この本の3部構成

- **第1部 講義編** 第1章～第5章。図解と確認問題つきで、範囲をひとつとり
- **第2部 解き方編** 2つまで絞れるのに決めきれない、を減らす4つの道具
- **第3部 総仕上げ** 全キーワード194語の一問一答。思い出せるか確かめる

⇒ ★ **おすすめの順番は、第1部 → 第3部 → 第2部です。**まず範囲を通し、一問一答で「思い出せない語」を洗い出し、そのうえで解き方を身につけると、いちばん短く済みます。

⚠️ ⚠️ **この試験に、公開された過去問はありません。**「過去問」を名乗るものはすべて推測です。本書に載せている問題も、すべて公式シラバスをもとにした自作の予想問題です。



解説動画（YouTube 再生リスト）

本書と同じ順番・同じ図で、全5章+難問対策編を無料公開しています。
章の扉ページにも、その章の動画のQRがあります。



練習問題537問・模擬試験（Web）

登録不要・広告なし。本番と同じ60問の模試、間違えた問題だけの復習、
一問一答の聞き流しまで、すべて無料です。 ai-shikaku-lab.pages.dev

目次

この本の中身

第1部	講義編	3
	第1章 AI（人工知能）	4
	第2章① 生成AI（仕組み編）	30
	第2章② 生成AI（ChatGPT編）	58
	第3章 現在の生成AIの動向	81
	第4章① 情報リテラシー（身を守る）	104
	第4章② 権利とAI社会原則	132
	第5章 プロンプト制作と実例	162
第2部	解き方編	188
	難問対策編	189
第3部	総仕上げ	206
	一問一答194語（聞き流し版）	207
巻末	総索引	244

第1部

講義編

全5章をひととおり

AI資格ラボ

第1章 AI（人工知能）



この章の解説動画

<https://youtu.be/C67nQet0JSs>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに この本の使い方

この本は、YouTubeチャンネル「AI資格ラボ」の動画「第1章 AI（人工知能）まるごと解説」に対応した、配布用のテキストです。動画を見ていない方でも、これ一冊で第1章が最後まで読み切れるように書いてあります。

この本の構成

章の中身は、**生成AIパスポートの公式シラバス**（出題範囲表）にそって5つのテーマに分けてあります。順番も公式のとおりです。飛ばさず頭から読んでいけば、出題範囲を1周したことになります。

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま問題になる所です。試験直前はここだけ拾い読みできます。
- **△マークの枠**（赤い枠）……**間違えやすい所・引っ掛けが仕込まれる所**です。
- **用語**（紫の帯）……シラバスに名前が載っている用語です。全部で39語あります。
- **確認問題**（緑の枠）……テーマの区切りに3問ずつ、合計9問あります。
- **チェックリスト**（巻末）……39語を自分で説明できるか、最後に確かめられます。

おすすめの進め方

1. まず**本文を最後まで通して読む**。分からない所で止まらず、飛ばして進んでかまいません。
2. 次に**確認問題9問**を、本文を見ずに解く。ここで落ちた所が、あなたの弱点です。
3. 落ちた所だけ本文に戻り、**自分の言葉で言い直す**。読むより言い直したほうが残ります。
4. 試験直前は、**試験ポイントの枠と、巻末の用語集だけ**を見返す。

→ 印刷するときは**A4・実物大（100%）**で。図の中の文字まで読める大きさに作ってあります。

第1章は、どういう章か

先に言っておきます。この章は、生成AIパスポートの中で**いちばん理不尽**な章です。

これから「AIとは何か」を学ぶのですが、実は**AIに定義はありません**。定義がないものについて、4択で問われる。それがこの試験です。

ただ、理不尽なぶん**やることは決まっています**。定義がないからこそ、「誰が」「いつ」「何と呼んだか」という**事実の暗記**が中心になります。考えて解く問題は多くありません。落ち着いて覚えれば、第1章は確実に取れる章です。

AI資格ラボ

目次

この本の中身

はじめに	この本の使い方	4
第1章	この章の全体像	7
1-1	AIの定義	8
	AIに定義はない／1956年 ダートマス会議／AIとロボットの区別	
1-2	AIに知能をもたらす仕組み	10
	ルールベースと機械学習／機械学習の4つの手法／ノーフリーランチ定理	
	脳とニューラルネットワーク／ディープラーニングと特徴量	
	過学習と、それを避ける手法／転移学習	
	確認問題①（1-1・1-2）	17
1-3	AIの種類	18
	AIの4つのレベル／弱いAI（ANI）と強いAI（AGI）	
1-4	AIの歴史	21
	第一次AIブームと1回目の冬／第二次AIブームと2回目の冬／第三次AIブーム	
	確認問題②（1-3・1-4）	23
1-5	シンギュラリティ	24
	技術的特異点／2045年問題／AI効果	
	確認問題③（総まとめ）	25
まとめ	第1章 これが言えれば合格ライン	26
用語集	第1章の重要用語 39語	28

動画と合わせて使う

- この本は、YouTube「AI資格ラボ」の動画「【生成AIパスポート】第1章 AI（人工知能）まるごと解説」と同じ順番・同じ図で作ってあります。
- 動画で流れをつかみ、この本で確かめる。あるいはこの本を先に読んでから動画を見る。どちらでも構いません。
- 印刷はA4・実物大（100%）で。図の中の文字まで読める大きさに作ってあります。

第1章 この章の全体像

第1章で問われるのは、次の5つです。公式シラバスの並びと同じで、この本の1-1～1-5に対応しています。

	テーマ	ここで問われること
1	AIの定義	AIという言葉の由来。 定義がないこと と、 ダートマス会議・1956年
2	AIに知能をもたらす仕組み	この章の本丸。 ルールベースと機械学習、4つの手法、ニューラルネットワーク、ディープラーニング、過学習、転移学習
3	AIの種類	4つのレベル と、弱いAI (ANI)・強いAI (AGI)
4	AIの歴史	ブーム3回・冬2回 。主役と、終わった理由
5	シンギュラリティ	技術的特異点。 人物2人 と2045年問題、AI効果

この章の読み方

- 2番の「**知能をもたらす仕組み**」が飛びぬけて多い。用語39語のうち21語がここに集中していて、2番さえ抜ければ第1章は終わったようなものです。
- 5つを貫く軸は**人間がどこまで手を出しているか**。人間がルールを全部書く（ルールベース／第一次・第二次ブーム）→ ルールはAIが見つける（機械学習／レベル3）→ **注目点もAIが決める**（ディープラーニング／レベル4）。**歴史もレベル分けも全部この一方向**です。
- 迷ったら「**人間の仕事がどれだけ減ったか**」で並べ直せば順番が出ます。

1-1 AIの定義

最初のテーマは「AIの定義」です。ところが、いきなり出鼻をくじくことを言います。**AIに、定義はありません。**

AI = 人工知能。ただし定義はない

まず言葉から。**AIは、日本語で「人工知能」**です。Artificial Intelligence（アーティフィシャル・インテリジェンス）の訳で、ここは素直です。

問題はここから。**AIには、専門家間で合意された定義が存在しません。**研究者ごとに言っていることが違います。そして、このこと自体が試験に出ます。

AI（人工知能） Artificial Intelligence

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

ここが引っかけ

- 「AIとは〇〇と定義されている」と書いてある選択肢は、**だいたい誤り**です。
- 逆に「AIには明確な定義がない」という趣旨の選択肢は、**正しいことが多い**。

なぜ定義がないのか

理由を聞くと、わりと納得します。**そもそも「知能とは何か」が決まっていないから**です。

知能が何かわかっていないのに、人工の知能だけ先に定義できるわけがない。順番が逆なのです。それでも人類は先に作ってしまい、作ってから「これは何だろうね」と議論している——今がその状態です。

⇒ この「定義がない」という事実は、章の最後に出てくる**AI効果**（1-5）にそのままつながります。伏線だと思って覚えておいてください。

1956年 ダートマス会議

定義はありませんが、**言葉が生まれた瞬間**ははっきりしています。ここは年号が出ます。

1956年、アメリカのダートマス大学で開かれた研究会が**ダートマス会議**です。この会議で、**ジョン・マッカーシー**が「Artificial Intelligence」という言葉を初めて使いました。つまり、**AIという言葉は1956年に生まれたこと**になります。

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。ジョン・マッカーシーがここで「Artificial Intelligence（人工知能）」という言葉の初めて使った。**AIという言葉が生まれた場**とされる。

試験ポイント

- 「ダートマス会議」と「1956年」は、必ずセットで覚える。片方だけ聞かれることも、両方セットで聞かれることもあります。
- 提唱者はジョン・マッカーシー。人名まで問われることがあります。

⇒ 定義はないのに、名前だけ先にある。第1章はこういう章です。

AIとロボットの区別

次に、混同する人が非常に多いところです。AIとロボットは別物です。

AIはソフトウェア、ロボットはハードウェア。AIが「頭」で、ロボットが「体」だと考えてください。

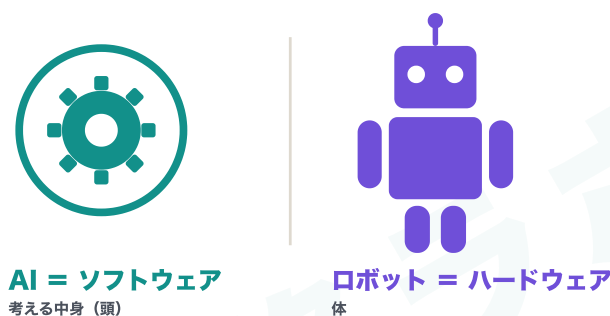


図1-1 AIは頭（ソフトウェア）、ロボットは体（ハードウェア）。両方そろっている必要はない

頭と体が別物だということは、**片方しか持っていないものが存在する**ということです。実際、両方あります。

AIが入っていないロボット

工場でずっと同じ動きを繰り返すアーム。決められた動作を正確に繰り返しているだけで、**考えてはいない**。

体を持たないAI

ChatGPTのような対話AI。あれだけ言葉を扱うのに、**手も足もない**。ソフトウェアだけで成立している。

ここが引っかけ

- 「**ロボットの研究こそAI研究のすべてである**」という選択肢 → **誤り**。体は必須ではありません。
- 「AIを実現するには必ず物理的な身体（ロボット）が必要である」といった言い回しも、同じく誤りです。

1-2 AIに知能をもたらす仕組み

ここが**第1章の本丸**です。シラバスに載っている用語39語のうち21語がこのテーマに集中していて、分量も出題も他とは桁が違います。逆に、ここさえ抜ければ第1章は終わったようなものです。

仕組みは2つしかない

まず全体の見取り図から。**AIに知能をもたらす仕組みは、2つあります。ルールベースと機械学習。**この2つだけです。

ルールベース

人間がルールを書く

機械学習

データからAIが見つける

⇒ 「仕組みは**2つ**」という数字自体が問われます。3つでも4つでもありません。

ルールベース

人間がルールを全部書く方式です。たとえば「もし体温が37度5分以上なら、発熱と表示する」。こうした条件を人間が延々と書き並べて、そのとおりに動かします。

長所

なぜその答えになったかを説明できる。書いてあるルールをたどれば、判断の理由がそのまま分かる。

短所

書ける量に限界がある。世の中のルールを人間が全部書くことはできない。

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

もう一つが機械学習です。**ルールを人間が書かず、データからコンピュータ自身に見つけさせる**方式です。大量のデータを与えると、その中のパターンをAIが自分で取り出してくれます。

そうやってできあがったものを**学習済みモデル**と呼びます。ここは用語として出ます。学習した**あとの**成果物が「学習済みモデル」です。一度作れば、あとはそれに新しいデータを入れて答えを出すだけ——工場という**かながた**金型のようなものだと考えてください。

大量のデータ



学習



学習済みモデル

成果物

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。

試験ポイント

- ルールベース＝人間がルールを書く。機械学習＝データからAIが見つける。
- この対比だけで1問取れます。どちらが「人間の仕事」かで見分けてください。

機械学習の4つの手法

機械学習には手法が**4つ**あります。第1章の中で、**丸暗記が必要な数少ない場所**です。

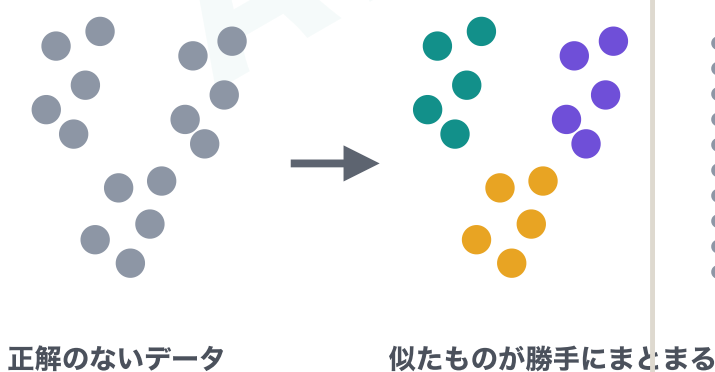
① 教師あり学習

正解つきのデータで学習します。「この写真は猫」「これは犬」と、答えを添えて大量に見せる方式です。迷惑メールの判定や、売上の予測など、**答えを知っている問題**に使います。

② 教師なし学習

正解がないデータを渡して、勝手に構造を見つけさせます。代表的なものが2つあり、**どちらも試験に出ます**。

クラスタリング



次元削減

100項目のアンケート

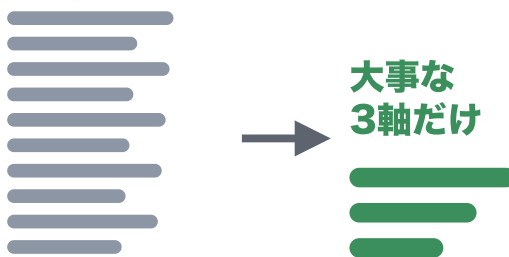


図1-2 教師なし学習の代表2つ。クラスタリングは「似たものの集め」、次元削減は「大事な軸だけ残す」

- **クラスタリング**……似たもののどうしを勝手にグループ分けする。顧客をタイプ別に分ける、など。
- **次元削減**^{じげんさくげん}……項目が多すぎる時、**大事な軸だけ残して減らす**。100項目のアンケートが、実は3つの軸で説明できる、といった話。

ここが引っかけ

- **クラスタリングと次元削減は、どちらも「教師なし学習」**です。ここをセットで問う問題が繰り返し出ます。
- 「クラスタリングは教師あり学習の一種である」という選択肢は **誤り**。

③ 強化学習

正解は教えません。かわりに**報酬**を与えます。うまくいったら点、失敗したら減点。試行錯誤を繰り返しながら、点の高い動きを覚えていく方式です。囲碁のAI、ロボットの歩行、ゲームのAIなどが代表例です。

④ 半教師あり学習

正解つきが少しだけ、正解なしが大量、という状況で使います。現実のデータはたいていこれです。データはあるけれど、全部に人手でラベルを貼るのは費用も時間も現実的でない。だから少しだけ貼って、残りは推測させます。

⇒ **「ラベル付けにお金がかかる」という現実の都合から生まれた手法**、と覚えると忘れません。

手法	正解データ	代表例・特徴
教師あり学習	あり	迷惑メール判定、売上予測
教師なし学習	なし	クラスタリング／次元削減
強化学習	なし（報酬）	囲碁AI、ロボットの歩行
半教師あり学習	少しだけあり	ラベル付けの費用を抑える

試験ポイント

- 4つの名前をそのまま言えるようにする。**教師あり／教師なし／強化／半教師あり**。
- **教師なしの中に、クラスタリングと次元削減**。この入れ子の関係まで含めて1セットです。

ノーフリーランチ定理

名前だけ変わった定理が1つ出ます。**ノーフリーランチ定理**。直訳すると「タダ飯はない」です。

意味は、**あらゆる問題に万能なアルゴリズムは存在しない**ということ。どんな手法にも得意な問題と苦手な問題があり、「この手法さえ使えば全部解ける」は原理的にありえません。だから現場では、問題ごとに手法を選び直しています。

ノーフリーランチ定理 No Free Lunch Theorem

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。

ここが引っかけ

「あらゆる問題に万能なアルゴリズムが存在する」という選択肢が出たら、**この定理で切ります**。考え込む必要はありません。

人間の脳とニューラルネットワーク

機械学習の中には、**人間の脳のしくみを真似た**やり方があります。これがのちのディープラーニングにつながる、いちばん大事な流れです。

人間の脳には**ニューロン**（日本語で神経細胞^{しんけいさいぼう}）があり、ニューロンどうしのつなぎ目を**シナプス**といいます。信号がシナプスを通して次のニューロンへ伝わる——これが脳の基本です。



これを数式で真似たものが**人工ニューロン**、別名**ノード**です。ノードをたくさん並べて層にし、つないだものが全体が**ニューラルネットワーク**です。

ここで最重要の言葉が**重み**。つなぎ目ごとに「この信号をどれだけ重視するか」という数字がついていて、それが重みです。そして**学習とは、この重みの数字を調整していく作業**のことです。「情報の重みづけを変える」と「学習する」は、同じことを言っています。

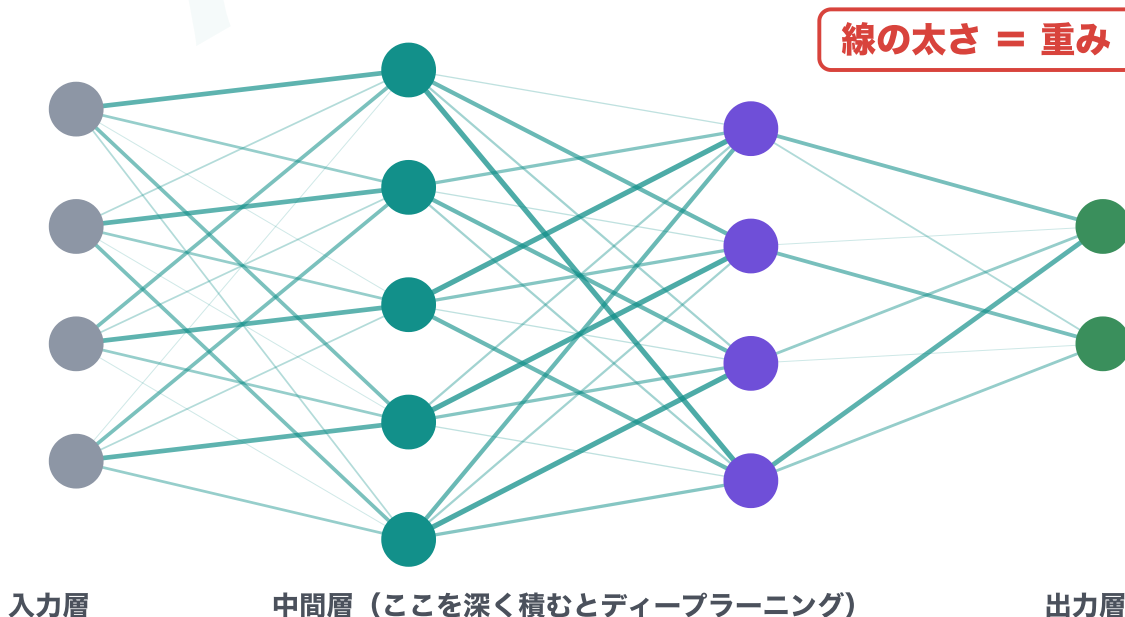


図1-3 ノードを層に並べたものがニューラルネットワーク。線の太さが「重み」で、学習とはこの太さを調整すること

人間の脳	AI側の言葉
ニューロン（神経細胞）	人工ニューロン（ノード）
ニューロンを層に並べた全体	ニューラルネットワーク
シナプス（つなぎ目）	重み

試験ポイント

- ニューロン → 人工ニューロン（ノード）、シナプス → 重み。この対応表がそのまま出ます。
- 学習 = 重み（情報の重みづけ）を調整すること。

ディープラーニングと特徴量

ニューラルネットワークの層を**たくさん深く積んだ**ものが**ディープラーニング**です。「ディープ」は深いという意味で、層が深い、それだけの意味です。名前がそのままなので、むしろ拍子抜けするかもしれません。

ディープラーニング 深層学習

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。

AIはどうやって画像を見分けているのか

猫の写真为例にします。入り口に近い層は、**線や角のような単純な形**を捉えます。次の層では、**目や耳のようなまとまり**になり、さらに奥の層で「**猫らしさ**」にたどりつきます。つまり、**単純なものから複雑なものへ、層を追うごとに組み上がって**いくわけです。

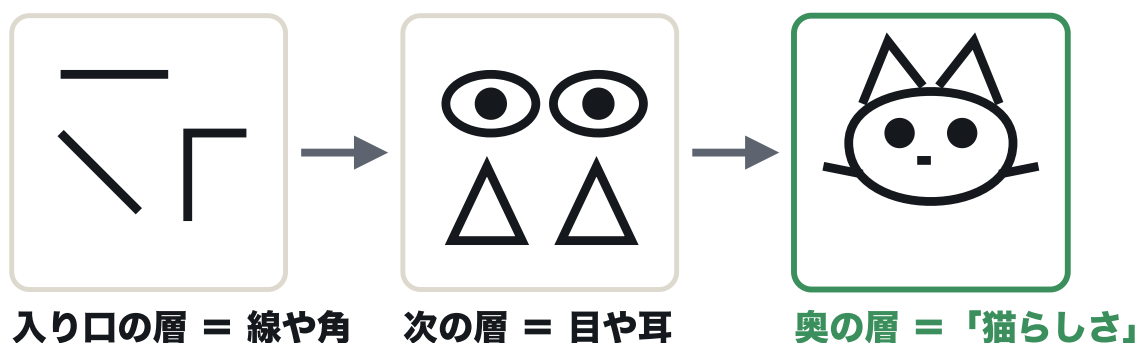


図1-4 層が深くなるほど、捉える特徴が単純なものから複雑なものへ組み上がっていく

そして、ここが**第1章の山場**です。このとき「**どこに注目すべきか**」を、**AIが自分で見つけている**という点です。人間は「耳の形を見ろ」とは教えていません。

この「注目すべきポイント」のことを^{とくちょうりょう}**特徴量**といいます。

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。

第1章で最も出る

- 特徴量を人間が決めるのか、AIが自分で見つけるのか。
- この違いが、次の1-3に出てくる「AIの4つのレベル」のレベル3とレベル4の境目になります。この3文字だけは絶対に持って帰ってください。

過学習と、それを避ける手法

AIが学習しすぎると、困ったことが起きます。^{かがくしゅう}**過学習**（オーバーフィッティング）です。ここも^{ひんしゅつ}頻出します。

どういう状態か。練習した問題は完璧に解けるのに、初めて見る問題がまるで解けないという状態です。過去問の答えを丸暗記して、本番で数字が変わった瞬間に解けなくなる——あの状態を、AIもやります。学習データの細かい癖^{くせ}まで覚え込んでしまい、一般化できなくなるのです。

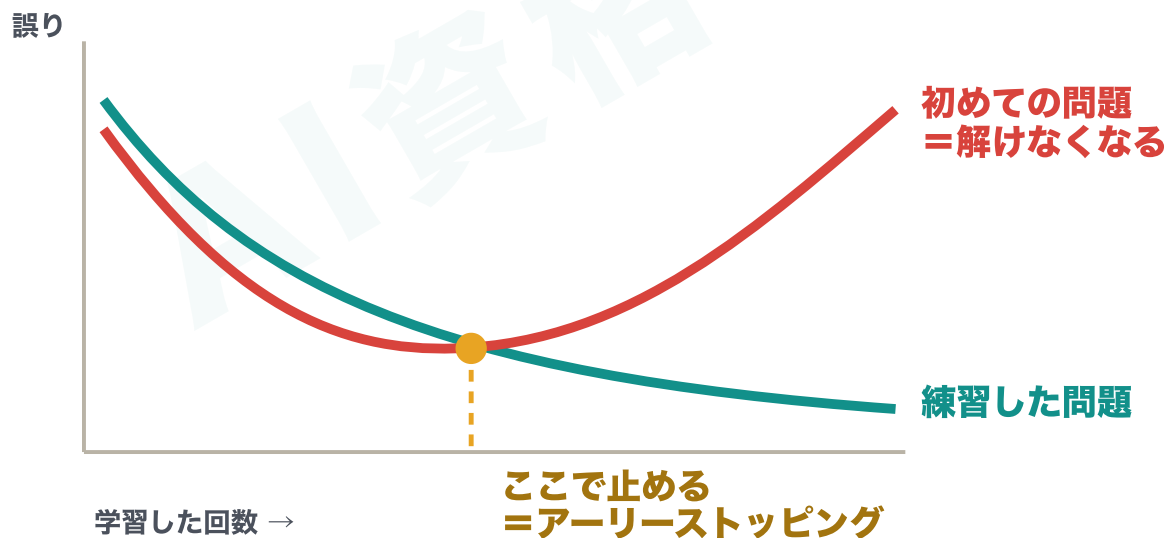


図1-5 学習を続けると、学習データの成績は上がり続けるのに、初めて見るデータの成績は途中から下がり始める

過学習 オーバーフィッティング

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。

これを避ける手法が、名前で問われます。まず**シラバスに名前が載っているのは、次の2つ**です。

① 正則化

学習のときに^{ばっそく}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。「そんなに細かく合わせに行くな」とブレーキをかけるイメージです。

② 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。毎回メンバーを入れ替えて練習させるようなもので、特定のノードに頼りきりにならないため、**丸暗記しにくくなります**。

③ アーリーストッピング（早期打ち切り）

学習を続けると、あるところから本番の成績が下がり始めます（図1-5の赤い線）。その手前で**学習をやめてしまう**方法です。身も蓋もないようですが、よく効きます。

⇒ ③は公式シラバスのキーワードには入っていませんが、過学習の対策としてセットで語られます。**まず①②を確実に、③は余裕があれば**という優先順位で覚えてください。

試験ポイント

- 過学習ときたら、正則化・ドロップアウト。この2語はシラバスのキーワードです。
- 過学習の説明は「**学習データには強いが、未知のデータに弱い**」という言い回しで出ます。

転移学習

最後に、実務でよく使われる手法を1つ。^{てんいがくしゅう}**転移学習**です。

すでに学習済みのモデルを持ってきて、別の目的に作り替える手法です。たとえば大量の画像で鍛えられたモデルがあるとして、それを土台にし、自分の欲しい判定だけを少ないデータで追加学習させます。ゼロから作れば膨大なデータと時間が要るところを、**少ないデータで済ませられる**のが利点です。

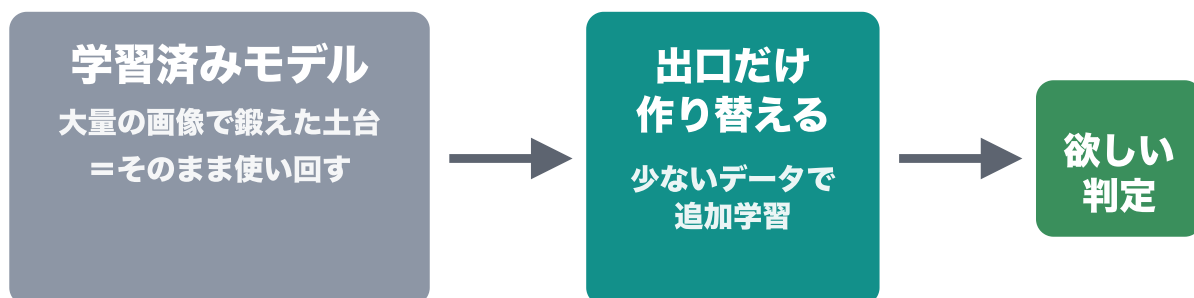


図1-6 学習済みモデルを土台にして、目的の部分だけを少ないデータで学習し直す

⇒ イメージは**中途採用**です。基礎ができている人を連れてきて、自社のやり方だけ覚えてもらう。新卒を一から育てるより速い、という発想です。

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。

試験ポイント

- 試験では「少ないデータで学習できる」「学習済みモデルを再利用する」という言い回しで出ます。

確認① ここまでの確認問題

1-1と1-2の範囲から3問です。まず本文を見ずに答えてください。頭の中で言い切ってから、下の答えを見ると定着します。

確認問題①

第1問 AIという言葉が生まれた会議の名前と、その年は？

答

ダートマス会議、1956年。

提唱したのはジョン・マッカーシー。「会議名・年・人名」の3点セットで押さえます。

確認問題①

第2問 AIに知能をもたらす2つの仕組みは？ またその違いは？

答

ルールベースと機械学習。

ルールベースは人間がルールを書く。機械学習はデータからAIが見つかる。

確認問題①

第3問 クラスタリングと次元削減は、機械学習の4つの手法のうちどれにあたるか？

答

どちらも教師なし学習。

「クラスタリングは教師あり」と思い込む人が多い、いちばんの落とし穴です。

1-3 AIの種類

ここからテーマが変わります。AIには定義がないので、かわりにレベル分けて整理しようということになりました。苦肉の策ですが、第1章でいちばん出るのがここです。

AIの4つのレベル

先に結論だけ言います。レベルが上がるほど、人間が教える量が減ります。この軸さえ持っていれば、4つを丸暗記しなくても選択肢を切れます。

人間が教える量が減る →



図1-7 AIの4つのレベル。上に行くほど、人間の仕事が減っていく

レベル1 | 単純な制御プログラム

決められたルールどおりに動くだけで、学習は一切していません。例はエアコンです。設定温度より暑ければ冷やし、寒ければ止める。それだけです。

⇒ 家電量販店で見かける「AI搭載」のシールは、多くがこのレベル1です。AIに定義がないので、そう名乗っても嘘にはならない——1-1の「定義がない」が、こんなところで効いてきます。

レベル2 | 古典的な人工知能

たんさく すいろん
探索と推論を使い、状況に応じて振る舞いを変えられます。代表例はお掃除ロボット（壁にぶつかったら向きを変える）や、症状を入れると候補を出す診断プログラムです。レベル1との違いは、扱える組み合わせの数です。

ここが引っかけ

レベル2でも、ルール自体は人間が書いています。人間が書いた分厚いルールブックを高速にめくっているだけで、AIが学習しているわけではありません。

レベル3 | 機械学習を取り入れたAI

ここでついに「学習」が出てきます。大量のデータから、**ルールをAI自身が見つける**段階です。例は検索エンジン。どのページが役に立つかを、人間が1件ずつ決めているわけではありません。

ただし、レベル3には人間の仕事が1つ残っています。「どこに注目すべきか」＝**特徴量は、人間が指定する**のです。

レベル4 | ディープラーニングを取り入れたAI

その特徴量まで、AIが自分で見つけるようになった段階です。「耳の形を見ろ」と教える必要すらありません。**今の生成AIは、すべてこのレベル4の延長線上**にあります。

	中身	例
1	単純な制御。ルールは人間が全部決める。学習はしない	エアコン
2	探索・推論。ルールは 人間が書く	お掃除ロボット
3	機械学習。ルールはAIが学ぶ。 特徴量は人間が指定	検索エンジン
4	特徴量もAIが自分で見つける	画像認識・生成AI

第1章で最も出る

- 問われるのは、ほぼ**レベル3とレベル4の違い**です。
- 答えは**特徴量**。**人間が決めるのがレベル3、AIが自分で見つけるのがレベル4**。この3文字だけは絶対に持って帰ってください。

弱いAI (ANI) と強いAI (AGI)

同じ「AIの種類」の中に、もう一つの分け方があります。**弱いAIと強いAI**です。哲学者の**ジョン・サール**が言い出した区別です。

強いAI = AGI

Artificial **General** Intelligence の略で、General は汎用^{はんよう}という意味です。**いろいろな課題に対応でき、人間と同じような意識や自我を持つAI**を指します。

弱いAI = ANI

Artificial **Narrow** Intelligence の略で、Narrow は**狭い**という意味です。**1つの課題に特化した、道具としてのAI**を指します。

弱いAI (ANI) Artificial Narrow Intelligence

特定の課題に**特化**したAI。現在実用化されているAIは**すべてこちら**。

強いAI (AGI) Artificial General Intelligence

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**。

⇒ 「汎用AI」「特化型AI」という言い方も同じ意味で使われますが、**公式シラバスに載っているのはANIとAGI**です。答案はこちらの表記で覚えてください。



図1-8 ChatGPTも画像認識も将棋AIも、分類上はすべて「弱いAI (ANI)」

ここで大事な事実です。**いま実用化されているAIは、すべて弱いAI (ANI)**です。ChatGPTも、画像認識も、将棋のAIも、全部そうです。あれだけのことができて、分類上は「弱い」のです。

ここが引っかけ

- 「**強いAI (AGI) はすでに実現している**」という選択肢 → **誤り**。考えずに切って構いません。
- 「**実用化されているAIはすべて弱いAI (ANI) である**」 → **正しい**。

1-4 AIの歴史

ここもテーマが変わります。押さえることは1つだけ——**ブーム3回、冬2回**。その3回それぞれについて、**何ができるようになって、何ができなくて終わったか**をセットで持つと、一気に頭に入ります。

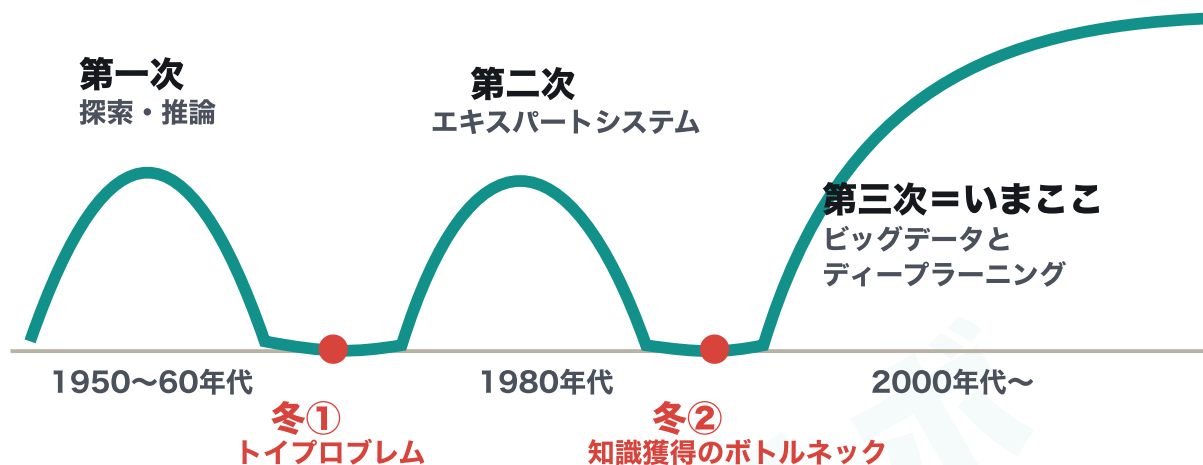


図1-9 ブームは3回、冬は2回。それぞれ主役が入れ替わっている

第一次AIブーム（1950年代後半～1960年代）

ダートマス会議（1956年）の直後にあたります。主役は**探索**と**推論**でした。

探索

可能な手を**しらみつぶし**に調べること。

推論

ルールをたどって**結論**を出すこと。

迷路を解く、パズルを解く。これがものすごく話題になりました。ところが、**現実の問題がまるで解けなかった**のです。迷路は解けるのに、病気の診断はできない。**おもちゃのような問題しか扱えない**ことから、これを**トイプロブレム**（おもちゃの問題）と呼びます。

期待が一気にしぼみ、1970年代に**1回目のAIの冬**が来ます。研究費が止まりました。

試験ポイント

- 第一次＝**探索と推論**。**トイプロブレム**で行き詰まり、**1970年代に1回目の冬**。

第二次AIブーム（1980年代）

今度の主役は**知識**です。「現実が解けないのは、AIが知識を持っていないからだ」と考え、**専門家の知識**をルールとして大量に詰め込みました。これが**エキスパートシステム**です。

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**。

医者¹の知識を詰め込んだ診断システム、設備の故障診断などが実際に企業で使われ、今度こそ本物だと言われました。ところが、また壁にぶつかります。

知識を人間が書き写す作業が、終わらないのです。専門家は「なんとなく分かる」ことを言葉にできません。言葉にできないものは、ルールに書けない。これを**知識獲得のボトルネック**^{かくとく}といいます。さらに、ルールが増えるほど矛盾^{むじゅん}が生じ、手入²れが追いつかなくなりました。

こうして1990年代、**2回目のAIの冬**を迎えます。2回も冬を越している分野は、そう多くありません。

試験ポイント

- 第二次＝**エキスパートシステム**（知識）。**知識獲得のボトルネック**で行き詰まり、**1990年代に2回目の冬**。

第三次AIブーム（2000年代～現在）

主役は**ビッグデータ**と機械学習、そして**ディープラーニング**です。変わったことが3つあります。

- インターネットの普及で**大量のデータ**が手に入るようになった＝**ビッグデータ**
- コンピュータの**計算力**が上がった
- **ディープラーニング**という手法が出てきた

つまり、**人間が知識を書き写す**という第二次の敗因を、**データから自分で学ばせる**ことで回避したわけです。決定打は**2012年**。画像認識のコンテストで、ディープラーニングを使ったチームが圧勝しました。ここから一気に現在の生成AIへつながります。

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。

AIの冬

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。

必ず出る

- **第一次＝探索と推論**。トイプロブレムで冬。
- **第二次＝エキスパートシステム**。知識を書き写せなくて冬。
- **第三次＝ビッグデータとディープラーニング**。いまここ。
- **ブームは3回、冬は2回**。この数字がそのまま問われます。

確認② AIの種類と歴史の確認問題

1-3と1-4の範囲から3問です。本文を見ずに答えてください。

確認問題②

第1問 AIの4つのレベル、レベル3とレベル4の違いは？

答

特微量を人間が決めるか、AIが自分で見つけるか。

レベル3は人間が指定、レベル4はAIが自分で見つける。第1章で最も問われる論点です。

確認問題②

第2問 実用化されているAIは、ANIとAGIのどちらか？

答

すべて**ANI（弱いAI）**。

AGI（強いAI）は**まだ実現していません**。「すでに実現している」は誤りです。

確認問題②

第3問 AIブームは何回、冬は何回あったか。それぞれのブームの主役も答えよ。

答

ブーム3回、冬2回。

第一次＝**探索と推論**／第二次＝**エキスパートシステム**／第三次＝**ビッグデータとディープラーニング**。

1-5 シンギュラリティ

第1章の最後のテーマです。**シンギュラリティ**、日本語で**技術的特異点**。AIが人間の知能を超え、それ以降が予測できなくなる転換点のことを指します。

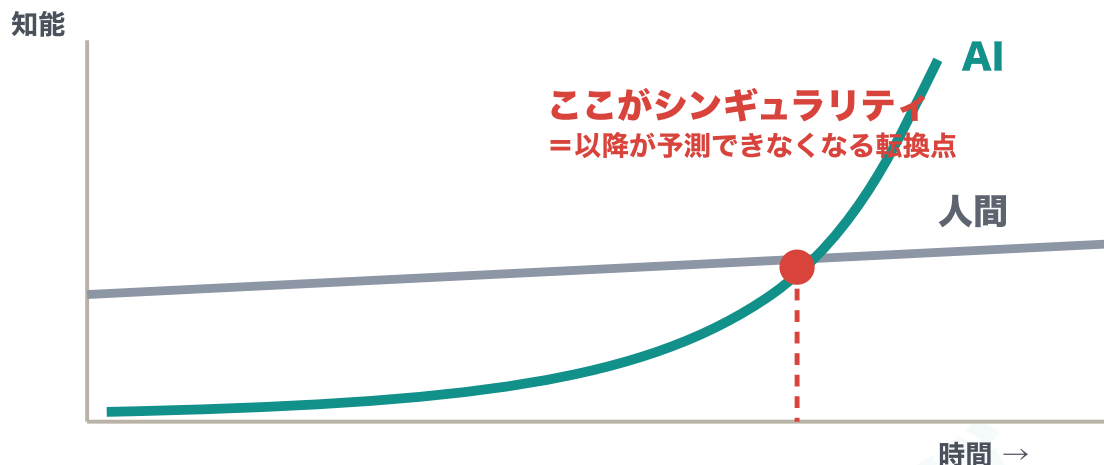


図1-10 シンギュラリティに関わる2人。役割が違うので、必ず分けて覚える

ここで**人物が2人**出ます。片方だけ覚えて落とす人が非常に多いところです。

1人目 | ヴァーナー・ヴィンジ

SF作家で、数学者でもあります。「**技術的特異点**」という**概念を最初に広めた人物**です。

2人目 | レイ・カーツワイル

シンギュラリティの**到来を2045年と予測**した人物です。ここから**2045年問題**と呼ばれています。

シンギュラリティ 技術的特異点

AIが人間の知能を超え、それ以降の変化が**予測できなくなる**とされる転換点。

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。

試験ポイント

- **概念を広めたのがヴィンジ、2045年と予測したのがカーツワイル。**役割が違います。
- 2人の名前と役割の**組み合わせを入れ替えた選択肢**が出ます。

ここが引っかけ

これは**予測であって、決定事項ではありません**。「来ない」と考える研究者もいます。「2045年にシンギュラリティが**到来する**」と**断定している**選択肢は疑ってください。

AI効果

第1章の最後に、もう1つ言葉が出ます。**AI効果**です。

これは、**AIが何かをできるようになった瞬間、人間が「それは知能ではない」と言い出す現象**を指します。

たとえば昔は「チェスで人間に勝てたら、それは知能だ」と言われていました。実際に勝った途端、「あれはただの計算だ」と言われるようになりました。将棋でも囲碁でも同じことが起き、翻訳でも絵でも、いまでも同じことが起きています。**できるようになると、その瞬間に「知能」の定義から外される**のです。

チェス

勝てたら知能だ
→「ただの計算だ」

将棋・囲碁

同じことが起きた

翻訳・絵

いまでも起きている

お気づきでしょうか。1-1の「AIに定義がない」に戻ってきました。定義がないからこそ、こういうことが起きるのです。

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。

試験ポイント

- 試験では「AI効果」という名前と、「**できるようになると知能と見なされなくなる現象**」という説明の対応で出ます。

確認③

総まとめの確認問題

第1章の全体から3問です。ここまで解けたら、第1章は仕上がっています。

確認問題③

第1問 ノーフリーランチ定理とはどういう定理か。ひと言で。

答

あらゆる問題に万能なアルゴリズムは存在しないという定理。

「万能なアルゴリズムが存在する」という選択肢は、この定理で切ります。

確認問題③

第2問 脳のニューロンとシナプスは、AI側の何にあたるか。また学習とは何をすることか。

答

ニューロン → **人工ニューロン（ノード）**、シナプス → **重み**。

学習とは、**重み（情報の重みづけ）を調整すること**。

確認問題③

第3問 シンギュラリティに関わる2人の人物と、それぞれの役割は？

答

ヴァーナー・ヴィンジ＝技術的特異点という概念を広めた。

レイ・カーツワイル＝到来を**2045年**と予測した（2045年問題）。

まとめ

これが言えれば合格ライン

第1章でやったことを、5つに畳みます。**それぞれ声に出して言えるか**を確かめてください。言えないものがあれば、その節に戻ります。

① AIの定義

- **AIに定義はない**。知能の定義が決まっていないから。
- 言葉が生まれたのは**1956年のダートマス会議**（ジョン・マッカーシー）。
- **AIはソフトウェア、ロボットはハードウェア**。体は必須ではない。

② 知能をもたらす仕組み

- 仕組みは2つ＝**ルールベース**（人間が書く）と**機械学習**（データからAIが見つける）。成果物が**学習済みモデル**。
- 手法は4つ＝**教師あり・教師なし・強化・半教師あり**。教師なしの中に**クラスタリングと次元削減**。
- 万能な手法は無い＝**ノーフリーランチ定理**。
- 脳を真似たのが**ニューラルネットワーク**（ニューロン→ノード、シナプス→**重み**）、層を深くしたのが**ディープラーニング**。
- 学習しすぎると**過学習**、対策は**正則化・ドロップアウト**。使い回すのが**転移学習**。

③ AIの種類

- 4つのレベルは、**上がるほど人間が教える量が減る**。
- **レベル3と4の境目は特徴量**（人間が決めるか、AIが見つけるか）。
- 実用化されているのは**全部、弱いAI（ANI）**。**強いAI（AGI）は未実現**。

④ AIの歴史

- **ブーム3回、冬2回**。
- 第一次＝**探索と推論**／第二次＝**エキスパートシステム**／第三次＝**ビッグデータとディープラーニング**。

⑤ シングularity

- 概念が**ヴィンジ**、**2045年**と予測したのが**カーツワイル**（2045年問題）。
- そして**AI効果**——できるようになると知能と見なされなくなる。

最後に

定義がないものについて、ここまで覚えました。第1章は「考える」より「思い出す」章です。忘れたら戻る、それだけで確実に得点になります。

覚えたかどうかのチェックリスト（39語）

意味を**説明できる**語にチェックを入れてください。読めるだけ・見たことがあるだけでは足りません。

- | | | |
|-------------------------------------|---------------------------------------|---|
| ☐ AI（人工知能） | ☐ シナプス | ☐ 第一次AIブーム |
| ☐ ダートマス会議 | ☐ 人工ニューロン（ノード） | ☐ 探索 |
| ☐ ルールベース | ☐ ニューラルネットワーク | ☐ 推論 |
| ☐ 機械学習 | ☐ ディープラーニング | ☐ 第二次AIブーム |
| ☐ 学習済みモデル | ☐ 重み | ☐ エキスパートシステム |
| ☐ 教師あり学習 | ☐ 情報の重みづけ | ☐ AIの冬 |
| ☐ 教師なし学習 | ☐ 過学習 | ☐ 第三次AIブーム |
| ☐ クラスタリング | ☐ 正則化 | ☐ ビッグデータ |
| ☐ 次元削減 | ☐ ドロップアウト | ☐ シングularity（技術的特異点） |
| ☐ 強化学習 | ☐ 転移学習 | ☐ ヴァーナー・ヴィンジ |
| ☐ 半教師あり学習 | ☐ 特徴量 | ☐ レイ・カーツワイル |
| ☐ ノーフリーランチ定理 | ☐ 弱いAI（ANI） | ☐ 2045年問題 |
| ☐ ニューロン | ☐ 強いAI（AGI） | ☐ AI効果 |

用語集

第1章の重要用語 39語

公式シラバス（2026年2月試験より適用）の第1章に**キーワード**として挙げられている**39語**です。増減はありません。試験直前は、ここだけを見返してください。

1-1 AIの定義（2語）

AI（人工知能） p.8

人間の知能を人工的に実現する技術。合意された定義は存在しない。

ダートマス会議 p.8

1956年。ジョン・マッカーシーが「AI」という語を初めて使った研究会。

1-2 AIに知能をもたらす仕組み（21語）

ルールベース p.10

人間がルールを全部書く方式。説明できるが、書ける量に限界。

機械学習 p.11

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.11

学習を終えた成果物。新しいデータを入れると答えを返す。

教師あり学習 p.209

正解付きのデータで学習する。迷惑メール判定、売上予測。

教師なし学習 p.209

正解のないデータから構造を見つける。

クラスタリング p.209

似たものどうしをグループ分けする。**教師なし学習**。

次元削減 p.209

項目を減らし、大事な軸だけ残す。**教師なし学習**。

強化学習 p.209

正解ではなく報酬を与え、試行錯誤で良い行動を学ばせる。

半教師あり学習 p.209

正解付きが少量、正解なしが大量という状況で使う。

ノーフリーランチ定理 p.13

あらゆる問題に万能なアルゴリズムは存在しない。

ニューロン p.209

人間の脳の神経細胞。

シナプス p.209

ニューロンどうしのつなぎ目。AI側の「重み」に対応。

人工ニューロン（ノード） p.210

ニューロンを数式で模したもの。

ニューラルネットワーク p.210

ノードを層に並べてつないだ全体。

ディープラーニング p.14

ニューラルネットワークの層を深く積んだ手法。

重み p.210

つなぎ目ごとの「どれだけ重視するか」の数字。学習＝重みの調整。

情報の重みづけ p.210

重みを変えること。すなわち「学習する」と同じ。

過学習 p.15

学習データに合わせすぎ、未知のデータに弱くなる状態。

正則化 p.210

罰則を加えてモデルが複雑になりすぎるのを抑える。過学習対策。

ドロップアウト p.210

学習のたびにノードを一部休ませる。過学習対策。

転移学習 p.17

学習済みモデルを別の目的に作り替える。少ないデータで済む。

1-3 AIの種類（3語）

特徴量 p.15

「どこに注目すべきか」の着目点。**レベル3と4の境目**。

弱いAI（ANI） p.20

特定の課題に特化したAI。実用化されているのは全部これ。

強いAI（AGI） p.20

汎用的で意識や自我を持つとされるAI。未実現。

1-4 AIの歴史（8語）

第一次AIブーム p.211

1950年代後半～60年代。探索と推論。

探索 p.211

可能な手をしらみつぶしに調べること。

推論 p.211

ルールをたどって結論を出すこと。

第二次AIブーム p.211

1980年代。知識＝エキスパートシステム。

エキスパートシステム p.22

専門家の知識をルールとして詰め込んだシステム。

AIの冬 p.22

期待がしぼんだ時期。1970年代と1990年代の2回。

第三次AIブーム p.212

2000年代～現在。ビッグデータとディープラーニング。

ビッグデータ p.22

インターネットの普及で得られるようになった大量のデータ。

1-5 シンギュラリティ（5語）

シンギュラリティ（技術的特異点） p.24

AIが人間の知能を超え、以降が予測できなくなる転換点。

ヴァーナー・ヴィンジ p.212

技術的特異点という概念を最初に広めたSF作家・数学者。

レイ・カーツワイル p.212

シンギュラリティの到来を2045年と予測した。

2045年問題 p.24

カーツワイルの予測に由来する呼び名。

AI効果 p.25

できるようになった途端、知能と見なされなくなる現象。

次回 第2章 生成AIそのものへ

第1章、おつかれさまでした。次は第2章、**生成AIそのもの**に入ります。生成モデルの誕生から、CNN、VAE、GAN、RNN、LSTM、そして **Transformer** まで。アルファベットが一気に増えますが、第1章と同じで「**何ができるようになったか**」で並べれば全部つながります。

第2章へ進む前に

- **確認問題9問**が何も見ずに解けて、**用語集39語**のうち説明できない語が3つ以下なら、そのまま第2章へ進んで大丈夫です。

⇒ 本書はYouTubeチャンネル「**AI資格ラボ**」の配布テキストで、**生成AIパスポート 公式シラバス（2026年2月試験より適用）**の第1章に対応しています。試験の最新情報・出題範囲は必ず公式サイトの原本をご確認ください。学習の補助を目的としたもので、合格を保証するものではありません。

第2章① 生成AI（仕組み編）



この章の解説動画

<https://youtu.be/UU6hqXCCqnY>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに この本の使い方

この本は、YouTubeチャンネル「AI資格ラボ」の動画「**第2章 生成AI ①仕組み編**」に対応した、配布用のテキストです。動画を見ていない方でも、これ一冊で第2章の前半が最後まで読み切れるように書いてあります。

この本が扱う範囲

生成AIパスポートの第2章は、公式シラバス（出題範囲表）で**62語**あります。第1章の39語の1.6倍で、1回で扱うには多すぎます。そこで**2冊に分けました**。

	扱う範囲	語数
①	生成AIの誕生まで（この本）	32語
②	ChatGPT（次の本）	30語

この本は**①の32語**を扱います。生成モデルがどう生まれて、どう進化して、**Transformer**にたどり着いたか。そこまでの範囲です。

→ **第1章でやった言葉は、やり直しません**。ニューラルネットワーク、隠れ層、ディープラーニング、特徴量。このあたりは第1章のテキストと動画にあります。まだの方は先にそちらを読んでください。

この本の構成

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま問題になる所です。試験直前はここだけ拾い読みできます。
- **△マークの枠**（赤い枠）……**間違えやすい所・引っ掛けが仕込まれる所**です。
- **用語**（紫の帯）……シラバスに名前が載っている用語です。全部で32語あります。
- **確認問題**（緑の枠）……テーマの区切りに3問ずつ、合計9問あります。
- **チェックリスト**（巻末）……32語を自分で説明できるか、最後に確かめられます。

おすすめの進め方

1. まず**本文を最後まで通して読む**。分からない所で止まらず、飛ばして進んでかまいません。
2. 次に**確認問題9問**を、本文を見ずに解く。ここで落ちた所が、あなたの弱点です。
3. 落ちた所だけ本文に戻り、**自分の言葉で言い直す**。読むより言い直したほうが残ります。
4. 試験直前は、**試験ポイントの枠と、巻末の用語集だけ**を見返す。

第2章の前半は、どういう章か

先に言っておきます。この章は、**アルファベットが襲ってきます**。CNN、VAE、GAN、RNN、LSTM、Transformer。名前だけ並べられると、それだけで帰りたくなります。

ただ、この6つには**共通する事情**があります。**全部が「前のやつの弱点をつぶした結果」**なのです。順番に出てきたことには理由があります。

だから、**理由で覚えると6つが1本の線になります**。逆に、名前だけで覚えようとする**確実に混ざります**。この本は、その線を引くために書いてあります。

この本・動画・練習問題の回し方

この本は**単体でも読めます**が、動画と練習問題と合わせると回転が速くなります。おすすめは次の4ステップです。

1. **この本（章ごとのPDF）を読む**。無料です。
2. **その章の動画を見る**。本と同じ順番・同じ図で説明しています。
3. **ホームページに戻って、その章の練習問題を解く**。登録もお金も要りません。
4. **間違えた問題の解説から、動画のその場面へ飛ぶ**。「どこで説明していたっけ」と探さずに済みます。

解けるようになったら、次の章のPDFを落として同じことを繰り返します。**練習問題は全5章ぶん、模擬試験まで用意してあります**ので、動画が追いつく前でも先に進めます。

- ホームページのリンクは、動画の**概要欄**にあります。
- 印刷するときは**A4・実物大（100%）**で。図の中の文字まで読める大きさに作ってあります。

AI資格ラボ

目次 この本の中身

はじめに	この本の使い方	30
第2章①	この章の全体像	34
2-1	生成AIとは何か 識別するか、生成するか／生成AIの中身も機械学習	36
2-2	生成モデルのはじまり ボルツマンマシン／制約付きボルツマンマシン／自己回帰モデル	38
	確認問題①（2-1・2-2）	39
2-3	画像の系譜 CNNと畳み込み／VAE／GAN	41
	確認問題②（2-3）	44
2-4	文章の系譜 RNN／LSTM／Transformer／Attention／位置エンコーディング	45
	確認問題③（2-4）	49
2-5	Transformerからの枝分かれ GPTモデル／BERTモデル／RoBERTaとALBERT	51
まとめ	第2章① これが言えれば合格ライン	54
用語集	第2章①の重要用語 32語	56

動画と合わせて使う

- この本は、YouTube「AI資格ラボ」の動画「【生成AIパスポート】第2章 生成AI ①仕組み編」と同じ順番・同じ図で作ってあります。
- 動画で流れをつかみ、この本で確かめる。あるいはこの本を先に読んでから動画を見る。どちらでも構いません。
- 印刷はA4・実物大（100%）で。図の中の文字まで読める大きさに作ってあります。

第2章① この章の全体像

最初に、これから通る道を1枚で見っておきます。細かい所は分からなくて構いません。**6つの名前が縦に並んでいる理由**だけ、頭に入れてください。

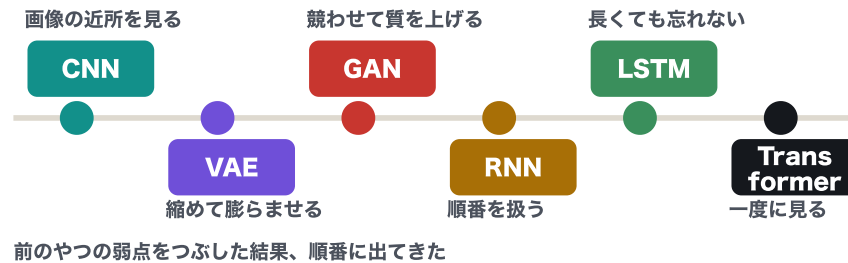


図2-1 6つは順番に出てきた。それぞれが「前のやつの弱点」をつぶしている

左から右へ、時代が進みます。**CNN**は画像の近所を見る仕掛け、**VAE**は縮めて膨らませる作り方、**GAN**は競わせる作り方、**RNN**は順番を扱う仕掛け、**LSTM**は長くても忘れない改良、そして**Transformer**が「順番に読む」こと自体をやめました。

この本は、この線を左から順にたどります。テーマの分け方は次のとおりです。

	テーマ	ここで問われること
2-1	生成AIとは何か	従来のAIとの違い。答え方は一つ
2-2	生成モデルのはじまり	ボルツマンマシンの立ち位置、何を制約したか
2-3	画像の系譜	畳み込みとは何か、VAEとGANの違い
2-4	文章の系譜	RNNの弱点、Transformerが得たものと失ったもの
2-5	枝分かれ	GPTとBERTの役割の違い、MLMとNSP

この章の勉強のコツ

- **名前を覚えるのではなく、順番と理由を覚える。**「前のやつのが不満だったのか」が言えれば、名前は後からついてきます。
- **画像の系譜と文章の系譜は、別の線として覚える。**途中で合流するのはTransformerからです。
- **似た名前は必ず対で覚える。**エンコーダとデコーダ、生成器と識別器、MLMとNSP、RoBERTaとALBERT。試験はこの対を入れ替えて出してくれます。

名前の読み方

この章はアルファベットの略語が続きます。読み方が分からないまま進むと頭に入りにくいので、先に確かめておきます。

書き方	読み方	日本語の名前
CNN	シー・エヌ・エヌ	畳み込みニューラルネットワーク
VAE	ブイ・エー・イー	変分自己符号化器
GAN	ガン	敵対的生成ネットワーク
RNN	アール・エヌ・エヌ	回帰型ニューラルネットワーク
LSTM	エル・エス・ティー・エム	長・短期記憶
Transformer	トランスフォーマー	(日本語名はなし)

→ **GAN**だけ「ガン」と一続きで読みます。ほかは1文字ずつです。

2-1 生成AIとは何か

まず言葉から入ります。**生成AI**。英語で**ジェネレーティブAI**。シラバスにはこの両方の書き方で載っているのですが、**どちらで出ても同じもの**だと思ってください。

生成AI（ジェネレーティブAI） Generative AI

文章・画像・音声・動画など、**もとは無かったものを作り出すAI**。従来の「見分けるAI」に対する呼び方。

識別するか、生成するか

では、何が「生成」なのでしょう。第1章で扱ったAIは、**見分けるAI**でした。この写真は猫か犬か。このメールは迷惑メールか。**答えを選ぶ**仕事です。

生成AIは違います。**無かったものを作ります**。文章、画像、音声、動画。選ぶのではなく、出す。ここが決定的な差です。



図2-2 従来のAIは決まった選択肢から選ぶ。生成AIは答えが無限にある

図の左を見てください。従来のAIは、**答えが決まった選択肢の中にあります**。猫か犬か、迷惑メールかどうか。選択肢の外の答えは出てきません。

右が生成AIです。出てくるのは**新しい文章、新しい画像**。**答えは無限にあります**。だから「作った」と言えるわけです。

試験ポイント

- 試験では「**従来のAIと生成AIの違いは何か**」という聞き方をされます。
- 答え方は一つで足ります。**識別するか、生成するか**。見分けるのが従来、無いものを作るのが生成です。

生成AIも、中身は機械学習

ここで一つ、勘違いされやすい所を潰しておきます。**生成AIも、中身は機械学習**です。第1章でやった「データから学ぶ」の延長線上にあります。

まったく別の生き物が急に現れたわけではありません。**学び方と出し方が変わっただけ**です。この感覚を持っておくと、このあと出てくる6つの名前が「機械学習のいろいろな作り方」として並んで見えてきます。

生成AIが作るもの

「生成」と聞くと文章を思い浮かべがちですが、扱う対象^{たいしやう}は文章だけではありません。シラバスでも、生成AIは**作り出すもの全般**を指す言葉として使われています。

作るもの	たとえば
文章	質問への答え、要約、翻訳、プログラム
画像	写真のような絵、イラスト、デザイン案
音声	読み上げ、歌声、効果音
動画	短い映像、アニメーション

この本で扱うのは、このうち**仕組みがどう生まれてきたか**です。どの種類を作る場合でも、土台になっている考え方は共通しています。

試験ポイント

- **生成AIとジェネレーティブAIは同じもの**。シラバスに両方の表記があるので、どちらで出ても迷わないこと。
- 「作り出す対象は文章だけ」という選択肢は**誤り**です。画像・音声・動画も含みます。

△ △ **引っかけ** 「生成AIは機械学習とは別の技術である」という選択肢は**誤り**です。生成AIは機械学習の一部です。

2-2 生成モデルのはじまり

ここから、生成モデルの歴史を順にたどります。最初に出てくるのは、**1985年**に提案されたモデルです。

ボルツマンマシン

ボルツマンマシン Boltzmann Machine

1985年に提案された生成モデルの祖先^{そせん}。データがどうやってできているかを確率で覚えようとした。ユニットが全部つながっているため計算量が大きく、実用にはならなかった。

名前はいかついですが、やっていることは一言で言えます。**データがどうやってできているかを、確率で覚えようとした**。それだけです。

第1章でやったニューラルネットワークは、入力から出力へ一方通行に信号が流れました。ボルツマンマシンは違います。ユニットどうしが全部つながっていて、信号が行ったり来たりします。その揺れ動きが落ち着いたところが「データらしい状態」だ、という考え方です。

ただし、ここが大事なところですよ。このモデルは実用になりませんでした。全部がつながっているで、**計算量が爆発する**のです。

制約付きボルツマンマシン

重すぎたなら、どうするか。つながりを減らせばいい。それが制約付きボルツマンマシンです。

制約付きボルツマンマシン Restricted Boltzmann Machine

ボルツマンマシンから同じ層の中のつながりを全部切ったもの。可視層と隠れ層の2層にし、層と層の間だけをつなぐ。計算が現実的な重さになり、ディープラーニングの実用化を支えた。

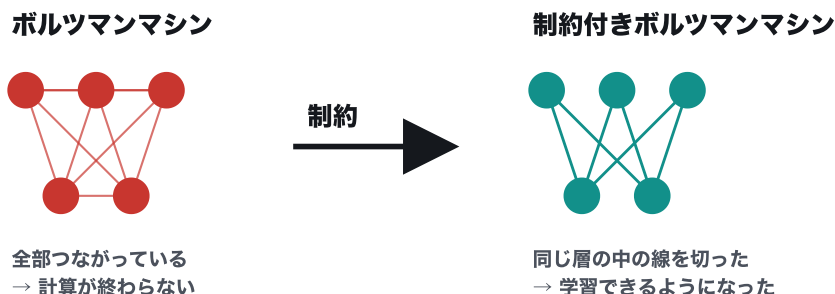


図2-3 同じ層の中の線を切っただけで、学習できるようになった

何を制約したのか。**同じ層の中のつながりを、全部切りました**。可視層と隠れ層の2層にして、層と層の間だけをつなぐ。層の中はつながりません。これだけで計算が現実的な重さになりました。

そして、これが効いてくるのが**ディープラーニング**です。層を深くすると学習がうまくいかない時期が長く続いたのですが、**制約付きボルツマンマシン**を積み上げて、先に軽く学習させておくという手が使われました。つまり、**深い層を実用にした立役者のひとり**です。

試験ポイント

- **ボルツマンマシンは重い。制約をつけたら動いた。**この2文で立ち位置を覚えます。
- 試験では「**制約とは何を制約したのか**」が問われます。答えは**層の中のつながり**です。
- ボルツマンマシンは「生成モデルの**祖先**」として問われます。実用になったかどうかまで含めて覚えてください。

自己回帰モデル

次は**自己回帰モデル**です。これは仕組みの名前というより、**やり方の名前**です。

自己回帰モデル Autoregressive Model

自分が出した出力を、次の入力として使うやり方。1つずつ順番に生成していく。GPTがこの方式。

どういうやり方か。**自分が出した答えを、次の入力にする**。それを繰り返します。文章で考えると一発で分かります。

1語ずつ、前を見ながら次を決める

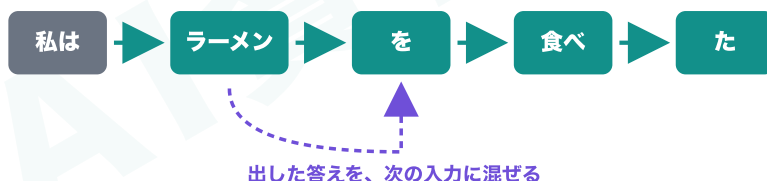


図2-4 1語出しては、それも読み込んで、また次を1語出す

「私は」と出したら、それも読み込んで次に「ラーメン」を出す。さらにそれも読み込んで「を」を出す。**1個出しては、それも読み込んで、また次を1個出す。**この繰り返しです。

……これ、どこかで見た動きではないでしょうか。**ChatGPTが、文字を1つずつ出していくあの動き。**あれです。**GPTは自己回帰モデルです。**次の本への伏線ふくせんですが、ここでつないでおきます。

試験ポイント

- 問われるのは**1つずつ順番に生成する**というところです。
- 逆に言うと、**まとめて一気にには出せません**。だから長い文章は時間がかかります。理屈が分かったら、あの表示の遅さにも納得がいきます。

確認問題① 2-1・2-2 のふりかえり

本文を見ずに、頭の中で答えてから下の答えを見てください。

確認問題①

第1問 従来のAIと生成AIの、いちばん大きな違いは何ですか。

答

識別するか、生成するか。

見分けるのが従来のAI、無かったものを作るのが生成AIです。

確認問題①

第2問 制約付きボルツマンマシンは、何を制約したのですか。

答

同じ層の中のつながり。

層と層の間だけを残したので、計算が現実的な重さになりました。

確認問題①

第3問 自分が出した出力を次の入力に使って、1つずつ生成していくモデルを何とといいますか。

答

自己回帰モデル。

GPTがこの方式です。

⇒ 3問ともできましたか。できなかった所は、いま本文に戻って**自分の言葉で言い直して**ください。読み返す
より残ります。

2-3 画像の系譜

ここから**画像**の話に入ります。この節には3つの名前が出てきます。**CNN・VAE・GAN**。最初のCNNは「見分ける」側、あとの2つは「作る」側です。

CNNと畳み込み

CNN（畳み込みニューラルネットワーク） Convolutional Neural Network

畳み込みを使って画像の特徴を取り出すニューラルネットワーク。画像を得意とする。

畳み込み Convolution

小さな窓を画像の上でずらしながら当て、**局所的な特徴を取り出す操作**。CNNの中心となる仕組み。

第1章で、AIが画像を認識する仕組みに触れました。その主役が**CNN**、日本語で**畳み込みニューラルネットワーク**です。

何が新しかったのでしょうか。それまでは、画像の点を**全部そのまま**ニューラルネットワークに入れていました。でも画像は、**隣り合った点どうしに意味がある**のです。目の形、輪郭の向き、模様。近所を見ないと分かりません。

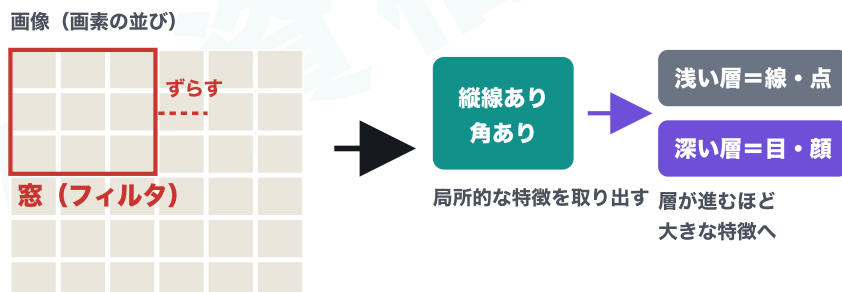


図2-5 小さな窓をずらして当て、局所的な特徴を取り出す。層が進むほど大きな特徴へ

そこで**畳み込み**です。小さな窓を、画像の上で少しずつずらしながら当てていきます。窓の中を見て、「ここに縦線があるな」「ここは角だな」と反応する。その反応を集めた地図を、また次の層で見る。

こうすると、浅い層は線や点、深い層は目や顔というふうに、だんだん大きな特徴になっていきます。第1章で**特徴量**をやりましたね。**CNNは、その特徴量を自分で見つけるための仕掛けだ**と思ってください。

試験ポイント

- 畳み込みは「局所的な特徴を取り出す操作」。この言い回しで出ます。
- CNNは画像を得意とする。この2つだけで十分戦えます。

△ △ 引っかけ 畳み込みは**画像を小さくするための操作ではありません**。あくまで**特徴を取り出す**操作です。

VAE（変分自己符号化器）

画像を**見分ける**話は済みました。ここからが**作る**話です。まずVAE、日本語で**変分自己符号化器**。

仕組みは、実はとても素直です。**入り口で縮めて、出口でふくらませる**。それだけです。

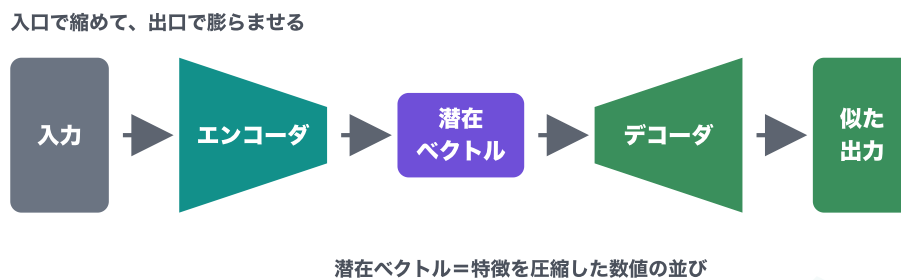


図2-6 縮める係がエンコーダ、ふくらませる係がデコーダ。間の短い数字の並びが潜在ベクトル

VAE（変分自己符号化器） Variational Auto Encoder

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。

エンコーダ Encoder

入力を**圧縮する側**（入口）。

デコーダ Decoder

圧縮されたものから**復元する側**（出口）。

潜在ベクトル Latent Vector

エンコーダが作る**短い数字の並び**。特徴を圧縮した表現。

縮める係が**エンコーダ**、ふくらませる係が**デコーダ**。そして、縮めた結果の**短い数字の並び**が**潜在ベクトル**です。

たとえば猫の写真を、**猫らしさの成分だけに絞った数字**^{しぼ}にします。その数字から、もう一度、猫の絵を組み立て直す。

……ここで、当然の疑問が出ます。**元に戻すだけなら、意味がないのでは？** そのとおりです。戻すだけなら意味がありません。

ノイズ Noise

潜在ベクトルに**わざと加える揺らぎ**。これがあるので、元と違う新しいデータが出てくる。

VAEが賢いのはここからです。**その数字を、わざとちょっと揺らします**。この揺らぎが**ノイズ**です。少しずらした数字から組み立てると、**元の猫ではない、新しい猫が出てきます**。これが「生成」です。

試験ポイント

- 覚え方は**エンコーダで縮めて、潜在ベクトルを揺らして、デコーダで戻す**。
- エンコーダ・潜在ベクトル・デコーダの3語はセットで出ます**。入口と出口を入れ替えた選択肢が定番です。

GAN（敵対的生成ネットワーク）

同じ「作る」でも、まったく違う発想で来たのが**GAN**です。**敵対的生成ネットワーク**。名前のとおり、**敵対**します。

GAN（敵対的生成ネットワーク） Generative Adversarial Network

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 Generator

GANの中で**偽物を作る側**。

識別器 Discriminator

GANの中で**本物か偽物かを見分ける側**。

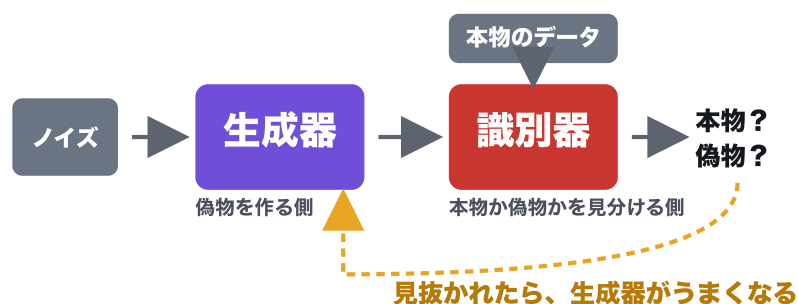


図2-7 生成器と識別器の追いかっこ。見抜かれるほど生成器がうまくなる

中に2人います。1人目が**生成器**。**偽物を作る係**です。2人目が**識別器**。**それが本物か偽物かを見破る係**です。この2人を**戦わせます**。

生成器は「バレないように」うまくなろうとする。識別器は「見逃さないように」厳しくなろうとする。**追いかっこをしているうちに、生成器の腕が上がっていく**。

よく偽札^{にせさつ}作りと警察にたとえられます。偽札がうまくなれば警察も目が肥^こえる。目が肥えれば偽札もさらにうまくなる。最終的に、人間には見分けがつかない偽札ができあがる。物騒^{ぶっそう}なたとえですが、試験ではこの構図がそのまま問われます。

試験ポイント

- **生成器と識別器。2つが競い合う。** この2語は必ず対で覚えます。
- VAEとの違いも聞かれます。**VAEは縮めて戻す。GANは2つを競わせる。** ここだけ押さえてください。

△ △ **引っかけ** 「生成器と識別器が**協力して**同じ目標に向かう」という選択肢は誤りです。**競わせる**のが要点です。

確認問題② 2-3 のふりかえり**確認問題②**

第1問 画像の中の、近くの点どうしの特徴を取り出す操作を何といいますか。またそれを使うネットワークは。

答

畳み込み/CNN。

畳み込みは「局所的な特徴を取り出す操作」です。

確認問題②

第2問 VAEで、入力を縮めた短い数字の並びを何といいますか。

答

潜在ベクトル。

縮めるのがエンコーダ、戻すのがデコーダです。

確認問題②

第3問 GANの中にいる2つのネットワークの名前は何ですか。

答

生成器と識別器。

偽札屋と警察、と覚えると混ざりません。

⇒ ここまでが**画像**の系譜です。次から**文章**の系譜に移ります。

2-4 文章の系譜

文章と、音声と、株価。この3つに共通するものは何でしょうか。**順番があること**です。ここからは、順番のあるデータをどう扱うかの話になります。

RNNと隠れ層・リカレント層

シーケンスデータ Sequence Data

文章・音声・株価など、**順番に意味があるデータ**。日本語では系列データ。

こういうデータを**シーケンスデータ**、日本語で系列データといいます。「私は／ラーメンを／食べた」を並べ替えたら意味が壊れますよね。**順番が意味を持つ**のです。

ところが、ここまでのニューラルネットワークは**入力を一度にまとめて受け取る**作りでした。順番という情報が、そもそも入りません。

RNN（回帰型ニューラルネットワーク） Recurrent Neural Network

隠れ層の出力を次の入力に混ぜて、順番のあるデータを扱えるようにしたモデル。

隠れ層 Hidden Layer

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 Recurrent Layer

前の状態を次へ戻す構造を持った層。RNNの中心となる仕組み。

順番のあるデータ（シーケンスデータ）を1語ずつ読む



前の状態を次へ渡す構造=リカレント層

図2-8 隠れ層の出力を次へ渡す。この「戻す」構造がリカレント層

そこで出てきたのが**RNN、回帰型ニューラルネットワーク**です。何をしたか。**隠れ層の出力を、次の入りに混ぜて、もう一度同じ層に戻した**。この「戻す」構造を持った層を**リカレント層**といいます。

1語読む。覚えておく。次の語を読むときに、さっきの記憶も一緒に見る。**メモを取りながら読んでいる**イメージです。第1章で出てきた**隠れ層**、つまり入力層と出力層のあいだの層が、RNNでは**記憶係**をやっていると思ってください。

LSTM（長・短期記憶）

そのRNNですが、**弱点がありました。長い文章になると、前のほうを忘れます。**

「私は、去年の夏に、家族と一緒に、沖縄の海に行って、そこで食べたソーキそばが……」。この時点で、**主語が誰だったか怪しくなります。**なぜ忘れるかというと、記憶をぐるぐる回しているうちに**信号が薄まっていく**からです。

LSTM（長・短期記憶） Long Short-Term Memory

RNNに**忘れるかどうかを決めるスイッチ**を足したもの。長い系列でも前のほうを覚えていられる。

RNN：長いと前を忘れる



「去年の夏に…沖縄で…食べた」の「去年」が薄れていく

LSTM：残すものと捨てるものを決める

覚えておく（長期） + いま使う（短期） = 長・短期記憶

図2-9 覚えておく（長期）と、いま使う（短期）を分けて持つ

これを解決したのが**LSTM、長・短期記憶**です。何を足したか。**忘れるかどうかを決めるスイッチ**です。大事なことは長く持つておく。どうでもいいことは捨てる。このスイッチのおかげで、**長い系列でも前のほうを覚えていられる**ようになりました。

試験ポイント

- RNNは順番を扱える。ただし長いと忘れる。忘れないようにしたのがLSTM。この3文で足ります。
- 名前がそのまま答えです。**長・短期記憶**=長い記憶と短い記憶を分けて持つ、と読めます。

Transformerモデル

そして、この章の主役です。**Transformerモデル**。2017年に出ました。

Transformerモデル Transformer

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

何がすごかったのか。……**RNNをやめました。**雑に聞こえますが本当です。**順番に読むのを、やめたのです。**

RNN : 1語ずつ順番に**Transformer : 全部を一度に見る**

2017年。いま名前を聞く生成AIは、ほぼ全部この子孫

図2-10 RNNは前が終わるまで次に進めない。Transformerは全部を一度に見る

RNNもLSTMも、**1語ずつ順番に読む**しかありませんでした。順番に読むということは、**前の語を読み終わるまで次に進めない**ということです。つまり**手分けができない**。学習にものすごく時間がかかります。

Transformerは、**文の全部の語を、一度に見ます**。そのうえで、語と語の関係を計算で出す。だから**手分けができる**。大量のデータを、現実的な時間で学習できるようになりました。

これが効きました。**いま名前を聞く生成AIは、ほぼ全部Transformerの子孫です**。GPTも、BERTも、Geminiも、Claudeも。

アーキテクチャ Architecture

モデルの構造・設計そのもの。建物の設計図と同じ意味。「Transformerというアーキテクチャ」のよ
うに使う。

ちなみに**アーキテクチャ**という言葉がシラバスに出てきますが、これは**モデルの構造・設計そのもの**と
いう意味です。建物の設計図と同じ言葉で、深い意味はありません。

Attention (自己注意力)

ここからがTransformerの**心臓部**^{しんそうぶ}です。**Attention**、日本語だと注意機構。

シラバスには**Attention層**、**自己注意力 (セルフ・アテンション)**、**Attention Mechanism** と、3
つの形で載っています。まとめて**どこに注目するかを決める仕組み**だと思ってください。

Attention層 Attention Layer

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) Self-Attention

同じ文の中の語**どうして**注目し合う仕組み。自分の中でやるので「自己」。

Attention Mechanism Attention Mechanism

Attentionの仕組みそのものを指す言い方。シラバスでは上の3つが並記されているが、**指しているものは同じ**。

文の中で、どの語とどの語が関係するかを計算する



「それ」 = 「犬」だと結び付ける

線の太さ=どれだけ注目するか（自己注意力）

図2-11 「それ」が指すのはどれか。語と語の関係を重みで計算する

具体的にいきます。「彼は犬を飼っている。それはとても人懐っこい。」……この「それ」は何を指しますか。**犬**ですね。「彼」ではありません。人間は当たり前にできますが、**機械にとっては難問**です。

Attentionは、「それ」を処理するとき、**文の中のどの語を強く見るかを、重み**として計算します。犬に強い重み、彼に弱い重み。これを**同じ文の中の語どうしてやるのが自己注意力**です。自分の中で注目し合うから「自己」です。

試験ポイント

- 試験では**文中のどの単語に注目すべきかを重みで決める仕組み**という言い回しで出ます。
- Attention層・自己注意力・Attention Mechanism は同じものを指します**。3つとも覚えておけば、どの表記で出ても迷いません。

位置エンコーディング

さきほど「Transformerは文の全部の語を一度に見る」と言いました。ここで、^{するど}鋭い人は気づきます。**一度に見たら、順番が分らないですか？**

そのとおりです。一度に見るということは、**語順の情報が消える**ということです。「犬が人を^か噛んだ」と「人が犬を噛んだ」が区別できなくなる。これは困ります。

位置エンコーディング Positional Encoding

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

一度に見ると、順番の情報が消える



これでは意味が壊れる

位置の情報を数値で足す



=位置エンコーディング

図2-12 語に「何番目です」という数値を足しておく。これが位置エンコーディング

そこで**位置エンコーディング**です。**それぞれの語に、あなたは何番目です、という情報を数字として足しておく**。これだけです。でも、これが無いと**Transformerは成立しません**。

試験ポイント

- **Transformerが語順を扱うための仕組み**として出ます。
- **Attentionとセットで覚えてください**。「一度に見る」ことで得たもの（手分け）と失ったもの（語順）が対になります。

文章の系譜を1枚で

ここまでの3つを並べます。**何ができるようになって、何が残った不満か**を横に見てください。順番に出てきた理由がそのまま表になります。

	できるようになったこと	残った不満
RNN	順番のあるデータを扱える	長いと前を忘れる
LSTM	長くても忘れない	1語ずつしか読めない＝学習が遅い
Transformer	全部を一度に見る＝手分けできる	語順が消える（位置エンコーディングで補う）

⇒ この表の右端が、次の行の左端になっています。「前のやつの弱点をつぶした結果」というのは、こういう意味です。

確認問題③ 2-4 のふりかえり

確認問題③

第1問 順番のあるデータを扱えるようにしたモデルは何ですか。またその弱点を直したモデルは。

答

RNN。弱点は長いと忘れること。直したのがLSTM。

RNNの記憶係が隠れ層とリカレント層です。

確認問題③

第2問 Transformerが、順番に読むのをやめたことで得たものは何ですか。

答

手分けができること。

全部の語を一度に見るので、並列に計算できて学習が速くなりました。

確認問題③

第3問 全部の語を一度に見ると失われるものは何で、それをどう補おぎないましたか。

答

語順。位置エンコーディングで補った。

得たものと失ったものは対で問われます。

2-5 Transformerからの枝分かれ

Transformerができました。ここから**枝分かれ**します。片方が**創る**ほう、もう片方が**読む**ほうです。

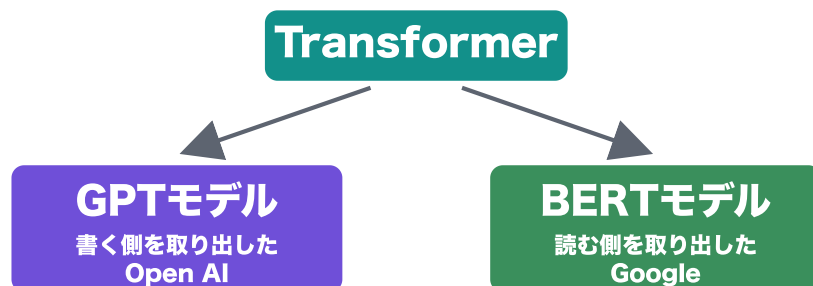


図2-13 Transformerから枝が2本。創るGPTと、読むBERT

GPTモデル (Open AI)

GPTモデル Generative Pre-trained Transformer

Transformerの**作る側**だけ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いていく。

Open AI OpenAI

GPTモデルを開発した組織。

まず一方の枝が**GPTモデル**。作ったのは**Open AI**です。GPTは、**Transformerの作る側だけを取り出した**モデルです。

やっていることは、2-2でやった**自己回帰**です。**次の語を予測して、1つ出す。それも読んで、また次を出す。**これをひたすら繰り返します。

……それだけで文章が書けるのか、と思いますよね。**書けたのです。**それが今のところの答えです。覚え方は**GPTは創るほう**。前から後ろへ、**一方向に書いていきます。**

BERTモデル (MLMとNSP)

BERTモデル Bidirectional Encoder Representations from Transformers

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**ので意味を読み取るのが得意。

もう一方の枝が**BERTモデル**です。こちらはGoogleが出しました。BERTは、**Transformerの読む側**を取り出したモデルです。

GPTが前から後ろへ一方向だったのに対して、**BERTは、前も後ろも同時に見ます。**だから**意味を読み取るのが得意**です。検索や分類に使われます。

そして、BERTの学習方法には名前がついていて、**ここが試験に出ます**。2つあります。

MLM (Masked Language Model) Masked Language Model

文の中の単語をわざと隠して、そこに何が入るかを当てさせる学習方法。穴埋め問題。

NSP (Next Sentence Prediction) Next Sentence Prediction

2つの文が続いているかどうかを当てさせる学習方法。続き当て。

1つ目、**MLM**。マスクド・ランゲージ・モデル。やり方は、**文の中の単語をわざと隠して、そこに何が入るかを当てさせる**。「今日は [] がいい」→「天気」。^{あなう}穴埋め問題を^{えんえん}延々と解かせるわけです。前と後ろの両方を見ないと解けないので、**双方向に読む力がつきます**。

2つ目、**NSP**。ネクスト・センテンス・プリディクション。**2つの文を見せて、これは続いていますか、と当てさせます**。文と文のつながりを学ばせるためのものです。

試験ポイント

- **GPTは創る、BERTは読む**。役割の違いをまず覚えます。
- **BERTの学習は、穴埋めのMLMと、続き当てのNSP**。この2つは入れ替えて出題されます。

RoBERTaとALBERT

最後に、BERTの**改良版が2つ**シラバスに載っています。**RoBERTaとALBERT**です。細かい中身は問われません。**方向性だけ**覚えてください。

RoBERTa Robustly Optimized BERT Approach

BERTを**もっと鍛えた版**。データを増やして長く学習させ、**NSPをやめた**。

ALBERT (a Lite BERT) A Lite BERT

BERTの**軽い版**。中身を工夫して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。

RoBERTa。こちらは**もっと鍛えた版**です。データを増やして、長く学習させた。そして、**さきほどのNSPをやめました**。「続き当ては、そんなに効いていなかった」という判断です。……さっき覚えたものが、すぐ消される。AIの世界はだいたいこうです。

ALBERT。正式には**ア・ライト・バート**。名前のとおり**軽い版**です。中身を工夫して、**パラメータを大幅に減らしました**。性能を保ったまま小さくした、というのが売りです。

試験ポイント

- RoBERTaは強くした版。ALBERTは軽くした版。
- 名前にライトが入っているほうが軽い。これで混ざりません。
- RoBERTaがNSPをやめたことは、そのまま問題になります。

枝分かれを1枚で

	GPTモデル	BERTモデル
役割	創る（文章を書く）	読む（意味を取る）
読む向き	前から後ろへ一方向	前も後ろも同時（双方向）
作ったところ	Open AI	Google
やり方	自己回帰（1語ずつ出す）	MLM（穴埋め）とNSP（続き当て）
改良版	（次の本で扱います）	RoBERTa（強く）／ALBERT（軽く）

△ △ 引っかけ **GPTとBERTの役割を入れ替えた**選択肢が定番です。**創るGPT、読むBERT**。ここだけは反射で言えるようにしてください。

まとめ 第2章① これが言えれば合格ライン

前半を、一本の線としてまとめます。ここが自分の言葉で言えれば、この範囲は^{だいじょうぶ}大丈夫です。

はじめり

- 生成AIは、識別ではなく生成をするAI。ジェネレーティブAIとも言う
- 祖先はボルツマンマシン。ただし全部つながっていて重すぎた
- 層の中のつながりを切ったのが制約付きボルツマンマシン。これで動いた
- 1つずつ順番に出すやり方が自己回帰モデル。GPTがこれ

画像の系譜

- 近所を見る畳み込みでCNN。畳み込みは局所的な特徴を取り出す操作
- 作るほうは2系統。縮めて揺らして戻すVAEと、生成器と識別器を戦わせるGAN
- VAEの部品はエンコーダ・潜在ベクトル・ノイズ・デコーダ

文章の系譜

- 順番を扱えるRNN。記憶係が隠れ層とリカレント層。扱うデータはシーケンスデータ
- 長いと忘れるので、忘却のスイッチをつけたのがLSTM
- そしてTransformerモデル。順番に読むのをやめて、全部を一度に見た
- どこに注目するかを決めるのがAttention層・自己注意力・Attention Mechanism
- 一度に見ると消える語順を補うのが位置エンコーディング。この設計そのものがアーキテクチャ

枝分かれ

- 創るGPTモデル。作ったのはOpen AI
- 読むBERTモデル。学習はMLMとNSP
- その改良が、強くしたRoBERTaと、軽くしたALBERT

最後に

- ……32語、全部通りました。
- アルファベットの列に見えたものが、弱点をつぶし続けた歴史だったと分かれば、もう混ざりません。
- 次は第2章の後半（ChatGPTの系譜）です。GPT-1からGPT-5まで、数字とアルファベットがさらに増えますが、やることは同じです。

入れ替えて出される「対」

この章でいちばん多い引っかけは、**対になっている2語を入れ替える**ものです。ここだけは反射で言えるようにしてください。

対	どちらがどちら
エンコーダ／デコーダ	エンコーダが 入口で縮める 、デコーダが 出口で戻す
生成器／識別器	生成器が 偽物を作る 、識別器が 見分ける
VAE／GAN	VAEは 縮めて戻す 、GANは 2つを競わせる
GPT／BERT	GPTは 創る・一方向 、BERTは 読む・双方向
MLM／NSP	MLMは 穴埋め 、NSPは 続き当て
RoBERTa／ALBERT	RoBERTaは 強くした版 、ALBERTは 軽くした版

32語チェックリスト

最後に、自分の言葉で説明できるか確かめてください。**1つでも詰まったら、その語のページに戻ります。**

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

用語集

第2章①の重要用語 32語

公式シラバスに名前が載っている語だけを並べてあります。p.○ は、その語が本文で説明されているページです。

生成AIとそのはじまり

生成AI（ジェネレーティブAI） p.36

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.38

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.38

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.39

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。

画像の系譜

CNN（畳み込みニューラルネットワーク） p.41

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.41

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.42

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.43

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.42

入力を圧縮する側（入口）。

デコーダ p.42

圧縮されたものから復元する側（出口）。

潜在ベクトル p.42

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.43

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.43

GANの中で偽物を作る側。

識別器 p.43

GANの中で本物か偽物かを見分ける側。

文章の系譜

RNN（回帰型ニューラルネットワーク） p.45

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.45

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.45

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.45

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.46

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.46

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。

Transformerの中身

Attention層 p.47

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.47

同じ文の中の語どうして注目し合う仕組み。

Attention Mechanism p.48

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.48

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.47

モデルの構造・設計そのもの。

枝分かれした先

GPTモデル p.51

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方向に書く。

Open AI p.51

GPTモデルを開発した組織。

BERTモデル p.51

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.52

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.52

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.52

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.52

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。

この本について

生成AIパスポート 対策テキスト 第2章 生成AI ①仕組み編

発行：YouTubeチャンネル「AI資格ラボ」／制作：AI資格ラボ

出題範囲は**生成AIパスポート 公式シラバス (2026年2月試験より適用)** にもとづいています。最新のシラバスは主催団体の公式サイトで確認してください。

⇒ 本書は試験対策のための私家版^{しかばん}テキストです。試験の実施団体とは関係がありません。内容には細心^{さいしん}の注意を払っていますが、**合否を保証するものではありません。**

動画版はこちら：**YouTube「AI資格ラボ」**／第1章のテキストも同じ作りで配布しています。

第2章② 生成AI（ChatGPT編）



この章の解説動画

<https://youtu.be/cssrTuzkY08>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに この本の使い方

この本は、YouTubeチャンネル「AI資格ラボ」の動画「第2章 生成AI ②ChatGPT編」に対応した、配布用のテキストです。動画を見ていない方でも、これ一冊で第2章の後半が最後まで読み切れるように書いてあります。

この本が扱う範囲

生成AIパスポートの第2章は、公式シラバス（出題範囲表）で**62語**あります。第1章の39語の1.6倍で、1回で扱うには多すぎます。そこで**2冊に分けました**。

	扱う範囲	語数
①	生成AIの誕生まで（前の本）	32語
②	ChatGPT（この本）	30語

この本は**②の30語**を扱います。**Transformer**から生まれたGPTが、どう育って、いまのChatGPTになったか。そこが範囲です。

⇒ **前の本でやった言葉は、やり直しません。**Transformer、Attention、自己回帰モデル、隠れ層。このあたりは①仕組み編のテキストと動画にあります。まだの方は、せめて**Transformerの所だけ**読んでから戻ってきてください。

この本の構成

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま問題になる所です。試験直前はここだけ拾い読みできます。
- **△マークの枠**（赤い枠）……**間違えやすい所・引っ掛けが仕込まれる所**です。
- **用語**（紫の帯）……シラバスに名前が載っている用語です。全部で30語あります。
- **確認問題**（緑の枠）……テーマの区切りに3問ずつ、合計9問あります。
- **チェックリスト**（巻末）……30語を自分で説明できるか、最後に確かめられます。

おすすめの進め方

1. まず**本文を最後まで通して読む**。分からない所で止まらず、飛ばして進んでかまいません。
2. 次に**確認問題9問**を、本文を見ずに解く。ここで落ちた所が、あなたの弱点です。
3. 落ちた所だけ本文に戻り、**自分の言葉で言い直す**。読むより言い直したほうが残ります。
4. 試験直前は、**試験ポイントの枠と、巻末の用語集だけ**を見返す。

第2章の後半は、どういう章か

先に、いちばん大事なことを言います。この章には**型番が16個**出てきます。GPT-1、GPT-2、GPT-3、GPT-3.5、GPT-4、GPT-4o、o1、o3、o4、GPT-4.1、GPT-5。それに道具の名前が4つ。

……全部を順番に覚えようとする、まず**続きません**。数字が近すぎて混ざります。

そこでこの本は、**型番ではなく「何ができるようになった時に番号が変わったのか」**で並べます。それだけなら**5つ**です。試験で聞かれるのも、そちらです。

この本・動画・練習問題の回し方

この本は**単体でも読めます**が、動画と練習問題と合わせると回転が速くなります。おすすめは次の4ステップです。

1. **この本（章ごとのPDF）を読む**。無料です。
2. **その章の動画を見る**。本と同じ順番・同じ図で説明しています。
3. **練習問題を解く**。全5章468問と、本番と同じ60問の模擬試験があります。登録も費用もありません。
4. **間違えた問題の解説から、動画のその場面へ飛ぶ**。どこを見直せばいいか探さなくて済みます。

動画と合わせて使う

- この本は、YouTubeチャンネル**AI資格ラボ**の動画 **第2章 生成AI ②ChatGPT編** と同じ順番・同じ図で作っております。
- **動画で流れをつかみ、この本で確かめる。**あるいは**この本を先に読んでから動画を見る。**どちらでも構いません。
- 印刷は**A4・実物大（100%）**で。図の中の文字まで読める大きさに作っております。

AI資格ラボ

はじめに	この本の使い方	58
第2章②	この章の全体像	62
2-6	ChatGPTとは何か	63
	話し言葉で使える／名前の意味	
2-7	GPTの誕生と、大きくする時代	64
	GPT-1と自然言語処理／GPT-2とパラメータ／GPT-3とデータセット	
	確認問題①（2-6・2-7）	66
2-8	人に合わせる	67
	InstructGPTとRLHF／アライメント／GPT-3.5とChatGPT／ファインチューニング／ハルシネーション	
	確認問題②（2-8）	70
2-9	入口が増え、考えるようになった	71
	GPT-4とマルチモーダル／Code InterpreterとGPTs／GPT-4o／推論モデル／GPT-4.1とGPT-5	
	確認問題③（2-9）	74
2-10	Open AIの道具と、他社のAI	75
	Sora・Image Generation・Codex・Operator／Gemini・Claude・Copilot	
まとめ	5つの流れと、覚え方	77
用語集	重要用語30語	79

第2章② この章の全体像

先に地図を配ります。この章は**ひとつの物語**です。「賢くしたら、言うことを聞かなくなった。人に合わせたら、当たった。そのあと、できることが増えていった」。それだけの話です。



どの型番を出されても、この5つのどこかに置ける

図2-14 型番16個は、この5つのどこかに置ける

型番を出されたら、まずこの5つの**どこの話か**を考えてください。それが分かれば、細かい数字を覚えていなくても答えられます。

流れ	起きたこと	代表
①	大きくして賢くする	GPT-1・GPT-2・GPT-3
②	人に合わせる	InstructGPT・GPT-3.5
③	入口を増やす	GPT-4・GPT-4o
④	考えさせる	o1・o3・o4
⑤	ひとつにまとめる	GPT-5

⇒ この表をいま覚える必要はありません。本文を読んだあと、まとめの所で戻ってくると、ちょうど埋まっているはずです。

2-6 ChatGPTとは何か

まず主役そのものから。**ChatGPT**は、**Open AI**が**2022年11月**に公開した対話型のAIです。

ChatGPT Chat GPT

Open AI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。

性能よりも、入口の話

ChatGPTの何がすごかったのか。性能もありますが、試験的に**いちばん大事なここ**です。

専門家でなくても、話し言葉で使えるようになった。

それまでのAIは、使うのに知識が要りました。設定を書き、形式を合わせ、うまくいかなければ調べ直す。ChatGPTは、**日本語で話しかけるだけ**です。

だから、たった**2か月で1億人**が使いました。当時としては史上最速の普及です。

試験ポイント

- **ChatGPT=2022年11月・Open AI**。年月と会社はセットで覚える。
- 画期的だったのは**話し言葉で使えるようにしたこと**。性能そのものより、入口を下げたことが問われる。

名前の意味が、そのまま前半の復習になる

ChatGPTの**Chat**は会話。**GPT**は **Generative Pre-trained Transformer** の頭文字です。

	意味
Generative	生成する
Pre-trained	あらかじめ学習した
Transformer	①仕組み編でやった、あのTransformer

……**全部、前の本でやった言葉**です。名前がそのまま復習になっています。GPTは**Transformerを土台にした、あらかじめ大量に学習させた、生成するモデル**、という意味です。

⇒ ①仕組み編では、この**GPTモデル**まで話が届いていました。この本は**その続き**です。

2-7 GPTの誕生と、大きくする時代

ここからが流れの①**大きくして賢くする**です。GPT-1で「型」ができ、GPT-2とGPT-3で「大きくすれば賢くなる」と分かりました。

GPT-1（2018年）が作ったのは「型」

GPT-1

2018年にOpen AIが発表した最初のGPT。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。

GPT-1でまず押さえる言葉が**自然言語処理**、英語で**NLP**です。

自然言語処理（NLP） Natural Language Processing

私たちが普段しゃべっている言葉（自然言語）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、きっちりした言葉の反対側にある。

自然言語というのは、**私たちが普段しゃべっている言葉**のことです。あいまいで、ゆらぎ、同じことを何通りにも言えます。その扱いにくい言葉をコンピュータに処理させる分野が、NLPです。

では、GPT-1は何をしたのか。



図2-15 GPT-1が作った2段階。いまの生成AIの基本の型になった

先に大量の文章を読ませて、言葉の使い方を覚えさせておく。そのあとで、目的ごとに少し追加で学習させる。

この2段階が、**いまの生成AIの基本の型**になりました。GPT-1自体の性能は、正直たいしたことはありません。**でも型を作った**。覚えるのはそこです。

GPT-2 (2019年) とパラメータ

GPT-2

2019年。パラメータ数を約15億まで増やしたモデル。**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。

GPT-2で出てくるのが**パラメータ**という言葉です。

パラメータ Parameter

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章でやった**重み**が、全部で何個あるか、という数え方。

第1章で**重み**という言葉をやりました。ニューラルネットワークの中にある、学習で調整される数字でしたね。**その数字の総数が、パラメータの数**です。

GPT-1が約1億。**GPT-2は約15億**。15倍です。そして、ここで大事なことが分かりました。

大きくすると、賢くなる。

理屈で工夫したわけではありません。**でかくしたら性能が上がった**のです。身も蓋もありませんが、これがこの10年でいちばん大きな発見でした。

⇒ GPT-2は「悪用されると危ない」という理由で、**段階的にしか公開されませんでした**。偽ニュースが作れてしまう、と。当時はそれくらいの衝撃だった、ということです。

GPT-3 (2020年) とデータセット

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。

GPT-3は2020年。パラメータは**1750億**です。GPT-2の15億から**100倍以上**。

学習に使ったのが、巨大な**データセット**でした。

データセット Dataset

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。

インターネット上の文章、本、Wikipedia。そういうものを大量に読ませました。すると、何が起きたか。

例をいくつか見せるだけで、教えていない仕事ができるようになった。翻訳しろと訓練していないのに翻訳する。要約しろと言っていないのに要約する。

大きくしただけで、勝手に賢くなった。GPT-2で見た法則が、GPT-3で決定的になったわけです。

試験ポイント

- 自然言語処理（NLP）＝人間の言葉をコンピュータに扱わせる技術分野。
- 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」の総数）。
- データセット＝学習に使うデータのまとまり。
- 数字が出るとすれば **GPT-3＝1750億パラメータ**。ここは覚えてよい。

確認問題①（2-6・2-7）

確認問題①

第1問 人間が普段しゃべる言葉を、コンピュータに扱わせる技術分野を何といいますか。

答

自然言語処理（NLP）。

プログラミング言語のような、きっちりした言葉の反対側にある分野です。

確認問題①

第2問 ニューラルネットワークの中で調整される数字の総数を、何といいますか。

答

パラメータ。

第1章でやった「重み」が全部で何個あるか、という数え方です。

確認問題①

第3問 GPT-2からGPT-3にかけて分かった、いちばん大きな法則は何ですか。

答

大きくすると賢くなる。

理屈の工夫ではなく、規模を上げたことで性能が伸びました。

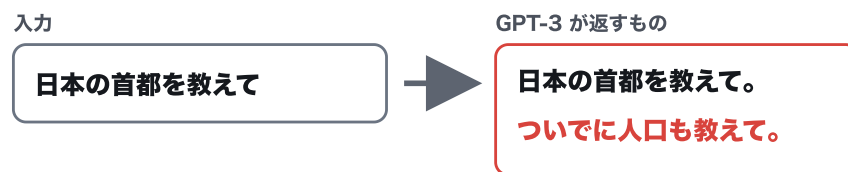
2-8 人に合わせる

流れの②人に合わせるです。ここが、ChatGPTが当たった理由そのものです。**賢いだけでは足りなかった**、という話をします。

賢くなったのに、答えてくれない

GPT-3は賢くなりました。でも、使いにくかったのです。なぜか。

GPT-3は「続きを書く」機械だからです。



質問の続きを書いてくる＝答えない

図2-16 続きを書く機械は、質問に答えてくれない

「日本の首都を教えて」と入力すると、「日本の首都を教えて。ついでに人口も教えて。」と、**質問の続きを書いてくる**。……答えないのです。

InstructGPTとRLHF

これを直したのがInstructGPTです。2022年。**どうやったかが試験に出ます**。

InstructGPT

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



RLHF=人間のフィードバックによる強化学習

図2-17 人が良い方を選び、その選び方をAIに学ばせる＝RLHF

人間に、良い答えと悪い答えを並べて選ばせた。その選び方をAIに学ばせて、**人間が好む答えを出す**ように調整した。この方法の名前がRLHFです。

RLHF (Reinforcement Learning from Human Feedback) 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。

第1章でやった**強化学習**を思い出してください。報酬をもらいながら試行錯誤する学習でしたね。**その報酬を、人間が決めている**。それがRLHFです。

アライメントは「目的」の名前

こうやってAIの振る舞いを、人間の意図や価値観に沿わせることを**アライメント**といいます。

アライメント (Alignment) Alignment

AIの振る舞いを**人間の意図や価値観に沿わせること**。RLHFはそれを実現する手段のひとつ。

⇒ RLHFは手段。アライメントは目的。この関係で覚えてください。**2つを並べて「同じ意味か」と聞く問題**が作れる所です。

GPT-3.5、そしてChatGPTへ

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**（2022年11月）。

GPT-3.5はInstructGPTの流れをくむ改良版です。そして、**これを対話用にして公開したのがChatGPT**。2022年11月のことです。



賢さだけでは当たらなかった

図2-18 賢くなり、人の言うことを聞くようになり、チャットの形で出した

ここは**順番が大事**です。GPT-3が賢くなり、RLHFで人の言うことを聞くようになり、それをチャットの形で出した。**だから当たった**のです。

逆に言うと、**GPT-3のままで当たらなかった**。賢さだけでは足りなくて、**アライメントが要った**ということです。

試験ポイント

- RLHF=人間のフィードバックによる強化学習（手段）／アライメント=人の意図に沿わせること（目的）。
- ChatGPTのベースはGPT-3.5。公開は**2022年11月**。この組み合わせはセットで覚える。

ファインチューニング

ここで**ファインチューニング**という言葉を押さえます。意味は、**学習済みモデルに追加で学習させて、特定の用途に合わせる**ことです。

ファインチューニング Fine-tuning

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。ゼロから作り直すのではなく、できあがったものを少し寄せる。

……聞き覚えがありませんか。**第1章の転移学習**です。学習済みモデルを別の目的に作り替える、あの話。**ファインチューニングは、その調整のやり方**だと思ってください。

たとえば医療の言葉を大量に追加で学習させれば医療に強いモデルになり、自社の問い合わせ履歴を学習させれば自社向けの応答ができます。**ゼロから作らない。できあがったものを、少し寄せる**。試験では「**追加学習で特定用途に合わせる**」という言い回しで出ます。

ハルシネーション

そして、生成AIの話で**絶対に出る言葉**がこれです。**ハルシネーション**。日本語では**幻覚**と訳されます。

ハルシネーション (Hallucination) 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**。仕組み上どうしても起きるので、人間の確認が要る。

意味は、**もっともらしい嘘を、堂々と言う**ことです。存在しない論文を挙げる。実在しない条文を引く。間違った数字を自信満々に出す。

なぜ起きるのか。ここが本質です。

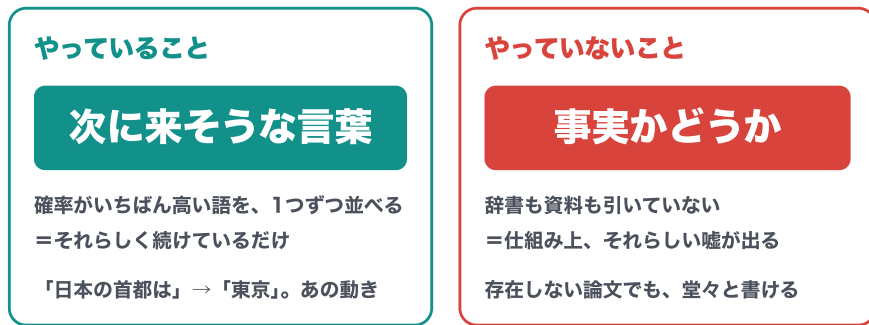


図2-19 やっていることと、やっていないこと

前の本でやったとおり、**GPTは、次に来そうな言葉を選んでいるだけです。**それが本当かどうかは、**そもそも見ていません。**辞書を引いているのではなく、それらしく続けている。だから**それらしい嘘**が出ます。試験では「**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**」という言い方で出ます。

⇒ 実務でもいちばん大事なのがここです。**生成AIの答えは、必ず人間が確かめる。**試験にも仕事にも出てきます。

確認問題② (2-8)

確認問題②

第1問 人間のフィードバックを使った強化学習の略称は？ その目的を表す言葉は？

答

RLHF。目的はアライメント。

RLHFが手段、アライメントが目的、という関係です。

確認問題②

第2問 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を何といいますか。

答

ファインチューニング。

第1章の転移学習の仲間で、その調整のやり方にあたります。

確認問題②

第3問 もっともらしい嘘を生成してしまう現象を何といいますか。なぜ起きるのですか。

答

ハルシネーション。

次に来そうな言葉を選んでいるだけで、事実を確かめていないからです。

2-9 入口が増え、考えるようになった

ここからは流れの③入口を増やす、④考えさせる、⑤ひとつにまとめるです。型番が一気に増えますが、**出てくる順に「何が新しくなったか」だけ**を拾えば足ります。

GPT-4（2023年）とマルチモーダル

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）。

GPT-4で新しくなったのは、大きく2つです。

	新しくなった点
①	ハルシネーションが 減った （無くなってはいない）
②	マルチモーダル になった ← こちらが試験に出る

マルチモーダル Multimodal

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」。

GPT-4 より前



入口は文字だけだった

GPT-4



入口が増えた

図2-20 文字だけだった入口に、画像が加わった

モーダルは「様式」、**マルチ**は「複数」。つまり、**文章だけでなく、画像も扱える**ということです。写真を見せて「これは何？」と聞ける。グラフを見せて説明させられる。

それまでは**文字**だけでした。**入口が増えた**。ここがGPT-4のいちばんの変化です。覚え方は**GPT-4=マルチモーダル**。これだけで足ります。

Code InterpreterとGPTs

GPT-4の時代に、使い方を広げる機能が2つ出しました。

Code Interpreter

ChatGPTが**プログラムを書いて、その場で実行する**機能。表計算ファイルを渡すと、自分で集計してグラフにして返す。

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。

Code Interpreterは、ChatGPTが**プログラムを書いて、その場で実行する**機能です。表計算のファイルを渡すと、自分でコードを書いて、集計して、グラフにして返してきます。文章を書くAIが、計算までやり始めた、ということです。

GPTsは、最後に s が付きます。**目的別のChatGPTを、自分で作れる仕組み**です。「うちの就業規則にだけ答えるやつ」「英会話の練習相手」のようなものを、プログラムを書かずに作れます。

⇒ 試験では**この区別だけ**押さえてください。**実行するのがCode Interpreter、作るのがGPTs。**

GPT-4o（2024年） = o は Omni

GPT-4o

2024年。o は **Omni（すべて）** の頭文字。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。

読み方は「ジーピーティー・フォー・オー」。このオーは**オムニ**の頭文字で、「すべて」という意味です。

GPT-4：中では別々に処理していた

文章

音声

画像

別々に渡していたので、返るまでが遅い

GPT-4o：1つのモデルでまとめて扱う

文章・音声・画像

o = Omni
だから速い

図2-21 別々に処理していたものを、1つのモデルにまとめた

GPT-4は画像も見られるようになりましたが、中では別々に処理していました。GPT-4oは、**文章・音声・画像を、1つのモデルでまとめて扱います**。だから**速い**。音声で話しかけると、人と話しているくらいの速さで返ってきます。

つまり、**GPT-4は「入口が増えた」**。GPT-4oは「同じ入口を、まとめて速くした」。この違いです。

⇒ **GPT-4oの「オー」は、ゼロではありません**。アルファベットのオーです。地味ですが、読み間違える人が本当に多い所です。

推論モデル (o1・o3・o4)

2024年から2025年にかけて、**性質のちがうモデル**が出てきます。シラバスの表記は**GPT-o1・GPT-o3・GPT-o4**です。

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・プログラム・複雑な理屈に強い。

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、返ってくるまでが遅い。

GPT-o4

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

何が違うのか。一言で言えます。**答える前に、考えるようになった**。

それまでのGPTは、思いついた順に書いていく感じでした。このシリーズは、**内側で手順を組み立ててから、答えを出します**。だから**数学・プログラム・複雑な理屈**に強い。そのかわり、**返ってくるのが遅い**。考えているからです。

⇒ **速さのGPT-4o、考えるo1系**。この対比で覚えてください。**oシリーズの「オー」も、ゼロではありません**。しかも**オム二のオーとは別物**です。ここは混ざります。

GPT-4.1とGPT-5 (2025年)

GPT-4.1

2025年。**開発者向け**の色が濃いモデル。長い文章を読ませられる、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**モデル。

GPT-4.1は開発者向けの色が濃いモデルです。一度にたくさんの文章を読ませられるようになった、指示に正確に従うようになった、といった実務寄りの改良です。

そして**GPT-5**。ここまでの流れが、**1つにまとまりました**。速く答える。必要なときは考える。文章も画像も扱う。**別々に育ってきたものが、1つのモデルの中に入った**という位置づけです。

試験ポイント

- **GPT-4=マルチモーダル**（入口が増えた）。
- **GPT-4oのoはOmni**（まとめて速く）。**o1・o3・o4は考えてから答える**（そのぶん遅い）。
- **Code Interpreter=書いて実行／GPTs=自分専用を作る**。
- 型番の細かい違いは追わなくてよい。**5つの流れのどこかが分かればよい**。

確認問題③ (2-9)

確認問題③

第1問 GPT-4で新しくなった、いちばん大きな点は何ですか。

答

マルチモーダルになったこと。

文章だけでなく、画像も扱えるようになりました。

確認問題③

第2問 GPT-4oの「オー」は何の頭文字ですか。

答

オムニ (Omni)。

すべてを1つのモデルでまとめて扱い、応答が速くなりました。ゼロではありません。

確認問題③

第3問 o1・o3・o4のシリーズが、それまでのモデルと違う点は何ですか。

答

答える前に考える（推論する）こと。

そのぶん、返ってくるのは遅くなります。

2-10 Open AIの道具と、他社のAI

最後は名前を覚えるだけの所です。まとめて「何を作るものか」だけ押さえれば足ります。

Open AIの、目的を絞った道具4つ

Open AIは、ChatGPT以外にも目的を絞った道具を出しています。シラバスに名前が載っているのは4つです。

名前	何を作るもの
Sora	動画を作る
Image Generation	画像を作る
Codex	プログラムを書く
Operator	ブラウザをAIが自分で操作する

Sora

Open AIの動画生成AI。文章で指示すると映像が出てくる。

Image Generation

Open AIの画像生成の機能。文章で指示して絵を作る。

Codex

Open AI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。

Operator

AIがブラウザを自分で操作するもの。検索し、ページを開き、フォームに入力して用事を代わりに済ませる。

このうちOperatorだけ毛色が違います。ブラウザを、AIが自分で操作します。検索して、ページを開いて、フォームに入力して、用事を代わりにやる。

⇒ 答えるAIから、動くAIへ。このOperatorの発想は、第3章の「AIエージェント」につながります。この本ではここまでです。

ChatGPT以外の主役3つ

シラバスは、**Gemini・Claude・Copilot**の3つに項目をまるまる1つ割いています。**会社名とセットで覚える**のがいちばん速いです。

名前	会社	強み
Gemini	Google	最初から文章・画像・音声・動画をまとめて扱う設計
Claude	Anthropic	長い文章を読むのが得意。安全性と受け答えの丁寧さ
Copilot	Microsoft	Word・Excel・Teamsの中でそのまま使える

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。Google検索やGmailと組み合わせられる。

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。名前のとおり「副操縦士」。Word・Excel・Teamsの中に組み込まれている。

試験ポイント

- 道具4つ＝**Sora（動画）／Image Generation（画像）／Codex（プログラム）／Operator（ブラウザ操作）**。
- 他社3つ＝**Gemini（Google）／Claude（Anthropic）／Copilot（Microsoft）**。会社名とセットで。

まとめ 5つの流れと、覚え方

第2章の後半を、**流れとして**まとめます。型番16個が、5つになります。

① 型ができて、大きくなった

GPT-1で型ができた。先に大量に読ませて、あとで用途に合わせる。分野の名前は**自然言語処理 (NLP)**。

GPT-2・GPT-3で「大きくすると賢くなる」と分かった。数字の総数が**パラメータ**、読ませたものが**データセット**。GPT-3で**1750億**。

② 人に合わせて、ChatGPTになった

賢くなったのに言うことを聞かない。それを直したのが**InstructGPT**。手段が**RLHF**、目的が**アライメント**。

その改良版**GPT-3.5**を対話用にして出したのが、**2022年11月のChatGPT**。用途に寄せる追加学習が**ファインチューニング**。仕組み上どうしても出るのが**ハルシネーション**で、**必ず人が確かめる**。

③④⑤ 入口が増え、考え、ひとつになった

GPT-4でマルチモーダル。**Code Interpreter**で実行、**GPTs**で自分専用。**GPT-4o**のオーは**オムニ**でまとめて速く。**o1・o3・o4**は考えてから答える。**GPT-4.1**は開発者向け、**GPT-5**でひとつにまとまった。

道具は4つ、他社は3つ。**Sora = 動画**、**Image Generation = 画像**、**Codex = プログラム**、**Operator = ブラウザ操作**。**Gemini = Google**、**Claude = Anthropic**、**Copilot = Microsoft**。



どの型番を出されても、この5つのどこかに置ける

図2-22 型番16個が、5つの流れになった

⇒ 試験で見たことのない型番が出て、**この5つのどこの話か**を考えれば位置が分かります。それで足りません。

チェックリスト

最後に、30語を**自分の言葉で説明できるか**確かめてください。言えなかったものだけ、本文に戻ります。

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

用語集

重要用語30語

ChatGPT p.63

Open AIが2022年11月に公開した対話型の生成AI。話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.64

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.64

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。

InstructGPT p.67

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.68

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.71

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）。

ファインチューニング p.69

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.69

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.71

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。

GPT-o1 p.73

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.73

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.73

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-2 p.65

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.65

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.65

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。

データセット p.65

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.68

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.68

AIの振る舞いを人間の意図や価値観に沿わせること。RLHFはそれを実現する手段のひとつ。

Code Interpreter p.72

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.72

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.72

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-4.1 p.73

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.73

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.75

Open AIの動画生成AI。文章で指示すると映像が出てくる。

Operator p.75

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.75

Open AI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.75

Open AIの画像生成の機能。文章で指示して絵を作る。

Gemini p.76

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 p.76

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 p.76

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。

AI資格ラボ

第3章 現在の生成AIの動向



この章の解説動画

<https://youtu.be/NxbgXzMgCrM>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに この本の使い方

この本は、YouTubeチャンネル「AI資格ラボ」の動画「**第3章 現在の生成AIの動向**」に対応した、配布用のテキストです。動画を見ていない方でも、これ一冊で第3章が最後まで読み切れるように書いてあります。

この本が扱う範囲

生成AIパスポートの第3章は、公式シラバス（出題範囲表）で**20語**です。第1章の39語、第2章の62語にくらべると**いちばん少ない章**で、1冊で足ります。

	テーマ	語数
1	生成AIが出来ることと主なサービス	10語
2	ディープフェイク（深層偽造）技術	2語
3	RAG	3語
4	AIエージェント	5語

⇒ ただし**新しい言葉が多い**のがこの章です。RAG、チャック、ベクトルデータベース、AIエージェント、MCP、GenSpark、Manus、Skywork AI、Veo3。**聞いたことのない名前が並びます**。そこだけ覚悟してください。

この章は、何の話か

第1章と第2章は、**AIの中身**の話でした。ニューロンがあつて、重みがあつて、Transformerがあつて。正直、しんどかったと思います。

第3章は変わります。中身の話は終わりです。ここからは、**それで何ができるのか。何が危ないのか**。そして、**AIが自分で動き出す**という話です。

前の章とのつながり

この章は、第2章を見た方に**効く仕掛け**が2つあります。単独でも読めますが、知っていると刺さります。

- **RAG**は、第2章で出てきた**ハルシネーション**（もっともらしい嘘）への**答え**です。
- **AIエージェント**は、第2章の最後に出てきた**Operator**の**正体**です。

この本の構成

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま問題になる所です。試験直前はここだけ拾い読みできます。
- **△マークの枠**（赤い枠）……**間違えやすい所・引っ掛けが仕込まれる所**です。
- **用語**（紫の帯）……シラバスに名前が載っている用語です。全部で20語あります。
- **確認問題**（緑の枠）……テーマの区切りに3問ずつ、合計9問あります。
- **チェックリスト**（巻末）……20語を自分で説明できるか、最後に確かめられます。

おすすめの進め方

1. まず**本文を最後まで通して読む**。分からない所で止まらず、飛ばして進んでかまいません。
2. 次に**確認問題9問**を、本文を見ずに解く。ここで落ちた所が、あなたの弱点です。
3. 落ちた所だけ本文に戻り、**自分の言葉で言い直す**。読むより言い直したほうが残ります。
4. 試験直前は、**試験ポイントの枠と、巻末の用語集だけ**を見返す。

この本・動画・練習問題の回し方

この本は**単体でも読めます**が、動画と練習問題と合わせると回転が速くなります。おすすめは次の4ステップです。

1. **この本（章ごとのPDF）を読む。** 無料です。
2. **その章の動画を見る。** 本と同じ順番・同じ図で説明しています。
3. **練習問題を解く。** 全5章468問と、本番と同じ60問の模擬試験があります。登録も費用もありません。
4. **間違えた問題の解説から、動画のその場面へ飛ぶ。** どこを見直せばいいか探さなくて済みます。

動画と合わせて使う

- この本は、YouTubeチャンネル**AI資格ラボ**の動画 **第3章 現在の生成AIの動向** と同じ順番・同じ図で作ってあります。
- **動画で流れをつかみ、この本で確かめる。** あるいは**この本を先に読んでから動画を見る。** どちらでも構いません。
- 印刷は**A4・実物大（100%）** で。図の中の文字まで読める大きさに作ってあります。

はじめに	この本の使い方	81
第3章	この章の全体像	85
3-1	学習の前に、データを整える	86
	画像のリサイズ／正規化	
3-2	データを増やす、質を上げる	87
	データの水増し（augmentation）／データ拡張技術／リマスタリング	
	確認問題①（3-1・3-2）	88
3-3	主なサービス	89
	Claude／Gemini／Sora／Veo3／自己回帰モデル	
3-4	ディープフェイクと偽情報	90
	ディープフェイク（深層偽造）技術／偽情報（ディスインフォメーション）	
	確認問題②（3-3・3-4）	91
3-5	RAG 答える前に、資料を探す	92
	RAGとは／何が嬉しいか／ファインチューニングとの違い	
3-6	チャンクとベクトルデータベース	94
	資料を切る／意味を数字にする	
3-7	AIエージェントとMCP	96
	AIエージェント／4つの繰り返し／GenSpark・Manus・Skywork AI／MCP	
	確認問題③（3-5～3-7）	98
まとめ	20語を、流れで覚える	99
用語集	重要用語20語	101

第3章 この章の全体像

先に地図を配ります。この章は**4つのテーマ**でできています。前半は「作る前の準備」と「危ない使い方」、後半が本命の**RAG**と**AIエージェント**です。

	何の話か	この本での場所
1	学習させる前の下ごしらえと、主なサービス	3-1・3-2・3-3
2	ディープフェイクと偽情報（危ない使い方）	3-4
3	RAG（答える前に資料を探す）	3-5・3-6
4	AIエージェント（AIが自分で動く）	3-7

ここが、この章いちばんの引っかけ

テーマ1は、公式シラバスに「学習項目」として**テキスト生成AI／画像生成AI／音楽生成AI／音声生成AI／動画生成AI**の5つが並んでいます。5つ並んでいると、**それぞれの仕組みを覚えるのかな**と思いますよね。

……違うんです。

→ 同じ項目の「キーワード」に並んでいるのは、**5種類の名前ではありません**。画像のリサイズ、正規化、データの水増し、データ拡張技術、リマスタリング。**全部、学習させる「前」のデータの整え方**です。試験に出るのは**そちら**なので、この本もそこから始めます。

覚え方の芯

細かい語は後から入ります。まず、この章全体を貫く一文だけ持って帰ってください。

第3章は、AIが「答える」から「動く」に変わった章。それだけです。

3-1 学習の前に、データを整える

AIに学習させるとき、**データはそのままでは使えません**。たとえば画像を1万枚集めたとします。大きさもバラバラ、明るさもバラバラ。このまま食べさせると、AIは混乱します。

そこで、**2つの下ごしらえ**をします。

画像のリサイズ Resize

画像の**大きさ（縦横のピクセル数）を揃えること**。たとえば全部を256×256にする。AIは決まった大きさの入力しか受け取れないため、揃えないとそもそも入らない。

正規化 Normalization

データの**値の範囲を揃えること**。画素の明るさは0～255の数字だが、これを0～1に直すのが代表例。数字の桁が大きいままだと学習が不安定になるため。



図3-1 リサイズは大きさを揃える。正規化は値を揃える

なぜ、それぞれ必要なのか

リサイズが要る理由は単純です。ニューラルネットワークは、入口の数があらかじめ決まっています。大きさがバラバラだと**そもそも入らない**のです。

正規化が要る理由は、学習の安定です。0～255のまま計算すると、数字が大きすぎて**学習が暴れます**。桁を揃えてやると、素直に学習してくれる。

試験ポイント | この2語は対で聞かれる

- **リサイズは「大きさ」を揃える。**
- **正規化は「値」を揃える。**
- どちらも**学習させる前**の処理。生成AIそのものの仕組みではない。

⇒ ⚠ **「正規化」は日常語と紛らわしい語です。**ここでは**値の範囲を0～1などに揃えること**を指します。データベースの正規化（表の設計を整えること）とは別物です。

3-2 データを増やす、質を上げる

下ごしらえができました。でも、まだ問題があります。**データが足りない**のです。

AIの学習には、とにかく大量のデータが要ります。でも、1万枚集めるのだって大変です。10万枚なんて、まず無理です。

データの水増し (augmentation) Data Augmentation

手持ちのデータを**加工して、枚数を増やすこと**。オーグメンテーションと読む。1枚の写真を左右反転・回転・明るさ変更・切り取りすることで、10枚ぶんの学習データにできる。

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方 (総称)**。水増しが「やること」、データ拡張技術が「その技術の呼び名」。

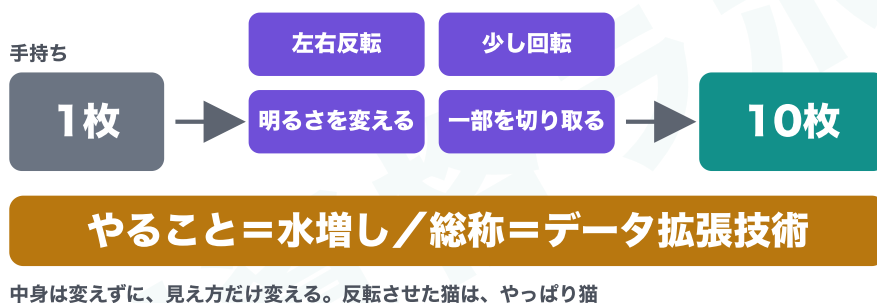


図3-2 1枚を加工して枚数を増やす。総称がデータ拡張技術

これは「水増し」であって「捏造」ではない

大事なのは、**中身は変えずに、見え方だけ変える**という点です。反転させた猫は、やっぱり猫です。データを**でっち上げているわけではありません**。

⇒ ⚠️ 「水増し」という日本語に引きずられないこと。不正な水増しの意味ではありません。同じ中身を、違う見え方で何枚も用意するという、正当な学習の工夫です。

もうひとつ、リマスタリング

リマスタリング Remastering

古い・粗いデータを、**作り直して質を上げること**。もとは音楽や映像の言葉で、古い録音を今の技術できれいにすること。生成AIでは解像度を上げる、ノイズを取る、色を戻す、といった処理を指す。

昔のレコードが「リマスター版」としてCDになる、あれです。生成AIの文脈では、**古い・粗いデータをAIで作り直して質を上げることを指します**。



図3-3 データの水増しは「数」、リマスタリングは「質」

試験ポイント | 混ざりやすい2語

- データの水増し=数を増やす。1枚を10枚に。
- リマスタリング=質を上げる。粗いものをきれいに。
- 迷ったら「数か、質か」で分ける。それだけで切り分けられる。

確認問題① (3-1・3-2)

確認問題①

第1問 画像の大きさを揃えることと、値の範囲を揃えること。それぞれ何とといいますか。

答

画像のリサイズと正規化。

リサイズは「大きさ」、正規化は「値」。どちらも学習させる前の下ごしらえです。

確認問題①

第2問 手持ちのデータを加工して枚数を増やすことを何とといいますか。その技術の総称は何とといいますか。

答

データの水増し (augmentation)。総称がデータ拡張技術。

水増しが「やること」、データ拡張技術が「その技術の呼び名」。中身は変えずに見え方だけ変えます。

確認問題①

第3問 古い・粗いデータをAIで作直して質を上げることを何とといいますか。

答

リマスタリング。

水増しは「数」、リマスタリングは「質」。この対比で覚えると混ざりません。

3-3 主なサービス

公式シラバスは、この項目で**サービスの名前**も挙げています。といっても、ほとんどは第2章で出てきた名前です。ここではやり直さず、**並べて確かめる**だけにします。

Claude

Anthropic（アンソロピック）の生成AI。長い文章を読ませるのが得意で、受け答えの丁寧さが評価されている。

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になっている。

Sora

Open AIの動画生成AI。文章で指示すると映像が出てくる。

Veo3

Googleの動画生成AI。ヴィオ・スリーと読む。**SoraのGoogle版**と思えばよい。この章で新しく出てくる語。

試験ポイント | 対で覚える

- **Sora=Open AI。Veو3=Google。**どちらも**動画**生成。
- **Claude=Anthropic。Gemini=Google。**
- 会社名と結び付けられれば十分。細かい性能差は問われない。

自己回帰モデルが、なぜここに載っているのか

自己回帰モデル Autoregressive Model

自分が出した答えを**次の入力**にして、**1つずつ生成**していくやり方。文字が1つずつ出てくる、あの動きの正体。

第2章でやった語が、第3章のこの項目にも載っています。なぜか。**生成AIが「作れる」理由そのもの**からです。

次に来そうな語を1つ出す。それを入力に足して、また次を出す。これを繰り返すから、文章が伸びていきます。**サービスの名前と並べて「どうやって作っているのか」を確かめる**位置づけだと考えてください。

3-4 ディープフェイクと偽情報

ここから**危ない話**です。技術そのものより、**危険性**のほうが問われます。

ディープフェイク（深層偽造）技術 Deepfake

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。実在する人の顔を別の映像に貼り替えたり、実在する人の声で言っていないことを喋らせたりできる。

何ができてしまうのか。実在する人の顔を、別の映像に貼り替える。実在する人の声で、言っていないことを喋らせる。しかも、**見分けが付きません**。

なぜ、見分けがつかないのか

第2章のGANを思い出してください。偽物を作る側と、それを見破る側を競わせる仕組みでした。見破られなくなるまで鍛えるわけですから……そりゃあ、見分けがつかなくなります。

実際に事件も起きています。経営者の声をまねた電話で**送金させる**。政治家が言っていないことを**言ったこと**にする。

⇒ ⚠ **技術そのものは中立です**。映画の吹き替えや、故人の声の再現にも使われます。試験で問われるのは**使い方の危険性**のほうです。「ディープフェイクは悪いものか」ではなく「何ができてしまうか」で考えてください。

セットで出る、偽情報

偽情報（ディスインフォメーション） Disinformation

意図的に流される、嘘の情報。うっかり間違えて広がったもの（誤情報／ミスインフォメーション）とは区別される。分かれ目は**悪意があるかどうか**。

誤情報

うっかり間違えて広がったもの
悪意なし

偽情報

わざと人をだますために流すもの
＝ディスインフォメーション・悪意あり

分かれ目は「悪意があるかどうか」

図3-4 分かれ目は「悪意があるかどうか」

ここは「**意図的に**」がポイントです。うっかり間違えて広がったものは**誤情報**。わざと人をだますために流すものが**偽情報＝ディスインフォメーション**。

なぜ、生成AIの章に出てくるのか

生成AIのせいで、偽情報を作るコストがほぼゼロになったからです。昔は、嘘の記事を100本書くのは大変でした。いまは一瞬です。しかも、それらしい。

試験ポイント | ハルシネーションとの違い

- ハルシネーションは、AIが勝手に間違える。悪意はない。仕組み上そうなる（第2章）。
- 偽情報は、人間がわざとやる。悪意がある。
- 誤情報は、人間がうっかり間違える。悪意はない。
- 誰が／わざとか、の2点で3つが切り分けられる。

確認問題② (3-3・3-4)

確認問題②

第1問 動画生成AIのうち、Open AIのものとGoogleのものを、それぞれ何といいますか。

答

Open AIが**Sora**、Googleが**Veo3**。

対で覚えるのが早いです。どちらも動画生成。Veo3はこの章で新しく出る語です。

確認問題②

第2問 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を何といいますか。

答

ディープフェイク（深層偽造）技術。

第2章のGAN（作る側と見破る側を competing させる仕組み）が下地です。試験では危険性が問われます。

確認問題②

第3問 意図的に流される嘘の情報を何といいますか。うっかりの間違いとの分かれ目は何ですか。

答

偽情報（ディスインフォメーション）。分かれ目は**悪意の有無**。

うっかりは誤情報。AIが勝手に間違えるのはハルシネーションで、こちらは人間がわざとやる話です。

3-5 RAG | 答える前に、資料を探す

ここからが、この章の**本命**です。実務でいちばん耳にする語でもあります。

RAG (Retrieval-Augmented Generation) Retrieval-Augmented Generation

ラグと読む。日本語では**検索拡張生成**。答える前に**資料を検索し、見つけた資料を読ませてから答えさせる**仕組み。ハルシネーションへの代表的な対策。

いつ生まれたのか

RAGという言葉は、**2020年**にMeta（当時のFacebook AI Research）の研究チームが出した論文で提案されました。**検索してきた文章を手がかりにして文を生成する**という枠組みで、公開当時から質問応答の成績を伸ばすことが示されています。

広く使われるようになったのは**大規模言語モデルが普及してから**です。**学習した時点までしか知らない**という限界と、**もっともらしい嘘をつく**という弱点。この2つに現場で効く答えがRAGだったため、一気に広まりました。

⇒ ⚠ 年号そのものが問われることは少ないと思われますが、**RAGはLLMより先にあった考え方で、LLMの弱点に効いたから広まった**——この流れは押さえておいてください。

前の章の話から入ります

第2章でハルシネーションをやりました。AIがもっともらしい嘘をつく。なぜか。**次に来そうな言葉を選んでだけで、それが本当かどうかは見えていない**からでした。辞書を引いていないのです。

……では、**引かせればいいんじゃないか**。答える前に、資料を検索させる。見つけた資料を読ませてから、答えさせる。**それがRAGです。**

順番でいうと、こうです。①**質問が来る**。②**先に資料を探す**。③**見つけた資料を、質問と一緒にAIへ渡す**。④**AIはそれを読んでから答える**。



この「探してから答える」がRAG

図3-5 この「探してから答える」がRAG

何が嬉しいのか | 3つ

	嬉しいこと	なぜ
①	嘘が減る	手元に資料があるので、でっち上げなくて済む
②	新しい情報に答えられる	AIは学習した時点までしか知らないが、資料は今日のものでいい
③	社内の情報に答えられる	インターネットに無い、自社の規則やマニュアルでも渡せる

そして、いちばん効くのが**モデルを作り直さなくていい**という点です。

ファインチューニングとの違い

第2章のファインチューニングは、**モデルそのものを追加学習させる**話でした。RAGは違います。**モデルは触らず、渡す資料を変えるだけ**です。

前回：ファインチューニング

モデルを作り直す

追加学習させる＝重みが変わる
高い・時間がかかる

今回：RAG

資料を渡すだけ

モデルそのものは触らない
安い・速い・すぐ差し替えられる

だから実務でいちばん使われる

図3-6 モデルを触るか、資料を変えるか

試験ポイント | RAGの言い回し

- **RAG＝検索拡張生成**。答える前に資料を探し、読ませてから答えさせる。
- **ハルシネーション対策**として出題されやすい。
- **モデルは触らない**。ここがファインチューニングとの決定的な違い。
- だから**安くて速い**。実務でいちばん使われている理由もそこ。

3-6 チャンクとベクトルデータベース

では、どうやって「探す」のか。ここに**2つの語**が出ます。RAGとセットで問われるので、必ず押さえてください。

チャンク Chunk

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、という意味。1冊の本を丸ごと渡すことはできないので、先に切っておく。

資料をそのまま渡すことはできません。**1冊の本を丸ごと渡されても困る**のと同じです。だから、あらかじめ細かく切っておきます。この切った一つひとつが**チャンク**です。

ベクトルデータベース Vector Database

チャンクを**意味を表す数字の並び（ベクトル）に変換して、しまっておく保管庫**。数字にしておくと、意味が近いかどうかを計算で比べられる。

なぜ、数字にするのか

言葉のままだと、**意味が近いかどうか比べられない**からです。

たとえば「犬」と「わんこ」。**文字は全然違うのに、意味は近い**。文字を見比べているかぎり、この2つは別物です。ところが**数字にしておくと、その「近さ」が計算できる**。

その数字をしまっておく保管庫が、**ベクトルデータベース**です。質問が来たら、**質問も数字にして、近いチャンクを引っばってきます**。

① 資料を細かく切っておく



② 意味を数字にして、しまっておく



質問も数字にして、近いものを探す

図3-7 切って、数字にしてしまう

試験ポイント | RAGの中身は3段

- ①資料をチャンクに切る。
- ②意味を数字にして、ベクトルデータベースへ入れる。
- ③質問が来たら近いものを探して、AIに読ませる。
- RAG・チャンク・ベクトルデータベースは3点セットで出る。

よくある取り違い

- ⇒ ⚠ ベクトルデータベースは「AIそのもの」ではありません。資料を数字にしてしまっておく**保管庫**です。考えるのはあくまでAIの側で、こちらは引き出しの役目。
- ⇒ ⚠ RAGを使っても、AIが賢くなるわけではありません。モデルは何も変わりません。**渡す資料が増えるだけ**です。だから、資料に無いことは相変わらず答えられません。

AI資格ラボ

3-7 AIエージェントとMCP

お待たせしました。第2章の最後に出てきた**Operator**の、正体です。

AIエージェント AI Agent

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。聞かれたことに答えるだけの従来のAIと、そこが違う。

これまでのAIと、何が違うのか

これまでのAIは、聞かれたことに答えるだけでした。「メールの文面を書いて」と言えば、文面が返ってくる。……**送るのは、あなた**です。

AIエージェントは違います。「取引先にお礼のメールを送っておいて」と言うと、**誰宛か調べて、文面を書いて、送信ボタンまで押します**。



図3-8 答えるAIと、動くAI

第2章の**Operator**が、まさにこれでした。ブラウザを開いて、検索して、フォームに入力する。**答えるAIから、動くAIへ**。その「動くAI」の呼び名が、**AIエージェント**です。

仕組みは、4つの繰り返し

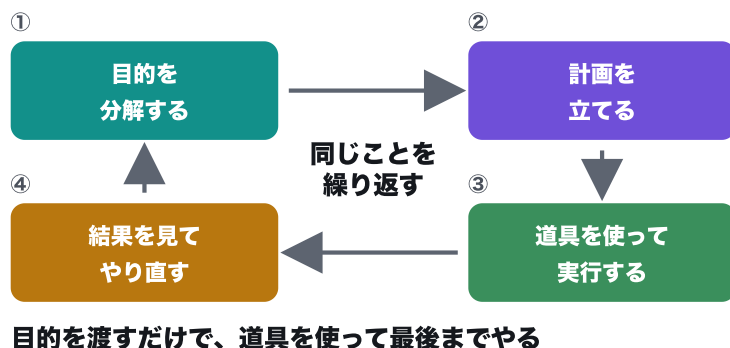


図3-9 目的を分解し、計画を立て、道具を使い、結果を見てやり直す

ざっくり言うと、**目的を分解する→計画を立てる→道具を使って実行する→結果を見て、やり直す**。これを繰り返します。……**人間の仕事の進め方と同じ**です。そこが新しい。

ツール事例は3つ

公式シラバスには、エージェントの実例が3つ載っています。**名前だけ押さえてください**。どれも新しいサービスなので、細かい機能は問われません。

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**。聞かれたことに対して、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**。指示すると、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。こちらも汎用型で、**資料やスライドを作るのが得意**とされている。

⇒ ⚠ **3つとも新しいサービスです**。できることは変わっていきます。試験では**名前と「エージェント系」だけ結び付けば十分**。機能の細部を追いかける必要はありません。

最後に、MCP

MCP Model Context Protocol

モデル・コンテキスト・プロトコル。**Anthropic**が出した決まりごとで、**AIが外の道具やデータにつながるための共通の差し込み口**。AIエージェントとセットで出る。

たとえ話をします。昔、**充電ケーブルが機種ごとに違って、全部バラバラでしたよね**。メーカーが違くと、挿さらない。それが、**規格が統一されて、どれでも挿さるようになった**。

MCPは、そのAI版です。MCPに合わせておけば、ChatGPTだろうがClaudeだろうがGeminiだろうが、同じやり方で、同じ道具やデータにつながります。



充電ケーブルの規格が統一されて、どれでも挿さるようになったのと同じ

図3-10 AIが外の道具につながるための共通の差し込み口

試験ポイント | エージェントとMCP

- AIエージェント＝目的を渡すと、自分で考えて、道具を使って、最後までやるAI。
- 実例はGenSpark・Manus・Skywork AIの3つ。名前だけでよい。
- MCP＝外の道具やデータにつなぐための共通の決まりごと（Anthropic）。
- エージェントが「手足を増やす」ための規格と考えると、セットで思い出せる。

— 確認問題③ (3-5～3-7)

確認問題③

第1問 答える前に資料を検索して、それを読ませてから答えさせる仕組みを何とといいますか。日本語では何とといいますか。

答

RAG (Retrieval-Augmented Generation)。日本語では**検索拡張生成**。

ハルシネーション対策として出ます。ファインチューニングと違い、モデルそのものは触りません。

確認問題③

第2問 RAGで、資料を細かく切ったかたまりを何とといいますか。それをしまっておく保管庫は何とといいますか。

答

チャンク。保管庫は**ベクトルデータベース**。

意味を数字にしてしまうので、「犬」と「わんこ」のように文字が違ってても近さを計算できます。

確認問題③

第3問 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を何とといいますか。どこが出したのですか。

答

MCP (Model Context Protocol)。**Anthropic**が出した。

充電ケーブルの規格統一のAI版。AIエージェントとセットで問われます。

まとめ 20語を、流れで覚える

最後に、20語を**流れ**で並べ直します。バラバラに覚えるより、この順で1本につないだほうが残ります。

① 学習させる前の下ごしらえ（5語）

語	何をする
画像のリサイズ	大きさを揃える
正規化	値を揃える
データの水増し	数を増やす (augmentation)
データ拡張技術	水増しの技術の総称
リマスタリング	質を上げる

⇒ 「数か、質か」で水増しとリマスタリングを分ける。ここは**毎回まざる**所です。

② 主なサービス（5語）

Claude=Anthropic。Gemini=Google。Sora=Open AIの動画。Veo3=Googleの動画。そして**自己回帰モデル**=自分が出した答えを次の入力にして1つつ生成していくやり方。

③ 危ない話（2語）

ディープフェイク（深層偽造）技術=本物そっくりの偽物を作る技術。偽情報（ディスインフォメーション）=わざと流す嘘。**AIが勝手に間違えるのがハルシネーション、人間がわざとやるのが偽情報。**

④ RAG（3語）

RAG=答える前に資料を探して、読ませてから答えさせる。ハルシネーションへの答え。**チャンク**=資料を切ったもの。**ベクトルデータベース**=意味を数字にしてしまう保管庫。

⑤ AIエージェント（5語）

AIエージェント=目的を渡すと、自分で考えて、道具を使って、最後までやる。第2章の**Operator**の正体。実例が**GenSpark・Manus・Skywork AI**。外の道具につなぐ共通の決まりごとが**MCP**。

この章から持って帰る一文

- **第3章は、AIが「答える」から「動く」に変わった章。**
- 前半（下ごしらえ）は**作る前**の話、後半（RAG・エージェント）は**使うとき**の話。
- どちらも**生成AIそのものの仕組み**ではなく、**その周りで起きていること**を問う章です。

次の章は**第4章「情報リテラシー・基本理念とAI社会原則」**です。いちばん語数が多く、技術の話がほとんど出てきません。法律、権利、ルールの章です。ただ、今日やった**ディープフェイクと偽情報**が、そのまま次につながります。**危ないと分かったから、ルールを作った**。そういう流れです。

AI資格ラボ

用語集

重要用語20語

公式シラバス第3章の詳細キーワード20語です。説明できるかを確かめてください。読めるだけでは足りません。

画像のリサイズ p.86

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.86

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.87

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。

Claude p.76

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

Gemini p.76

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。

Sora p.75

Open AIの動画生成AI。文章で指示すると映像が出てくる。

ディープフェイク（深層偽造）技術 p.90

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.90

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。

データ拡張技術 p.87

データの水増しをするための手口をまとめた呼び方（総称）。水増しが「やること」、こちらが「その技術の呼び名」。

リマスタリング p.87

古い・粗いデータを作り直して質を上げること。解像度を上げる、ノイズを取る、色を戻すなど。水増しが「数」、こちらは「質」。

Veo3 p.89

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。

自己回帰モデル p.39

自分が出した答えを次の入力にして、1つずつ生成していくやり方。文字が1つずつ出てくる動きの正体。

RAG (Retrieval-Augmented Generation) p.92

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.94

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.94

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.96

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.97

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.97

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.97

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.97

Model Context Protocol。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。

AI資格ラボ

自分の言葉で説明できたら✓を入れてください。読めるだけでは足りません。

- ☐ 画像のリサイズ
- ☐ 正規化
- ☐ データの水増し (augmentation)
- ☐ データ拡張技術
- ☐ リマスタリング
- ☐ Claude
- ☐ Gemini
- ☐ Sora
- ☐ 自己回帰モデル
- ☐ Veo3
- ☐ ディープフェイク (深層偽造) 技術
- ☐ 偽情報 (ディスインフォメーション)
- ☐ RAG (Retrieval-Augmented Generation)
- ☐ チャンク
- ☐ ベクトルデータベース
- ☐ AIEージェント
- ☐ GenSpark
- ☐ Manus
- ☐ Skywork AI
- ☐ MCP

第4章① 情報リテラシー（身を守る）



この章の解説動画

<https://youtu.be/w2xB5OFMuwM>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに この本の使い方

この本は、YouTubeチャンネル「AI資格ラボ」の動画「第4章 情報リテラシー ①身を守る」に対応した、配布用のテキストです。動画を見ていない方でも、これ一冊で第4章の前半が最後まで読み切れるように書いてあります。

この本が扱う範囲

生成AIパスポートの第4章は、公式シラバス（出題範囲表）で**65語**あります。第1章の39語、第2章の62語、第3章の20語とくらべて**全5章でいちばん多い章**です。1回で扱うには多すぎるので、第2章と同じように**2冊に分けました**。

	扱う範囲	語数
①	自分の身を守る（この本）	25語
②	作ったものの権利と社会のルール（次の本）	40語

この本は**①の25語**を扱います。ネットで何が起きるのか、どうだまされるのか、個人情報はどう守られているのか。**全部、自分の身を守る話**です。

→ この章は、技術の話がほとんど出てきません。第1章から第3章までは、ニューロンや Transformer や RAG といった**中身**の話でした。第4章で出てくるのは**法律と、手口の名前と、ルール**です。頭の使い方が変わります。

この章の、いちばんしんどい所

先に正直に言います。この本の山は**だます手口の名前**です。フィッシング、スミッシング、ヴィッシング、スピアフィッシング、ベイト、ブラックメール、プレテキスト。**11語がほぼ全部カタカナ**で並びます。

名前だけを並べて覚えようとする、まず残りません。そこでこの本では、**名前を先に出しません。何が起きるのかを先に書いて、そのあとに名前を付けます**。そのほうが残ります。

この本の構成

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま問題になる所です。試験直前はここだけ拾い読みできます。
- **△マークの枠**（赤い枠）……**間違えやすい所・引っ掛けが仕込まれる所**です。
- **用語**（紫の帯）……シラバスに名前が載っている用語です。この本で25語あります。
- **確認問題**（緑の枠）……テーマの区切りに3問ずつ、合計9問あります。
- **巻末**……25語のチェックリストと、用語集を付けました。

進め方

1. まず**本文を最後まで通して読む**。分からない所で止まらず、飛ばして進んでかまいません。
2. 次に**確認問題9問**を、本文を見ずに解く。ここで落ちた所が、あなたの弱点です。
3. 落ちた所だけ本文に戻り、**自分の言葉で言い直す**。読むより言い直したほうが残ります。
4. 試験直前は、**試験ポイントの枠と、巻末の用語集だけ**を見返す。

動画と合わせて使う

- この本は、YouTubeチャンネル**AI資格ラボ**の動画 **第4章 情報リテラシー ①身を守る** と同じ順番・同じ図で作ってあります。
- **動画で流れをつかみ、この本で確かめる**。あるいは**この本を先に読んでから動画を見る**。どちらでも構いません。
- 印刷は**A4・実物大（100%）**で。図の中の文字まで読める大きさに作ってあります。

はじめに	この本の使い方	104
第4章①	この本の全体像	107
4-1	インターネットリテラシーとは	108
	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	
4-2	だます手口① どの道で来るか	110
	フィッシング詐欺／スミッシング／ヴィッシング	
4-3	だます手口② 何につけこむか	112
	ソーシャルエンジニアリング攻撃／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	
	確認問題① (4-2・4-3)	115
4-4	悪いソフトと、その備え	116
	マルウェア／ランサムウェア／アンチウイルスソフトウェア	
4-5	個人情報保護法の骨組み	119
	個人情報保護法／改正個人情報保護法／個人情報保護委員会／個人情報取扱事業者	
4-6	何が「個人情報」なのか	121
	個人識別符号／要配慮個人情報／機微（センシティブ）情報	
	確認問題② (4-5・4-6)	123
4-7	加工して、使えるようにする	124
	匿名加工情報／マスキング	
4-8	生成AIに個人情報を入れない	126
	入れてはいけない2つの理由／対策2つ	
	確認問題③ (4-7・4-8)	127
まとめ	25語、全部出ました	128
巻末	25語チェックリスト	129
巻末	用語集（25語）	130

第4章① この本の全体像

25語を、**3つのかたまり**に分けて進みます。かたまりごとに確認問題が3問ずつ入ります。

	テーマ	語数
1	インターネットリテラシー	5語
2	セキュリティとプライバシー (だます手口・悪いソフト)	11語
3	個人情報保護の観点	9語

1つ目は言葉の定義だけです。5語しかなく、すぐ終わります。**2つ目が山**で、11語が全部カタカナ。**3つ目**は法律の言葉が並びます。

3つのかたまりの関係

ばらばらに見えますが、**順番につながっています**。

- まず**何ができれば身を守れるのか**という枠組みを置く (4-1)。
- 次に**実際にどう攻められるのか**を見る (4-2～4-4)。攻める側を知らないと守れません。
- 最後に**法律が何を守っているのか**を見る (4-5～4-8)。個人情報という、いちばん狙われるものの話です。

手口の話と法律の話は、別々の暗記ではありません。盗まれると困るものがあるから法律があり、その盗み方が4-2～4-4だ、という順で読んでください。

⇒ この本には**親子の形**が3回出てきます。①インターネットリテラシーとその中身4つ ②ソーシャルエンジニアリング攻撃とその手口4つ ③マルウェアとその一種。**親を先に置いて、子どもをぶら下げる**。この形で覚えると崩れません。

4-1 インターネットリテラシーとは

まず**インターネットリテラシー**。意味は、**インターネットを適切に使うために必要な力**です。それだけです。ひねりはありません。

インターネットリテラシー Internet Literacy

インターネットを適切に利用するために必要な力。この本の**親になる言葉**で、次に出る4つはすべてこの中身。

シラバスは、この力の中身を**4つ**に分けています。ここが5語のうちの残り4語です。

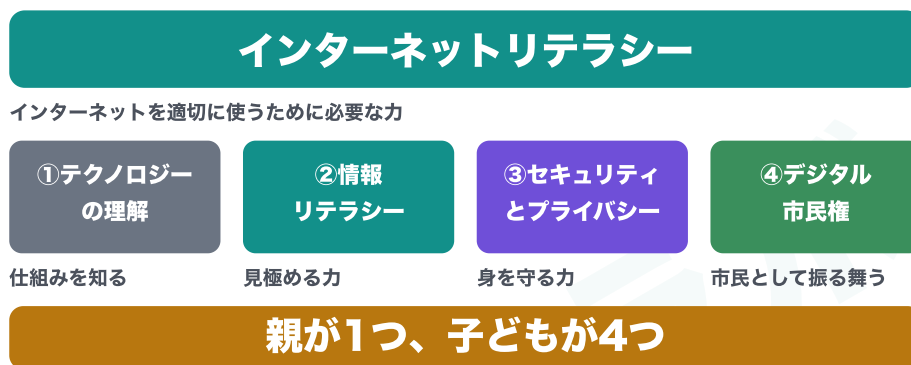


図4-1 インターネットリテラシーが親。中身が4つ

1つ目 テクノロジーの理解

仕組みを知っていることです。なぜ必要か。**知らないものは、疑えない**からです。仕組みが分かっていると、おかしいことが起きても気づけません。

テクノロジーの理解

インターネットや情報technology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。

2つ目 情報リテラシー

情報を見極める力です。その情報は本当か。誰が言っているのか。**確かめる力**のことです。

情報リテラシー

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

3つ目 セキュリティとプライバシー

身を守る力です。この本ではこのあと、テーマを丸ごと1つ使って扱います。11語ぶんあります。

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、設定で自分の情報を守れること。第4章①の山になるテーマ。

4つ目 デジタル市民権

聞き慣れない言葉だと思います。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です。権利があり、責任もある。匿名だから何をしてもいい、の反対です。

デジタル市民権 Digital Citizenship

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する立場で、匿名性を理由に無責任に振る舞うことを否定する。

なぜ、この章の先頭にこれが来るのか

4つを並べると、順番に意味があることが見えてきます。仕組みを知り（テクノロジーの理解）、情報を見極め（情報リテラシー）、身を守り（セキュリティとプライバシー）、そのうえでどう振る舞うかを決める（デジタル市民権）。

この本の残りは、3つ目のセキュリティとプライバシーを丸ごと開いたものです。11語あります。そして個人情報の話が9語。つまり4-1で出た4つのうち、1つを深く掘るのが第4章①だと思ってください。

試験ポイント | 5語は「親と子」で出る

- インターネットリテラシーが親。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つが子ども。
- 情報リテラシーとインターネットリテラシーを取り違えないこと。情報リテラシーは4つのうちの1つで、親ではありません。
- デジタル市民権は名前が難しいだけで、中身はネットでも市民として振る舞う。それだけです。

4-2 だます手口① どの道で来るか

ここから手口です。まず、**何が起きるのか**から書きます。名前はその後です。

起きること

メールが来ます。**銀行です。口座が停止されました。こちらから手続きを**——リンクを踏むと、**本物そっくりの偽サイト**が開きます。IDとパスワードを入れると、そのまま盗まれる。

これが**フィッシング詐欺**です。偽物のメールやサイトで、個人情報を**釣り上げる**。

フィッシング詐欺 Phishing

偽物の**メールやWebサイト**で本物を装い、ID・パスワード・カード番号などの個人情報を入力させて盗む詐欺。

同じことを、別の道で

やっていることは同じで、**来る道だけが違う**手口が2つあります。ここが名前を覚える勘所です。

ショートメッセージ（SMS）で来たら——**スミッシング**です。**SMSとフィッシングを足した言葉**。宅配便の不在通知を装うものが、いちばん多い形です。

スミッシング SMiShing

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

電話で来たら——ヴィッシング。**voice、つまり声のフィッシング**です。カード会社を名乗る電話がかかってきて、暗証番号を聞き出す。

ヴィッシング Vishing

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



図4-2 やることは同じ。違うのは来る道だけ

道はメール・SMS・電話だけではない

シラバスには、キーワードは付いていないが学習項目として挙がっている入口が3つあります。名前を答えさせる問題にはなりにくいものの、**どういう危険があるかは問われます**。

紙から来る 悪意のあるQRコード

QRコードは**人間の目では中身が読めません**。だからリンク先を確かめられない、という弱点があります。掲示物や請求書に貼られた正規のQRコードの上に、**偽のQRコードを重ねて貼る**手口が実際に起きています。

対策は単純で、**読み取ったあとに表示されるURLを見る**こと。そして**シールが二重に貼られていないか**を確かめることです。

回線から来る Wi-Fiに潜む罠

暗号化されていない公衆Wi-Fiでは、**通信の中身が第三者に読まれるおそれ**があります。さらに、店舗名に似せた**偽のアクセスポイント**を立てて接続させ、そこを通る情報を抜く手口もあります。

対策は、**提供元がはっきりしない電波につながらない**こと。つなぐ場合も、**ID・パスワードやカード番号を入力しない**ことです。

ファイル置き場から来る アップロードサービスに潜む詐欺

ファイル共有サービスを装ったメールで、**ファイルを見るにはログインを**と偽のログイン画面へ誘う手口です。**中身はフィッシング詐欺と同じ**ですが、業務で日常的に使うサービスの見た目をしているぶん、警戒が緩みます。

対策は、**メールのリンクからログインしない**こと。いつも使っているブックマークや公式アプリから入り直せば、偽の画面に情報を入れずに済みます。

試験ポイント | 道で覚える

- メールならフィッシング詐欺。SMSならスミッシング。電話ならヴィッシング。
- 3つとも**やっていることは同じ**で、違うのは**来る道**だけ。ここを一度つかむと、二度と混ざりません。
- **スミッシング=SMS+Phishing**、**ヴィッシング=Voice+Phishing**。語源で聞かれることがあります。

4-3 だます手口② 何につけこむか

さきほどの3つは**道**の違いでした。ここからは**やり口**の違いです。まず親になる言葉から置きます。

親 ソーシャルエンジニアリング攻撃

長い名前ですが、意味は単純です。**技術ではなく、人間の心理につけこんで情報を盗む**こと。パスワードを解析するのではなく、**本人に喋らせる**。

ソーシャルエンジニアリング攻撃 Social Engineering

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**。以下の4つはすべてこの中に入る。

ここが**親**です。これから出る4つは、全部この子どもです。



図4-3 親はひとつ。何につけこむかで4つに分かれる

1つ目 スピアフィッシング

スピアは、**槍**。ばらまきではなく、**特定の相手を狙い撃ち**します。**経理部の〇〇さん**と名指しで、社内の事情を混ぜて送ってくる。**引っかけやすさが桁違い**です。

スピアフィッシング Spear Phishing

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、通常のフィッシングより信じやすい。

2つ目 ベイト攻撃

ベイトは、**えさ**。**無料でもらえます、限定公開の動画です**——**欲しがる気持ちを、えさにする**手口です。USBメモリをわざと落としておく、という古典的な手口もこれにあたります。

ベイト攻撃 Baiting

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。

3つ目 ブラックメール

黒いメール、ではありません。脅迫のことです。あなたの秘密を握った。ばらされたくなければ払え——弱みにつけこむ手口です。

ブラックメール Blackmail

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** (blackmail = ゆすり)。

4つ目 プレテキスト

作り話、という意味です。システム部の者ですが、確認のためパスワードを——役になりきって、信じ込ませてから聞き出す。

プレテキスト Pretexting

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。

手口を知ったうえで、自分の設定を見る

手口の名前と並んで、シラバスには**守りの側の学習項目**も挙がっています。ここも名前を問う問題にはなりにくいものの、**考え方は問われます**。

プライバシー設定

SNSやスマートフォンの**公開範囲**の設定です。誰に見せるのか、位置情報を付けるのか、過去の投稿はどうするのか。**初期設定のまま使うと、意図せず全体公開になっていることがあります。**

4-3で見たスパイフィッシングは、**相手の所属や事情を調べたうえで**送られてきます。その材料は、たい
てい**公開されたままの情報**から集められます。設定を見直すことが、そのまま手口への備えになりま
す。

不適切なコンテンツへのWebアクセス

業務用の端末や、子どもが使う端末で、**見せたくないサイトに行かせない**という考え方です。**フィルタリング**と呼ばれる仕組みで、あらかじめ分類されたカテゴリごとに遮断します。

生成AIの技術的發展に潜む脅威

この章が生成AIパスポートに載っている理由のひとつです。生成AIが使えるようになったことで、**だます側の道具も良くなりました。**

- **文面が自然になった**——以前のフィッシングメールは日本語が不自然で見分けられました。いまは**その手がかりが消えつつあります。**
- **声や映像をまねられる**——本人そっくりの音声で電話をかけ、送金を指示する事件が起きています。フィッシングの精度が上がった形です。
- **下調べが速くなった**——公開情報を集めて相手に合わせた文面を作る作業が、短時間でできるようになりました。スパイフィッシングと相性が悪い変化です。

手口そのものは、これまでと同じです。変わったのは**質と量**だけ。だからこそ、名前と形で覚えておく意味があります。

試験ポイント | 何につけこむかで分かれる

- **親**がソーシャルエンジニアリング攻撃。**狙い撃ち**がスパイフィッシング、**えさ**がベイト攻撃、**脅し**がブラックメール、**作り話**がプレテキスト。
- **ブラックメール**を「黒いメール」と読まないこと。**脅迫**です。ここは引っ掛けになります。
- **4-2は道の違い、4-3はやり口の違い**。この2本の軸で整理すると、7語がきれいに並びます。

確認問題① だます手口

確認問題①

第1問 SMSで送られてくる偽メッセージの詐欺を何といいますか。電話の場合、メールの場合はどうですか。

答

スミッシング。電話なら**ヴィッシング**、メールなら**フィッシング詐欺**。

3つとも中身は同じで、違うのは来る道だけです。

確認問題①

第2問 技術ではなく、人の心理につけこんで情報を盗む攻撃の総称を何といいますか。

答

ソーシャルエンジニアリング攻撃。

これが親で、スパイフィッシング・ベイト攻撃・ブラックメール・プレテキストが子どもです。

確認問題①

第3問 特定の相手を狙い撃ちするフィッシングは何ですか。えさで釣る手口は何ですか。

答

スパイフィッシングと**ベイト攻撃**。

スパイは槍、ベイトはえさ。語源で覚えると取り違えません。

4-4 悪いソフトと、その備え

手口の次は**もの**です。ここも親子の形が出てきます。

親 マルウェア

悪意のあるソフトウェアの総称です。**mal**は「悪い」という意味で、ソフトウェアと足してマルウェア。ウイルス、ワーム、トロイの木馬——全部この中に入ります。

マルウェア Malware

悪意のあるソフトウェアの総称。malicious（悪意のある）+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。

その中で、試験に出るのが1つ

ランサムウェア。**ランサムは、身代金**。ファイルを勝手に暗号化して開けなくし、**元に戻す代わりに金を要求する**。病院やインフラが止まる事件は、たいていこれです。

ランサムウェア Ransomware

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ありません**。復旧できなかった例も、払ったうえで再び要求された例もあります。だから**備えのほう的大事**です。

備えの側の語

アンチウイルスソフトウェア。**マルウェアを見つけて、止めて、取り除くソフト**です。

アンチウイルスソフトウェア Antivirus Software

マルウェアを**検出・遮断・駆除**するソフトウェア。⚠ 名前は「アンチウイルス」だが、**ウイルスだけを見るわけではない**。

⇒ ⚠ 名前は**アンチウイルス**ですが、**ウイルスだけを相手にするわけではありません**。マルウェア全般が相手です。ここは地味に引っかけになります。



図4-4 マルウェアが親。ランサムウェアがその一種。備えがアンチウイルスソフトウェア

マルウェアの中身

マルウェアは総称なので、中に何が入っているかも押さえておくと、問題文の言い回しに強くなります。

	どう広がるか	特徴
ウイルス	他のファイルに寄生する	単独では動けず、宿主が要る
ワーム	自分で複製して広がる	宿主が要らず、勝手に増える
トロイの木馬	役に立つソフトを装う	入り込むまで悪さをしない
ランサムウェア	経路はさまざま	暗号化して身代金を要求する

名前を答えさせる問題として出るのはランサムウェアだけですが、選択肢に他が並ぶことはあります。

備えは、ソフトだけではない

アンチウイルスソフトウェアは**入口**の備えです。ランサムウェアに対しては、それだけでは足りません。

- **バックアップ**——暗号化されても、別に控えがあれば戻せます。**払わずに済む唯一の道**です。
- **更新**——古いままのソフトの穴から入られます。
- **怪しいファイルを開かない**——結局ここに戻ります。4-2と4-3で見た手口が入口になります。

ランサムウェアの被害は、こう進む

名前だけでは実感が湧きにくいので、起きることを順に書きます。

1. メールの添付や、古いままの機器の穴から**入り込む**。
2. 気づかれないうちに**社内へ広がる**。
3. ある朝いつせいに**ファイルが開かなくなる**。画面に要求の文面が出る。
4. **業務が止まる**。病院なら診療、工場なら生産、自治体なら窓口が止まります。

近年は**暗号化する前に中身を盗み出しておき、払わなければ公開すると迫る**やり方も増えました。バックアップがあっても、**公開されると困る**という別の圧力をかける形です。

⇒ ⚠ 試験で問われるのは**定義**です。**ファイルを暗号化して、復旧と引き換えに金銭を要求する**——ここを一文で言えるようにしてください。

試験ポイント | 3語セットで出る

- マルウェアが親。ランサムウェアがその一種。アンチウイルスソフトウェアが対策。この3語はセットで問われます。
- ランサム=身代金。暗号化して金を要求する、が定義。**払っても戻る保証はない**ことまで押さえます。
- アンチウイルスソフトウェアはウイルス専用ではない。マルウェア全般が対象です。

4-5 個人情報保護法の骨組み

ここからテーマが変わります。個人情報です。法律の言葉が並びますが、**誰が何をする立場なのか**で整理すると入ります。

法律の名前

個人情報保護法。個人情報を扱うときのルールを決めた法律です。

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。

そして**改正個人情報保護法**という語も、別に載っています。**3年ごとに見直す**と決められていて、時代に合わせて何度も改正されているからです。

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、技術や社会の変化に合わせて内容が更新される。

⇒ ⚠ 改正〇年版の中身までは問われません。**改正され続けている、という事実**を押さえてください。

誰が見張っているのか

個人情報保護委員会です。**国の機関**で、監督や指導をするところ。**委員会**という名前ですが、**国の組織**です。ここは覚えるだけです。

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**。⚠ 民間の団体ではない。

誰がルールを守る側か

個人情報取扱事業者。**個人情報をデータベースにして事業に使っている者**、という意味です。

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。⚠ **規模の下限はない**。

⇒ ⚠ ここに**引っ掛け**があります。**大企業だけでしょ**と思いますよね。**違います**。1件でも扱っていれば、個人事業主でも当てはまります。**規模の下限はありません**。

個人情報保護委員会

国の機関。監督や指導をする
委員会という名前だが、国の組織

個人情報取扱事業者

個人情報をデータベースにして事業に使う者
ルールを守る側

1件でも扱えば当てはまる

図4-5 見張る側が委員会。守る側が事業者

そもそも「個人情報」とは何か

法律は、個人情報を**2つの形**で定義しています。

1. **生存する個人に関する情報**であって、氏名や生年月日などにより**特定の個人を識別できるもの**。他の情報と**容易に照合できて**識別できるようになるものも含まれます。
2. **個人識別符号が含まれるもの**。こちらは4-6で扱います。

⇒ ⚠ **生存する**という条件が付いています。また**容易に照合できれば含む**という書き方なので、**単体では誰か分からない情報でも、社内の名簿と突き合わせれば分かる場合は個人情報にあたります**。ここは考え方として問われます。

試験ポイント | 立場で覚える

- 見張るのが個人情報保護委員会（国の機関）、守るのが個人情報取扱事業者。
- 個人情報取扱事業者に規模の下限はない。**1件でも扱えば当てはまる**——ここが最頻出の引っ掛けです。
- **改正個人情報保護法**は、3年ごとの見直しがある、まで。改正年ごとの中身は問われません。

4-6 何が「個人情報」なのか

ここが**この章でいちばん細かい**所です。3語ありますが、うち2つはほぼ同じ意味なので、実質2つと
思っ
てか
ま
い
ま
せ
ん。

それ単体で個人が分かるもの

個人識別符号。**それ単体で個人を特定できる符号**のことです。

ひとつは**体の特徴をデータにしたもの**。指紋、顔、虹彩^{こうさい}など。もうひとつは**公的な番号**。マイナンバー、旅券番号、運転免許証番号など。

名前が無くても、これだけで誰か分かる。だから個人情報として扱われます。

個人識別符号

それ**単体で特定の個人を識別できる**符号。①**身体の特徴**をデータ化したもの（指紋・顔・虹彩など）②**公的な番号**（マイナンバー・旅券番号・運転免許証番号など）。

特に配慮が要るもの

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、特に配慮が要る情報**です。人種、信条、社会的
的身分、病歴、犯罪の経歴、犯罪被害の事実。

扱いが特別で、**原則、本人の同意なしに取得できません**。**そこが普通の個人情報との違い**です。

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**。

ほぼ同じ意味の、もうひとつの言い方

機微（センシティブ）情報。**要配慮個人情報と、ほぼ同じ意味**です。

違いはどこか。**要配慮個人情報**が**法律の言葉**で、**機微情報**のほうが、**もっと広い一般的な言い方**です。
両方載っているので、**両方覚える**。それだけです。

機微（センシティブ）情報

取り扱いに配慮が要る情報の**一般的な呼び方**。要配慮個人情報とほぼ同じ意味だが、**こちらは法律用語**
ではなく、より広い言い方。

個人識別符号

それ単体で個人を特定できる符号
指紋・顔・虹彩／マイナンバー・旅券番号・運転免許証番号

要配慮個人情報

不当な差別や偏見を受けないよう配慮が要る
人種・信条・社会的身分・病歴・犯罪の経歴
原則、本人の同意なしに取得できない

機微情報

= センシティブ情報
要配慮個人情報と、ほぼ同じ意味
こちらは、もっと広い一般的な言い方

違いは、法律の言葉か、一般の言い方か

図4-6 単体で分かるものと、特に配慮が要るもの

具体例で確かめる

言葉の定義だけでは判断できないので、代表的なものを並べておきます。

例	どれにあたるか	理由
氏名	個人情報	それだけで特定できる
指紋・顔・虹彩のデータ	個人識別符号	体の特徴で特定できる
マイナンバー・旅券番号	個人識別符号	公的な番号で特定できる
病歴・信条・犯罪の経歴	要配慮個人情報	差別や偏見につながりうる
社員番号だけ	場合による	社内名簿と照合できるなら個人情報

⇒ ⚠ 要配慮個人情報は、個人情報の一種です。別のものではありません。個人情報のうち、特に扱いが厳しいものという関係になります。ここも親子の形です。

試験ポイント | 違いを一言で

- 個人識別符号＝それ単体で個人が分かるもの（体の特徴／公的な番号）。
- 要配慮個人情報＝差別につながりうるので特に配慮が要る情報。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＝要配慮個人情報とほぼ同義。法律の言葉か、一般的な言い方かの違い。

確認問題②

個人情報の骨組み

確認問題②

第4問 指紋や顔、マイナンバーのように、それ単体で個人を特定できるものを何とといいますか。

答

個人識別符号。

体の特徴をデータにしたものと、公的な番号の2種類があります。

確認問題②

第5問 病歴や信条のように、差別につながりうるので特に配慮が要る情報を何とといいますか。

答

要配慮個人情報。

原則、**本人の同意なしに取得できません**。そこが普通の個人情報との違いです。

確認問題②

第6問 個人情報保護法を監督する国の機関は何ですか。ルールを守る側の呼び名は何ですか。

答

個人情報保護委員会と個人情報取扱事業者。

事業者は**1件でも扱えば当てはまります**。規模の下限はありません。

4-7 加工して、使えるようにする

個人情報、**守るだけでは終わりません**。使いたい。でも、本人が分かってはいけない。**そこで出てくるのが、次の2語**です。

匿名加工情報

特定の個人を識別できないように加工して、元に戻せないようにした情報です。

ここが大事です。「戻せない」ことが条件。**名前を消しただけでは足りません**。

何が嬉しいか。**本人の同意なしに、第三者へ提供できる**ようになります。だから統計や研究、AIの学習データとして使えます。

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**。この条件を満たせば**本人の同意なしに第三者提供ができる**。

マスキング

情報の一部を隠すことです。名前を山田**にする、カード番号の下4桁だけ残す、顔にぼかしを入れる。

マスキング Masking

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。

⇒ ⚠ **この2つを混ぜないこと**。マスキングは**隠す手口**。匿名加工情報は**加工した結果できあがったもの**で、法律上の区分です。**手口か、区分か**。そこで分けてください。



図4-7 手口としてのマスキング、区分としての匿名加工情報

どうやって作るのか

匿名加工情報は、**気持ちの上で匿名にした**というものではありません。決められた基準にしたがって加工します。

- **氏名などを削除する**——特定の個人が分かる記述を消す。
- **個人識別符号を削除する**——指紋データやマイナンバーは、残っていると4-6のとおり単体で特定できてしまいます。
- **珍しい値をならす**——たとえば年齢116歳のように**1人しかいない値**は、消すか丸めます。
- **元に戻すための情報を残さない**——対応表を持っていたら、それは戻せるということです。

何に使えるのか

条件を満たすと、**本人の同意なしに第三者へ提供できます**。統計、研究、そして**AIの学習データ**。ここが生成AIの話につながります。

ただし、提供する側には**作ったことを公表する**などの義務があり、受け取った側は**元に戻そうとしてはいけません**。自由に使ってよい、という意味ではありません。

混ざりやすいのは、ここ

マスキングは**手口**で、匿名加工情報は**区分**——と書きましたが、実際には**マスキングは匿名加工情報を作るときの手口のひとつ**でもあります。

ただし**マスキングをすれば匿名加工情報になる、わけではありません**。**戻せないこと**という条件を満たして初めて匿名加工情報です。名前の一部を伏せただけなら、それはただのマスキングです。

試験ポイント | 条件と効果

- **匿名加工情報**の条件は**元に戻せないこと**。名前を消しただけでは足りません。
- 満たすと**本人の同意なしに第三者提供ができる**。ここが効果です。
- **マスキングは手口、匿名加工情報は区分**。取り違えを狙った問題が出ます。

4-8 生成AIに個人情報を入れない

最後に、この章が生成AIパスポートに載っている理由です。**生成AIに、個人情報を入力してはいけない**。理由は2つあります。

理由1 学習に使われることがある

入力した内容が**学習に使われることがある**からです。そうすると、**他の人への回答に出てくるおそれ**があります。

理由2 社外に渡していることになる

そもそも**社外に渡していることになる**からです。**個人情報を第三者に提供している**と見なされうる。

やりがちな例

- 顧客名簿を貼り付けて「分析して」とやる。
- 履歴書を丸ごと入れて「要約して」とやる。
- 社外の生成AIに、社内の問い合わせ記録をそのまま入れる。

やりがちです。でも**アウト**です。

対策は2つ

1. 個人が分かる部分を消してから入れる——4-7で出た**マスキング**です。
2. 学習に使わない設定のサービスを選ぶ。

生成AIに個人情報を入力しない

① 学習に使われる

入力した内容が学習に使われることがある
他の人への回答に出てくるおそれ

② 社外に渡している

そもそも社外へ渡していることになる
第三者提供とみなされうる

マスキング

個人が分かる部分を消してから入れる

設定で防ぐ

学習に使わない設定のサービスを選ぶ

図4-8 入れてはいけない理由と、その対策

試験ポイント | 技術ではなく運用

- 理由は2つ。①**学習に使われうる** ②**第三者提供にあたりうる**。
- 対策も2つ。①**マスキングしてから入れる** ②**学習に使わない設定を選ぶ**。
- **技術の話ではなく、運用の話**です。試験でもそう聞かれます。

確認問題③ 加工と、生成AI

確認問題③

第7問 個人を識別できないよう加工し、元に戻せないようにした情報を何といいますか。

答

匿名加工情報。

戻せないことが条件です。名前を消しただけでは足りません。

確認問題③

第8問 情報の一部を隠すことを何といいますか。

答

マスキング。

こちらは手口です。匿名加工情報は区分でした。

確認問題③

第9問 ファイルを暗号化して身代金を要求するマルウェアを何といいますか。

答

ランサムウェア。

ランサムは身代金。マルウェアが親で、これはその一種です。

まとめ 25語、全部出ました

前半で出た25語を、**3つのかたまり**で並べ直します。ここだけ見返せば全体が戻ります。

1 インターネットリテラシー（5語）

インターネットリテラシーが親。中身が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。

2 セキュリティとプライバシー（11語）

道の違いが3つ——メールならフィッシング詐欺、SMSならスミッシング、電話ならヴィッシング。

やり口の違いが4つ——親がソーシャルエンジニアリング攻撃で、狙い撃ちがスパイフィッシング、えさがベイト攻撃、脅しがブラックメール、作り話がプレテキスト。

ものの話が3つ——マルウェアが親、ランサムウェアがその一種、アンチウイルスソフトウェアが対策。

3 個人情報保護の観点（9語）

法律が個人情報保護法と改正個人情報保護法。見張る側が個人情報保護委員会、守る側が個人情報取扱事業者。

何が個人情報かが個人識別符号／要配慮個人情報／機微（センシティブ）情報。加工して使う話が匿名加工情報とマスキング。

⇒ 後半（②）へ続きます。次は**作ったものの権利**と**社会のルール**です。知的財産権、肖像権とパブリシティ権、不正競争防止法、そしてAI社会原則とAI新法。40語あります。

巻末

25語チェックリスト

自分の言葉で説明できたら✓を入れてください。読めるだけでは足りません。

- ☐ インターネットリテラシー
- ☐ 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
- ☐ フィッシング詐欺
- ☐ スミッシング
- ☐ ヴィッシング
- ☐ ソーシャルエンジニアリング攻撃
- ☐ 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
- ☐ マルウェア
- ☐ ランサムウェア
- ☐ アンチウイルスソフトウェア
- ☐ 個人情報保護法
- ☐ 改正個人情報保護法
- ☐ 個人情報保護委員会
- ☐ 個人情報取扱事業者
- ☐ 個人識別符号
- ☐ 要配慮個人情報
- ☐ 機微（センシティブ）情報
- ☐ 匿名加工情報
- ☐ マスキング

巻末 用語集

インターネットリテラシー p.108

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.108

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.108

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。

フィッシング詐欺 p.110

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.110

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.110

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。

ブラックメール p.113

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.113

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。

マルウェア p.116

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。

セキュリティとプライバシー p.109

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.109

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。

ソーシャルエンジニアリング攻撃 p.112

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.112

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.113

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。

ランサムウェア p.116

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。

アンチウイルスソフトウェア p.116

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。

個人情報保護法 p.119

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.119

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.119

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.119

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.121

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.121

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.121

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.124

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.124

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。

第4章② 権利とAI社会原則



この章の解説動画

<https://youtu.be/WTsZ-VMG5S0>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに この本の使い方

この本は、YouTubeチャンネル「AI資格ラボ」の動画「第4章 情報リテラシー ②権利とAI社会原則」に対応した、配布用のテキストです。動画を見ていない方でも、これ一冊で第4章の後半が最後まで読み切れるように書いてあります。

この本が扱う範囲

生成AIパスポートの第4章は、公式シラバス（出題範囲表）で**65語**あります。第1章の39語、第2章の62語、第3章の20語とくらべて**全5章でいちばん多い章**です。1回で扱うには多すぎるので、第2章と同じように**2冊に分けました**。

	扱う範囲	語数
①	自分の身を守る（前の本）	25語
②	作ったものの権利と社会のルール（この本）	40語

この本は**②の40語**を扱います。**作ったものに、誰の、どんな権利があるのか**。そして**AIを使う社会が、何を大事にすると決めたのか**。この2本立てです。

→ **前の本でやった言葉は、やり直しません**。個人情報保護法、要配慮個人情報、匿名加工情報、マスキング。このあたりは①身を守る編のテキストと動画にあります。ただし**考え方は同じ**なので、読んでいなくても進めます。

この本の、いちばんしんどい所

先に正直に言います。この本の山は**後半の原則と指針**です。**基本理念が3つ、AI社会原則が6つ、共通の指針が10**——並んだ言葉を覚える作業になります。

しかも**原則6つと指針10のうち、4語が同じ言葉**です。ここを避けて通ると必ず混ざるので、この本では**正面から並べて比べます**。

この本の構成

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま問題になる所です。試験直前はここだけ拾い読みできます。
- **△マークの枠**（赤い枠）……**間違えやすい所・引っ掛けが仕込まれる所**です。
- **用語**（紫の帯）……シラバスに名前が載っている用語です。この本で40語あります。
- **確認問題**（緑の枠）……テーマの区切りに3問ずつ、合計9問あります。
- **巻末**……40語のチェックリストと、用語集を付けました。

進め方

1. まず**本文を最後まで通して読む**。分からない所で止まらず、飛ばして進んでかまいません。
2. 次に**確認問題9問**を、本文を見ずに解く。ここで落ちた所が、あなたの弱点です。
3. 落ちた所だけ本文に戻り、**自分の言葉で言い直す**。読むより言い直したほうが残ります。
4. 試験直前は、**試験ポイントの枠と、巻末の用語集だけ**を見返す。

動画と合わせて使う

- この本は、YouTubeチャンネル**AI資格ラボ**の動画 **第4章 情報リテラシー ②権利とAI社会原則**と同じ順番・同じ図で作ってあります。
- **動画で流れをつかみ、この本で確かめる**。あるいは**この本を先に読んでから動画を見る**。どちらでも構いません。
- 印刷は**A4・実物大（100%）**で。図の中の文字まで読める大きさに作ってあります。

はじめに	この本の使い方	132
第4章②	この本の全体像	135
4-9	知的財産権は4つ	136
	知的財産権／著作権／特許権／商標権／意匠権	
4-10	人そのものにまつわる権利	138
	肖像権／パブリシティ権	
	確認問題④（4-9・4-10）	140
4-11	不正競争防止法	141
	不正競争防止法／営業秘密／限定提供データ	
4-12	AI生成物と、権利の侵害	143
	生成物／著作権侵害／名誉毀損	
	確認問題⑤（4-11・4-12）	145
4-13	AI社会の基本理念 3つ	146
4-14	AI社会原則 6つ	147
4-15	共通の指針 10	149
	原則6つとの重なり／指針にだけある3つ／高度なAIシステム	
	確認問題⑥（4-13～4-15）	152
4-16	AIガバナンスの構築	153
4-17	AIの事業活動を担う3つの主体	154
4-18	AI新法	155
まとめ	40語、全部出ました	157
巻末	40語チェックリスト	158
巻末	用語集（40語）	160

第4章② この本の全体像

40語を、**3つのかたまり**に分けて進みます。かたまりごとに確認問題が3問ずつ入ります。

	テーマ	語数
1	制作物に関わる権利	13語
2	AIを取り巻く理念と原則・ガイドライン	25語
3	AI新法	2語

3つのかたまりの関係

前半（①）が**守られる側**の話——個人情報という、狙われるものの話でした。この本は**作る側・使う側**の話です。

1. まず**作ったものに、誰のどんな権利があるのか**（4-9～4-12）。ここは**法律**の話です。
2. 次に**AIを使う社会が、何を大事にすると決めたのか**（4-13～4-17）。こちらは**法律ではない**——理念・原則・指針、つまり考え方です。
3. 最後に**それを国の制度にしたもの**（4-18）。2025年にできた新しい法律です。

⇒ ⚠ **2つ目が最大の山**です。25語あり、しかも**原則6つと指針10で4語が重なります**。この本では4-14と4-15で並べて比べるので、そこは飛ばさずに読んでください。

4-9 知的財産権は4つ

まず**知的財産権**。**人が頭で作り出したものを守る権利**の、**総称**です。総称、つまり**親**ですね。子どもを4つ見ます。

知的財産権

人が頭で作り出したもの（知的創作物）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。

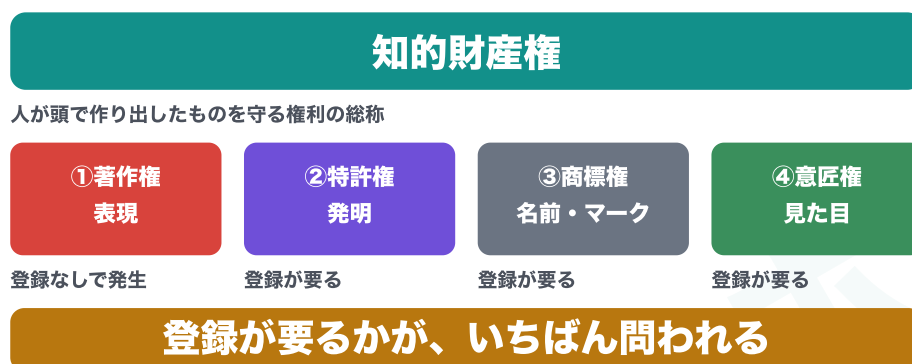


図4-9 総称が親。子どもが4つ

1つ目 著作権

表現したものを守る権利です。文章、絵、音楽、プログラム。

作った瞬間に発生します。登録は要りません。ここが他の3つと違う——この本でいちばん問われる違いです。

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**。

2つ目 特許権

発明を守る権利。**技術的なアイデア**です。こちらは**出願して、審査を通過して、登録して、はじめて発生**します。

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。

3つ目 商標権

商品やサービスの名前やマークを守る権利。**ロゴやブランド名**です。これも登録が要ります。

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。

4つ目 意匠権

見た目のデザインを守る権利。形、模様、色。**中身ではなく、姿かたち**です。これも登録が要ります。

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。

4つを表で見る

	守るもの	登録
著作権	表現したもの	要らない
特許権	発明（技術的なアイデア）	要る
商標権	商品・サービスの名前やマーク	要る
意匠権	見た目のデザイン	要る

著作権だけが浮いているのが見て取れると思います。ほかの3つは**役所に出して、審査を受けて、はじめて権利になる**。著作権は**書いた瞬間・描いた瞬間**です。

理由は守るものの性質が違うからです。**表現**は作った時点でそこにあります。一方**発明**や**マーク**は、誰が先に権利を持つのかを**はっきりさせないと争いになる**。だから登録という手続きが要ります。

試験ポイント | 「登録が要るか」で切る

- 表現が著作権。発明が特許権。名前とマークが商標権。見た目が意匠権。
- **著作権だけが登録なしで発生する**。残り3つは登録が要る。**この違いがいちばん問われます**。
- 知的財産権は総称。「知的財産権と著作権はどちらが広いか」という形でも聞かれます。

4-10 人そのものにまつわる権利

さきほどの4つは**作ったもの**の権利でした。ここからは**人そのもの**の権利です。2つあり、**対で問われます**。

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です。**誰にでもあります**。有名かどうかは関係ありません。

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

こちらは違います。**有名人の名前や顔がもつ、お客を呼ぶ力。その経済的な価値を守る権利**です。

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。有名人に生じる財産的な権利。



図4-10 誰にでもあるか／お金の話か

→ ⚠️ **ここ、対で問われます**。肖像権は**誰にでもある、人格の権利**。パブリシティ権は**有名人の、経済的な権利**。**誰にでもあるか／お金の話か**。そこで分けてください。

生成AIとのつながり

- 実在の人物そっくりの画像を作って公開すると、**肖像権**の問題。
- 有名人の顔を広告に使うと、**パブリシティ権**の問題。

第3章で出た**ディープフェイク**が、ここに効いてきます。技術の話が、そのまま権利の話になる所です。

撮ることと、公開すること

肖像権は**撮られない権利**と**公開されない権利**の両方を含みます。街の風景に偶然写り込む程度なら問題になりにくい一方、**特定の個人だと分かる形で撮って、無断で公開する**と問題になります。

SNSに友人の写真を上げるのも、厳密には同じ話です。**許可を取る**という当たり前の作法が、そのまま法律の話につながっています。

なぜ2つに分かれているのか

肖像権は**人格**を守るもの——**嫌な思いをしない**ための権利です。パブリシティ権は**財産**を守るもの——**その人の名前や顔が生む価値**を勝手に使わせないための権利です。

だから**有名人には両方あります**。無名の人には肖像権はあるが、パブリシティ権を主張する場面はまずない。**パブリシティ権が誰にでもある、という選択肢は誤り**です。

試験ポイント | 2つの軸で分ける

- **肖像権**＝誰にでもある／人格の権利／勝手に撮られない・公開されない。
- **パブリシティ権**＝有名人／経済的な権利／顧客を引きつける力を守る。
- 生成AIでは、**そっくりの画像を公開＝肖像権、有名人の顔を広告に＝パブリシティ権**。

確認問題④ 作ったものと、人の権利

確認問題④

第1問 知的財産権のうち、登録なしで発生するものはどれですか。

答

著作権。

特許権・商標権・意匠権は登録が必要です。ここがいちばん問われます。

確認問題④

第2問 見た目のデザインを守る権利は何ですか。名前やマークを守る権利は何ですか。

答

意匠権と商標権。

表現が著作権、発明が特許権、名前とマークが商標権、見た目が意匠権です。

確認問題④

第3問 誰にでもある人格の権利と、有名人の経済的な権利。それぞれ何といいますか。

答

肖像権とパブリシティ権。

誰にでもあるか／お金の話か、で分けます。

4-11 不正競争防止法

不正競争防止法。事業者どうしの、不正な競争を防ぐための法律です。他社の商品そっくりの見た目で売る。他社の技術情報を盗んで使う。そういうことを止めます。

不正競争防止法

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

この法律で守られるものが、シラバスに**2つ**載っています。

営業秘密 3つの条件を全部満たすこと

3つの条件を全部満たしたもののだけが営業秘密になります。

1. 秘密として管理されていること（秘密管理性）
2. 事業に役立つ情報であること（有用性）
3. 公然と知られていないこと（非公知性）

⇒ ⚠ **秘密として管理されている**が肝です。**誰でも見られる場所に置いてあったら、営業秘密になりません。**
「大事な情報だから営業秘密だ」ではなく、**大事に扱っているから営業秘密**です。

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**——この**3つを全部満たすもの**だけが該当する。

不正競争防止法

事業者どうしの、不正な競争を防ぐための法律

営業秘密

- ① 秘密として管理されている
 - ② 事業に役立つ情報である
 - ③ 公然と知られていない
- 3つ全部を満たしたものだけ

限定提供データ

特定の相手に限って提供するデータ
営業秘密ほど厳しく秘密にしない
でも、誰にでも渡すわけでもない
その中間を守る区分

社外の生成AIに貼ると、管理が崩れる

図4-11 3つ全部そろって、はじめて営業秘密

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のことです。営業秘密ほど厳しく秘密にしていなくても、誰にでも渡すわけではない。**その中間**を守るために作られた区分です。

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。営業秘密ほど厳格に秘密管理されていないが、誰にでも渡すものでもない中間の区分。

生成AIとのつながり

社外の生成AIに、営業秘密を貼り付けて質問する。——それだけで**秘密として管理されているが崩れかねません**。

前半(①)でやった**個人情報を生成AIに入れないと、まったく同じ理屈です**。**社外に渡した時点で、守られる条件を自分から壊している**。

2つを表で見る

	どんな情報	守られる条件
営業秘密	社外に出したくない情報	3条件を全部満たす
限定提供データ	特定の相手に限って渡すデータ	業として反復して提供している

たとえば**顧客名簿や製造の手順書**は営業秘密になりえます。一方、**複数の会社**に**有償で提供している業界データ**のようなものは、秘密ではないけれど誰にでも渡すわけでもない——そこを守るのが限定提供データです。

「大事な情報だから守られる」ではない

ここが誤解されやすい所です。**情報の中身が重要かどうかで決まるものではありません**。どう扱っているかで決まります。

鍵のかかる場所に置いてあるか。アクセスできる人を限っているか。**秘密だと分かる表示をしているか**。**扱い方が条件になっている**——だから生成AIに貼り付けた時点で崩れるのです。

試験ポイント | 条件と中間

- **営業秘密は3条件を全部**——秘密管理・有用性・非公知性。**1つでも欠けたら該当しません**。
- **限定提供データ**は、営業秘密と公開データの**中間**を守る区分。
- 生成AIに貼り付けると**秘密管理性が崩れる**。個人情報と同じ理屈です。

4-12 AI生成物と、権利の侵害

ここで、**AIが作ったもの**の話に入ります。シラバスの言葉は**生成物**です。

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。

AIが作ったものに、著作権はあるのか

いちばん聞かれることです。**日本では、原則、ありません。**

理由は、**著作権が「人間の創作的な表現」を守るもの**だからです。AIが自動で出ただけのものは、人間が作ったとは言えない。

⇒ ⚠ **ただし、例外があります。**人間が**指示や修正を重ねて、創作的に関わった**と言えるなら、認められる。**絶対に無い、ではありません。**関わり方による、が答えです。

AIが作ったものに著作権はあるか

原則、ありません

著作権は人間の創作的な表現を守るもの
AIが自動で出ただけでは創作ではない

認められうる

人間が指示や修正を重ねて
創作的に関わったと言えるなら

著作権侵害

名誉毀損

どちらも「AIが勝手にやった」は通らない。公開した人の責任

図4-12 原則は無い。ただし関わり方による

著作権侵害

AIが出したものが、既存の作品とそっくりだった場合です。

⇒ ⚠ **AIが勝手にやった、は言い訳になりません。**公開した人の責任になります。

著作権侵害

他人の著作物を、権利者の許諾なく利用すること。**AI生成物が既存の作品と類似していた場合も対象**で、責任は**公開した人が負う**。

名誉毀損

事実かどうかにかかわらず、人の社会的評価を下げると成立します。

AIが嘘の経歴を出して、それをそのまま載せた——第2章で出た**ハルシネーション**が、ここで法律の問題になります。

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。AIの誤った出力をそのまま公開した場合も問題になる。

使う前に、何を確認めるか

AI生成物をそのまま世に出すと、ここまでの3つ（著作権・肖像権・名誉毀損）すべてに触れる可能性があります。実務では次を確認めます。

1. 既存の作品に似ていないか——画像なら検索で近いものを探す。
2. 実在の人物が写っていないか——そっくりの顔は肖像権の問題になります。
3. 書かれている事実は本当か——経歴や実績はハルシネーションの温床です。

AI生成物は必ず事実確認をする。試験でもそう聞かれます。第2章で学んだハルシネーションが、この章では法律の問題として戻ってきます。

試験ポイント | 責任は使う人にある

- AI生成物に著作権は原則ない。ただし人間が創作的に関われば認められうる。
- 著作権侵害も名誉毀損も、公開した人の責任。AIのせいにはできません。
- だからAI生成物は必ず事実確認をする。試験でもそう聞かれます。

確認問題⑤ 秘密と、AI生成物

確認問題⑤

第4問 営業秘密になるための3つの条件を挙げてください。

答

秘密として管理されている・事業に有用・公然と知られていない。

3つ全部そろって、はじめて営業秘密です。

確認問題⑤

第5問 AIが作ったものに、日本では原則として著作権はありますか。

答

原則ありません。ただし人間が創作的に関わったと言えるなら認められうる。

絶対に無い、ではありません。関わり方によります。

確認問題⑤

第6問 AIが出した嘘の経歴をそのまま公開しました。何の問題になりますか。

答

名誉毀損。

事実かどうかにかかわらず、社会的評価を下げれば成立します。責任は公開した人にあります。

4-13 AI社会の基本理念 3つ

ここから後半です。日本には、**AIをどう使うかの、国としての考え方**があります。その土台が**AI社会の基本理念**。3つです。

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**にいるという理念。

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**社会という理念。

持続可能な社会 (Sustainability)

続けられる社会という理念。環境も、格差も、ここに入る。



図4-13 基本理念は3つ。英語もセットで

試験ポイント | 英語で覚えると速い

- **Dignity / Diversity & Inclusion / Sustainability**——尊厳、多様性、持続可能性。
- **3つセット**で出ます。1つだけ聞かれることは少なく、「3つ挙げよ」の形が多い。
- ⚠️ 正式な言い方は**多様な背景を持つ人々が多様な幸せを追求できる社会**です。

4-14 AI社会原則 6つ

理念の次がAI社会原則。6つです。理念が「目指す社会」、原則が「そのために守ること」という関係になります。

人間中心の考え方

AIは人間を助けるもので、人間が判断の主であるという原則。

安全性・公平性

AIが安全に動き、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。

セキュリティ確保

攻撃や事故に耐えること。

透明性

どう動いているかを説明できること。

アカウントビリティ

説明責任。結果に責任を持つこと。



図4-14 原則は6つ

⇒ ⚠️ 透明性とアカウントビリティを混ぜないでください。透明性は見えること。アカウントビリティは責任を負うこと。見えるだけでは、責任は果たせません。

理念と原則の関係

並べて覚えるだけでは残らないので、関係で見ます。

	何を言っているか	数
基本理念	目指す社会の姿	3
AI社会原則	そのために守ること	6
共通の指針	事業者が実際に取り組むこと	10

上から下へ、だんだん具体的になります。理念が**どんな社会にしたいか**、原則が**そのために何を守るか**、指針が**では何をするか**。

この順で並んでいると分かれば、「基本理念はどれか」「原則はどれか」を取り違えにくくなります。**数(3・6・10) もセットで覚えてください。**

試験ポイント | 6つを言えるように

- 人間中心の考え方／安全性・公平性／プライバシー保護／セキュリティ確保／透明性／アカウントビリティ。
- ⚠️ **安全性・公平性**でひとつの項目です。ここが次の「共通の指針」との違いになります。
- **透明性＝見える／アカウントビリティ＝責任を負う**。対で聞かれます。

4-15 共通の指針 10

ここが、**この章でいちばん混ざる所**です。いまの**AI社会原則**とは別に、**AI事業者ガイドライン**という文書があって、そこに**共通の指針**が10あります。

先に言います。4つ、さきほどと同じ言葉が出ます。でも、別物です。**そういうものだと思ってください。**



図4-15 共通の指針は10。色が付いた4つは原則と同じ言葉

原則6つと、指針10を並べる

	AI社会原則（6）	共通の指針（10）
1	人間中心の考え方	人間中心
2	安全性・公平性	安全性
3	（同上）	公平性
4	プライバシー保護	プライバシー保護
5	セキュリティ確保	セキュリティ確保
6	透明性	透明性
7	アカウントビリティ	アカウントビリティ
8	—	教育・リテラシー
9	—	公正競争確保
10	—	イノベーション

→ ⚠ **表記が違う所**——原則では**人間中心の考え方**、指針では**人間中心**。原則では**安全性・公平性でひとつ**、指針では**安全性と公平性で2つ**。そのまま覚えてください。

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。⚠ 原則では**安全性・公平性**でひとつの項目になっている。

公平性

共通の指針の1つ。⚠ 同上。指針では**安全性と公平性で2つ**に分かれる。

指針にだけある3つ

教育・リテラシー

AIを**使う人を育てる**こと。指針にだけある項目。

公正競争確保

特定の企業が**独占しない**ようにすること。指針にだけある項目。

イノベーション

新しいものが生まれる余地を残すこと。指針にだけある項目。



図4-16 指針にだけある3つ＝育てる・広げる

試験ポイント | 守りに、攻めが3つ足された

- 原則6つは「守る」話。指針の後ろ3つは「育てる・広げる」話。この見方で区別できます。
- 重なる4語は**プライバシー保護・セキュリティ確保・透明性・アカウントビリティ**。**同じ言葉でも別の文書の項目**です。
- ⚠ **人間中心の考え方（原則）と人間中心（指針）、安全性・公平性（原則）と安全性／公平性（指針）**。表記の違いが問われます。

誰が決めたものか

AI社会原則は国（内閣府）がまとめた考え方、**AI事業者ガイドライン**は総務省と経済産業省がまとめた文書です。**出どころが違うので、言葉が重なっても別物**——そう理解しておく、と、混ざったときに整理できます。

ガイドラインのほうは**事業者向け**です。だから「守る」だけでなく、**育てる・広げる**という項目（教育・リテラシー／公正競争確保／イノベーション）が足されています。

もう1つの指針 | 高度なAIシステム

シラバスには、もう1つ**高度なAIシステムに関する事業者に通の指針**という項目が載っています。

AI事業者ガイドラインの第2部は、**A. 基本理念／B. 原則／C. 共通の指針／D. 高度なAIシステム／E. AIガバナンスの構築**という並びです。いま見てきた順番そのものですね。

Dは、さきほどの10に**加えて**守るものです。大規模で影響の大きいAIを扱う事業者向けの指針で、**広島AIプロセス**の国際指針を反映しています。広島AIプロセスは、**2023年のG7**で立ち上がった国際的な枠組みです。

⇒ ⚠ 中身の12項目を丸暗記する必要はありません。**10の共通の指針に加えてある／高度なAIシステムを扱う事業者向け／元は広島AIプロセス**——この3点で足ります。

確認問題⑥ 理念・原則・指針

確認問題⑥

第7問 AI社会の基本理念、3つは何ですか。

答

人間の尊厳、多様性、持続可能性。Dignity、Diversity & Inclusion、Sustainability。

3つセットで出ます。英語で覚えると速いです。

確認問題⑥

第8問 「どう動いているか説明できる」のと、「結果に責任を負う」のは、それぞれ何ですか。

答

透明性とアカウンタビリティ。

見えるだけでは、責任は果たせません。

確認問題⑥

第9問 共通の指針10のうち、AI社会原則6つには無い3つは何ですか。

答

教育・リテラシー、公正競争確保、イノベーション。

原則は「守る」話、指針の後ろ3つは「育てる・広げる」話です。

4-16 AIガバナンスの構築

AIガバナンスとは、**AIを安全に使い続けるための、組織の仕組み**です。その作り方が**3段**で載っています。

環境・リスク分析

自分たちが**どんな環境**にいて、**どんなリスク**があるかを調べる段。**現状を知る**。

AIガバナンス・ゴール

どういう状態を目指すのかを決める段。**目標を置く**。

AIマネジメントシステム

それを**回し続ける仕組み**を作る段。**運用に落とす**。

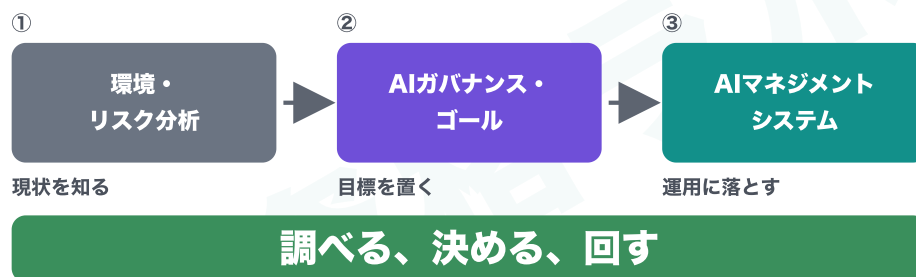


図4-17 調べる → 決める → 回す

試験ポイント | 順番で覚える

- 環境・リスク分析 → AIガバナンス・ゴール → AIマネジメントシステム。
- **調べる、決める、回す**。この順番です。順序を入れ替えた選択肢が出ます。
- 第3章の**AIエージェント**と同じで、分析して、目標を立てて、実行する。人間の仕事の進め方と同じです。

4-17 AIの事業活動を担う3つの主体

AIの事業活動を担う主体が、3つに分けられています。

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。⚠ 作った人と届ける人は別。

AI利用者 (AI Business User)

そのAIを事業に使う側。⚠ 英語はAI Business Userで、単なる User ではない。



図4-18 作る・届ける・使う

⇒ ⚠ **英語に注意。**利用者だけ **AI Business User** です。単なる **User** ではありません。事業として使う人、という意味だからです。

なぜ分けるのか。責任の置き場所が違うからです。作った人の責任、届けた人の責任、使った人の責任。それぞれ果たすことが違うので、3つに分けてあります。

試験ポイント | 英語がそのまま問われる

- AI開発者=AI Developer / AI提供者=AI Provider / AI利用者=AI Business User。
- **利用者だけ Business が入る**——ここが引っかけです。
- 分ける理由は**責任の置き場所が違うから**。

4-18 AI新法

最後、**AI新法**です。これは**2025年10月のシラバス改訂で入った、新しい範囲**です。古い教材には載っていません。

正式名称は**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**。長いので、通称が**AI新法**です。**2025年6月4日に公布**されました。この日付、問われます。

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。

何のための法律か

名前に**推進**と入っているとおり、**AIの研究開発と活用を進めるため**の法律です。

⇒ ⚠ **規制するための法律ではありません。ここ、引っかけになります。** 罰則を並べた法律ではなく、**国と事業者が何をするかを定めた法律**です。

中身と、ガイドラインとの関係

- 国が基本計画を作る。
- 事業者は国の施策に協力する。

そして**AI事業者ガイドライン**とつながっています。**法律で大枠を決めて、具体的な進め方はガイドラインで示す**。そういう関係です。

AI新法（通称）

正式名称：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

2025年6月4日 公布

2025年10月のシラバス改訂で追加

規制ではなく推進

罰則を並べた法律ではない

AI新法



AI事業者ガイドライン

法律で大枠を決めて、具体的な進め方はガイドラインで示す

図4-19 推進の法律。規制ではない

なぜ「推進」なのか

世界の流れを見ると分かりやすくなります。EUのAI法（AI Act）は、危険度に応じて**禁止や義務を課す規制型**です。一方、日本のAI新法は**研究開発と活用を進める推進型**を選びました。

国が基本計画を作り、事業者はそれに協力する。罰則を並べて縛るのではなく、方向を示して後押しする形です。具体的にどう振る舞うかは**AI事業者ガイドライン**が受け持ちます。

⇒ ⚠ **規制型と推進型の対比**で聞かれることがあります。AI新法を「AIを規制する法律」と説明する**選択肢は誤り**です。

試験ポイント | 日付と性格

- 2025年6月4日公布。日付が問われます。
- 推進のための法律で、規制のための法律ではない。罰則を並べたものではありません。
- 法律で大枠、ガイドラインで具体策。4-15の共通の指針とつながります。

まとめ 40語、全部出ました

後半で出た40語を、**3つのかたまり**で並べ直します。ここだけ見返せば全体が戻ります。

1 制作物に関わる権利（13語）

知的財産権が総称で、表現が著作権、発明が特許権、名前とマークが商標権、見た目が意匠権。**著作権だけが登録なしで発生する。**

人そのものの権利が2つ。**誰にでもあるのが肖像権、有名人の経済的価値がパブリシティ権。**

事業者どうしの話が**不正競争防止法**で、守られるものが**営業秘密**と**限定提供データ**。

AIが作ったものが**生成物**。そこで起きうるのが**著作権侵害**と**名誉毀損**。

2 理念と原則・ガイドライン（25語）

基本理念が3つ（尊厳・多様性・持続可能性）、**AI社会原則が6つ**、**共通の指針が10**。そこに教育・リテラシー、公正競争確保、イノベーションが足されます。

ガバナンスが3段（環境・リスク分析 → AIガバナンス・ゴール → AIマネジメントシステム）、**主体が3つ**（AI開発者・AI提供者・AI利用者）。

3 AI新法（2語）

AI新法と、その正式名称**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**。**2025年6月4日公布。推進の法律。**

→ **第4章は、AIを「使う側の責任」を問う章**です。前半は自分の身を守る話、後半は他人の権利と社会のルール。**それだけ持って帰ってください。**

巻末 40語チェックリスト

自分の言葉で説明できたら✓を入れてください。⚠ 原則と指針で同じ言葉が4つありますが、別の文書の項目です。

- ☐ 知的財産権
- ☐ 著作権
- ☐ 特許権
- ☐ 商標権
- ☐ 意匠権
- ☐ 肖像権
- ☐ パブリシティ権
- ☐ 不正競争防止法
- ☐ 営業秘密
- ☐ 限定提供データ
- ☐ 著作権侵害
- ☐ 名誉毀損
- ☐ 生成物
- ☐ 人間の尊厳が尊重される社会 (Dignity)
- ☐ 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)
- ☐ 持続可能な社会 (Sustainability)
- ☐ [原則] 人間中心の考え方
- ☐ [原則] 安全性・公平性
- ☐ [原則] プライバシー保護
- ☐ [原則] セキュリティ確保
- ☐ [原則] 透明性
- ☐ [原則] アカウンタビリティ
- ☐ [指針] 人間中心
- ☐ [指針] 安全性
- ☐ [指針] 公平性
- ☐ [指針] プライバシー保護

- ☐ [指針] セキュリティ確保
- ☐ [指針] 透明性
- ☐ [指針] アカウンタビリティ
- ☐ [指針] 教育・リテラシー
- ☐ [指針] 公正競争確保
- ☐ [指針] イノベーション
- ☐ 環境・リスク分析
- ☐ AIガバナンス・ゴール
- ☐ 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)
- ☐ 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

巻末 用語集

知的財産権 p.136

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.136

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.136

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.137

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。

不正競争防止法 p.141

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.141

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 p.142

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。

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.146

いろいろな背景の人が、それぞれの幸せを目指す社会という基本理念。

持続可能な社会 (Sustainability) p.146

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.147

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。

意匠権 p.137

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.138

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.138

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。

著作権侵害 p.143

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.144

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.143

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.146

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。

安全性・公平性 p.147

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.147

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.147

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.147

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.147

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.150

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.150

共通の指針の1つ。⚠ 原則では「安全性・公平性」でひとつの項目になっている。

環境・リスク分析 p.153

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.153

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.153

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.154

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。

公平性 p.150

共通の指針の1つ。⚠ 原則では「安全性・公平性」でひとつの項目になっている。

教育・リテラシー p.150

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.150

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.150

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。

AI提供者 (AI Provider) p.154

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.154

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

AI新法 p.155

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 p.155

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 プロンプト制作と実例



この章の解説動画

<https://youtu.be/HWYIKwqsWyo>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに この本の使い方

この本は、YouTubeチャンネル「AI資格ラボ」の動画「**第5章 テキスト生成AIのプロンプト制作と実例**」に対応した、配布用のテキストです。動画を見ていない方でも、これ一冊で第5章が最後まで読み切れるように書いてあります。

先に、いい知らせと悪い知らせ

いい知らせから。**この章のキーワードは16語**です。第1章が39語、第2章が62語、第4章が65語ですから、**全5章でいちばん少ない**。

章	扱う内容	語数
第1章	AI（人工知能）	39語
第2章	生成AIの仕組みとChatGPT	62語
第3章	現在の生成AIの動向	20語
第4章	情報リテラシー・AI社会原則	65語
第5章	テキスト生成AIのプロンプト制作と実例	16語

悪い知らせ。**キーワードは少ないのに、学習項目が34あります**。学習項目というのは、シラバスに「これができるようになること」として並んでいる箇条書きのことです。

⇒ つまりこの章は、**名前を覚える章ではなく、できることを覚える章**です。⚠️ **「キーワードが少ない＝出題が少ない」ではありません**。当サイトの練習問題では、**第5章の85問のうち50問**——およそ6割が、**名前の付いていない範囲**から出ています。

それでも、この章がいちばん楽しい

理由はひとつ。**明日から使えるから**です。第1章から第4章までは、中身の話とルールの話でした。第5章は**AIに何をどう頼むか**。手を動かす章です。試験のためだけの知識ではありません。

この本の構成

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま問題になる所です。試験直前はここだけ拾い読みできます。
- **△マークの枠**（赤い枠）……**間違えやすい所・引っ掛けが仕込まれる所**です。
- **用語**（紫の帯）……シラバスに名前が載っている用語です。この本で16語あります。
- **確認問題**（緑の枠）……テーマの区切りに3問ずつ、合計9問あります。
- **巻末**……16語のチェックリストと、用語集を付けました。

進め方

1. まず**本文を最後まで通して読む**。分からない所で止まらず、飛ばして進んでかまいません。
2. 次に**確認問題9問**を、本文を見ずに解く。ここで落ちた所が、あなたの弱点です。
3. 落ちた所だけ本文に戻り、**自分の言葉で言い直す**。読むより言い直したほうが残ります。
4. 試験直前は、**試験ポイントの枠と、巻末の用語集だけ**を見返す。

動画と合わせて使う

- 動画は約30分。この本と**同じ順番・同じ図**で進みます。
- ☆ **読んでから見ると**、動画が復習になります。**見ってから読むと**、この本が確認になります。どちらでもかまいません。
- 練習問題（無料・登録不要）は **ai-shikaku-lab.pages.dev** にあります。間違えた問題から、それを説明している動画の場面へ直接飛べます。

目次

この本の目次

はじめに	この本の使い方	162
第5章	この本の全体像	165
5-1	LMとLLM 言語モデル／n-gramモデル／ニューラル言語モデル／大規模言語モデル／プレトレーニング	166
5-2	出力を決める3つのつまみ ハイパーパラメータ／Temperature／Top-p 確認問題①（5-1・5-2）	169 171
5-3	プロンプトと、プロンプトエンジニアリング	172
5-4	プロンプトの4要素 Instruction／Context／Input Data／Output Indicator	173
5-5	Zero-Shot と Few-Shot 確認問題②（5-3～5-5）	175 177
5-6	文章を扱う 12の技法	178
5-7	仕事で使う 16の型	180
5-8	不得意なこと 3つ 確認問題③（5-6～5-8）	182 183
まとめ	16語、全部出ました	184
巻末	16語チェックリスト	186
巻末	用語集（16語）	187

第5章 この本の全体像

16語を、**3つのかたまり**に分けて進みます。かたまりごとに確認問題が3問ずつ入ります。

かたまり	テーマ	語数
1	言葉の仕組みと、出力のつまみ	10語
2	プロンプトの型	6語
3	何ができて、何ができないか	0語

3つのかたまりの関係

- まずAIが言葉をどう扱っているかと、出力を左右する設定（5-1・5-2）。ここは**仕組み**の話です。
- 次にどう頼むか（5-3～5-5）。**ここが試験の中心**です。4要素と、例を見せるかどうか。
- 最後に何ができて、何ができないか（5-6～5-8）。**ここには名前が付いていません**——でも出ます。

⇒ ⚠ **3つ目にキーワードは1語もありません**。だから軽く見られがちですが、当サイトの練習問題では**85問のうち50問がこの範囲から**出ています。**6割**です。落とさずにいきましょう。

5-1 LMとLLM

まず**言葉のモデル**の話です。LM——**Language Model、言語モデル**。次に来る言葉を予測する仕組みです。

「今日はいい」の次は？ ……「**天気**」。それをやっています。ただそれだけの仕組みが、いまの生成AIの土台です。

LM Language Model

言語モデル。次に来る言葉を予測する仕組みの総称。文章を1語ずつ「次はこれが来やすい」と選んでいく。

LM (Language Model) = 言語モデル

次に来る言葉を予測する仕組み。「今日はいい」の次は「天気」

LLM (Large Language Model) = 大規模言語モデル

とても大きな言語モデル。それだけ

大きいのは、データの量と、パラメータの数

図5-1 LMが親、LLMがその中の大きいもの

作り方が2つ載っています

シラバスには、言語モデルの**作り方が2つ**挙がっています。

1つ目 n-gramモデル

直前のいくつかの単語だけを見て、次を予測するやり方です。**n個のかたまり**、という意味。古いやり方です。

⚠ 弱点は**離れた言葉のつながりが見られない**こと。「私は、去年ロンドンで買った青い傘を、今日も——」の「私は」と「今日も」の関係は、直前の数語しか見ないので捉えられません。

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**。

2つ目 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予測するやり方です。第1章でやったニューラルネットワークが、ここから出てきます。

離れた言葉のつながりも扱える。語の意味の近さも捉えられるので、学習していない言い回しにも対応しやすい。**こちらが今の主流**です。

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や意味の近さを扱える。現在の主流。

① n-gramモデル

直前のいくつかだけ見る

n個のかたまり、という意味。古いやり方
離れた言葉のつながりは見られない

② ニューラル言語モデル

離れた言葉も扱える

第1章のニューラルネットワークを使う
こちらが今の主流

分かれ目は、離れた言葉を見られるか

図5-2 古いのがn-gram、いまの主流がニューラル

試験ポイント | 2つの分かれ目

- **n-gramモデル**＝直前のいくつかだけ。**古い**。離れた関係は**見られない**。
- **ニューラル言語モデル**＝ニューラルネットワークを使う。**今の主流**。離れた関係も**扱える**。
- ☆ どちらが古いか／どちらが主流か、が問われます。

そして、LLM

LLM——**Large Language Model、大規模言語モデル**。とても大きな言語モデル、それだけです。

☆ **何が大きいのか**。……**データの量**と、**パラメータの数**です。パラメータは第2章で出てきた、学習によってモデル内部で調整される値のことでした。

LLM Large Language Model

大規模言語モデル。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**。

→ ⚠ **LMとLLMの違いを聞かれます**。**LMが言語モデル全体**。**LLMはその中の、規模が大きいもの**。——親子ですね。🚫「別物」と書いてある選択肢は誤りです。

最後に、プレトレーニング

プレトレーニング。日本語では**事前学習**です。**大量の文章で、あらかじめ学習させておくこと**を指します。

第2章で**ファインチューニング**という言葉が出てきました。あれは**そのあと**の話です。

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。

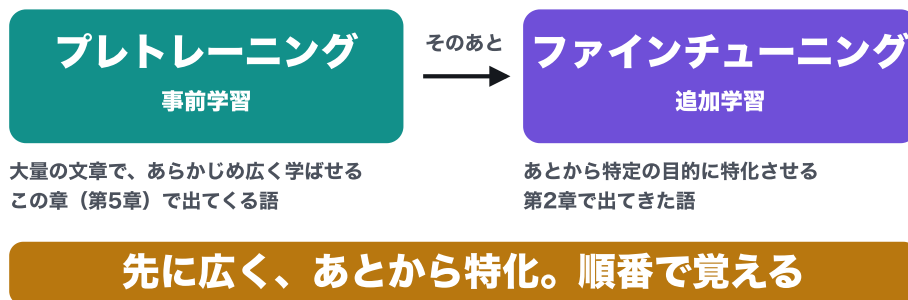


図5-3 事前学習が先、追加学習があと

試験ポイント | 順番で覚える

- 先に広く学ぶのが**プレトレーニング**（事前学習）。
- あとから特化させるのが**ファインチューニング**（追加学習・第2章の語）。
- ☆ **順番を入れ替えた選択肢**が出ます。「先」と「あと」で覚えてください。

5-2 出力を決める3つのつまみ

同じことを聞いても、AIの答えは毎回すこし違います。それを決めている設定の話です。

まず、ハイパーパラメータ

ハイパーパラメータ。人間があらかじめ決めておく設定値のことです。

⚠ AIが学習で決める値ではありません。学習で調整される内部の値はパラメータ（第2章）。ハイパーパラメータは人間が外から決める——ここが問われます。

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。⚠ 学習によってモデル内部で調整されるパラメータとは別物。

その中で、試験に出るのが2つあります。

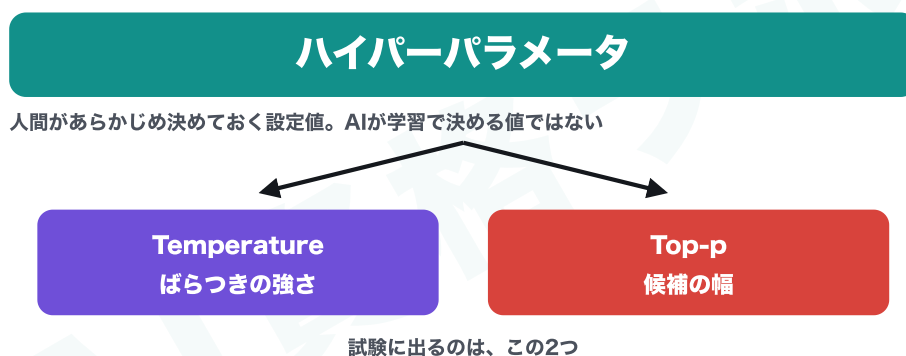


図5-4 親がハイパーパラメータ。子が2つ

1つ、Temperature

Temperature。温度、という意味です。答えのばらつきを決めます。

- 低くすると……いつも同じような、かたい答え。事実を聞くときは低く。
- 高くすると……ばらついた、意外な答え。アイデア出しは高く。

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。



図5-5 Temperatureは、ばらつきのつまみ

2つ、Top-p

Top-p。 次の単語を選ぶとき、候補をどこまで広げるかを決めます。確率の高いものから足していって、**合計が p になるまで**を候補にする、という決め方です。

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。

⇒ ⚠ **Temperatureと役割が似ています。** ☆ 分け方はこうです。 **Temperature はばらつきの強さ。** **Top-p は候補の幅。** ——どちらも「出力の多様さを調整するつまみ」で、どちらも**人間が事前に決める**点は同じです。

試験ポイント | また親子

- ハイパーパラメータが親。Temperature と Top-p が、その例。
- 3つとも**出力を調整するつまみ**で、人間が学習の前に決める値です。
- ⚠ 「AIが学習で決める」と書いてある**選択肢は誤り**です。

— 確認問題① (5-1・5-2)

本文を見ずに答えてみてください。答えは各問のすぐ下にあります。

確認問題①

第1問 直前のいくつかの単語だけを見て次を予測する、古いやり方の言語モデルは何ですか。

答

n-gramモデル。

いまの主流は**ニューラル言語モデル**です。離れた言葉のつながりを扱えるかどうかが目でした。

確認問題①

第2問 答えのばらつきを決める設定と、候補の幅を決める設定。それぞれ何といいますか。

答

Temperature と **Top-p**。

どちらも**ハイパーパラメータ**です。人間が学習の前に決める値でした。

確認問題①

第3問 大量の文章であらかじめ学習しておくことを何といいますか。

答

プレトレーニング（事前学習）。

そのあと目的に合わせるのが**ファインチューニング**。**順番**で覚えます。

5-3 プロンプトと、プロンプトエンジニアリング

プロンプト。AIへの指示文です。私たちが打ち込む、あの文章のことです。

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。

そして**プロンプトエンジニアリング**。欲しい答えが出るように、プロンプトを工夫することです。

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。

なぜ工夫が要るのか

同じことを聞いても、頼み方で答えが変わるからです。

頼み方	返ってくるもの
「まとめて」だけ	何行で、誰向けで、 どんな形か分からない
「小学生にも分かるように、3行でまとめて」	狙いどおりに返る

→ ☆ この章の残りは、全部この話です。どう頼めば、欲しいものが出てくるか。それを型にしていきます。

試験ポイント | 3つを混ぜない

- **プロンプト**＝指示文そのもの。
- **プロンプトエンジニアリング**＝その指示文を工夫すること。モデルには手を加えない。
- ⚠ **ファインチューニング**＝モデル自体を追加学習で調整すること（第2章）。❌ この3つを取り違える**選択肢**が出ます。

5-4 プロンプトの4要素

プロンプトは、**4つの要素**でできています。★ **この章でいちばん出ます。4つとも英語で覚えてください。**

1つ目 Instruction（指示）

何をしてほしいのか。「要約して」「翻訳して」「案を出して」。

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。

2つ目 Context（文脈・背景）

どういう状況で、誰に向けたものか。「新入社員向けの資料です」「相手は取引先の部長です」。

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。

3つ目 Input Data（入力データ）

AIに処理してほしい、中身そのもの。要約したい文章、翻訳したい文、分析したい数字。

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。

4つ目 Output Indicator（出力指示）

どんな形で返してほしいか。「3行の箇条書きで」「表にして」「敬語のメール文で」。

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。



図5-6 4つの要素と、それぞれの役割

もう一度、ひとことで

要素	日本語	ひとことで
Instruction	指示	何を
Context	文脈・背景	どんな状況で
Input Data	入力データ	何を材料に
Output Indicator	出力指示	どんな形で

⇒ ⚠ 4つ全部を入れる必要はありません。でも、**答えが思ったものと違うとき、たいていどれかが抜けています。**★ 抜けているものを足す——それがプロンプトエンジニアリングです。

試験ポイント | 英語と日本語を対で

- Instruction=指示/Context=文脈・背景/Input Data=入力データ/Output Indicator=出力指示。
- ★ 「3行の箇条書きで」は Output Indicator——具体例からどの要素かを当てさせる問題が出ます。
- 🚫 日本語だけ、英語だけでは足りません。両方セットで覚えてください。

5-5 Zero-Shot と Few-Shot

次は、**例を見せるか、見せないか**の話です。

Zero-Shot プロンプティング

例を1つも見せずに、いきなり頼むやり方です。「この文を要約して」。**ふだんやっているのは、だいたいこれ**です。

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、頼むやり方です。「こういう入力には、こう答えてほしい」と2つ3つ見せてから、本番を渡します。

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。



図5-7 例ゼロがZero-Shot、例ありがFew-Shot

何が違うのか

Few-Shot は、形をそろえたいときに効きます。**分類、言い換え、決まった書式**。こういう作業は、言葉で説明するより**例を見せたほうが速い**のです。

逆に、質問が単純で説明が要らないなら、例を足す手間のぶんだけ損になります。**使い分け**で覚えてください。

⇒ ⚠ Shot は「回数」ではなく「例の数」です。**何回やり取りしたか、ではありません。Zero = 例ゼロ。Few = 例が少しある。**——そこだけ押さえてください。

迷ったときの使い分け

こういうとき	どちらで頼むか
質問が単純で、説明が要らない	Zero-Shot
出力の形式（表・箇条書き）をそろえたい	Few-Shot
語調や言い回しをまねさせたい	Few-Shot
分類の基準を言葉で説明しにくい	Few-Shot

試験ポイント | 使い分け

- 単純な作業は **Zero-Shot**。形式や語調をそろえたいときは **Few-Shot**。
- ⚠️ 「常にFew-Shotが優れている」は誤りです。目的で使い分けます。
- ☆ **Shot=例の数**。回数と読み替える選択肢は誤りです。

— 確認問題② (5-3～5-5)

確認問題②

第1問 プロンプトの4要素を、英語で言えますか。

答

Instruction、Context、Input Data、Output Indicator。

日本語では、何を／どんな状況で／何を材料に／どんな形で、です。

確認問題②

第2問 例を見せずにいきなり頼むやり方は？ 例を見せてから頼むやり方は？

答

Zero-Shot プロンプティング と **Few-Shot プロンプティング**。

Shot は回数ではなく例の数でした。

確認問題②

第3問 「3行の箇条書きで」は、プロンプトの4要素のどれにあたりますか。

答

Output Indicator (出力指示)。

どんな形で返してほしいか、を伝える部分です。

5-6 文章を扱う 12の技法

ここからは**名前の無い範囲**です。**⚠でも出ます**。さっき言ったとおり、当サイトの練習問題では6割がここからでした。

文章を扱う技法が、シラバスに**12**並んでいます。1つずつ丸暗記すると必ず抜けるので、**★動詞で束ねます**。**5つのかたまり**です。

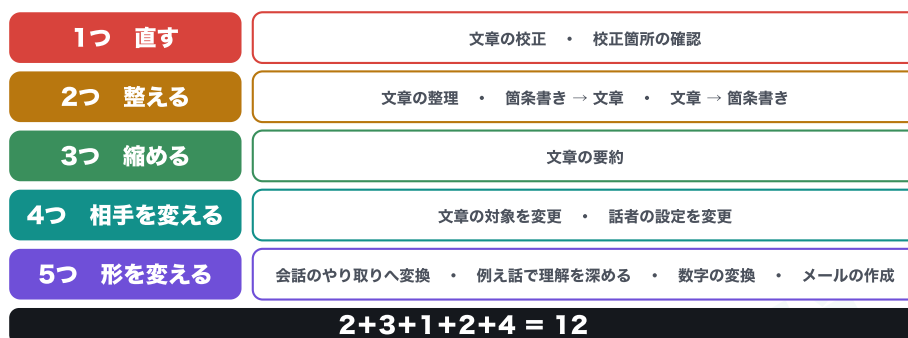


図5-8 12の技法は、全部この5つのどれか

1つ、直す

文章の校正——誤字や言い回しを直してもらおう。そして**校正箇所の確認**——**★どこを直したのかを、一覧で出させる**。

⇒ **⚠これが大事です**。**直った文だけ受け取ると、何が変わったか分かりません**。校正を頼むときは、**直した箇所も一緒に出させる**——実務でも試験でも、ここが問われます。

2つ、整える

文章の整理——ばらばらの内容を、筋の通った順に並べ直す。**箇条書きを文章に変換**と、**文章を箇条書きに変換**。**両方向**です。

メモから報告書へ。報告書からメモへ。どちらもよく使います。

3つ、縮める

文章の要約。……**★「要約して」だけでは足りません**。何行で、誰向けにを足します。さっきの**Output Indicator** ですね。

4つ、相手を変える

文章の対象を変更する——専門家向けを、初心者向けに書き直す。**話者の設定を変更する**——**★AIに役を与える**。「あなたは校長先生です」。

同じ内容でも、言葉づかいが変わります。**⚠読み手を変えるのと語り手を変えるのは別物**です。

5つ、形を変える

- **文章を会話のやり取りへ変換**——説明文を、対話の形に。
- **例え話で理解を深める**——「小学生に説明するなら」。
- **数字の変換**——単位や表記をそろえる。
- **メールの作成**——用件を渡して、メール文にしてもらう。

試験ポイント | 5つの動詞

- **直す・整える・縮める・相手を変える・形を変える。**
- 内訳は $2+3+1+2+4=12$ 。★ 数が合うことを確かめると、抜けに気づけます。
- ⚠ **校正には「校正箇所の確認」を必ずセットで**——ここは単独で問われます。

5-7 仕事で使う 16の型

次は、仕事で使う型が16。こちらも★束ねます。6つのかたまりです。

1つ 聞いて、読む	アンケート項目の作成 ・ アンケートの分析
2つ 作る	キャッチコピー ・ ビジネス書類のテンプレート ・ アジェンダ ・ 海外企業宛のメール
3つ 分解する	業務の手順を分解 ・ タスクの抽出
4つ 言葉を渡す	外国語の翻訳 ・ 英単語から英文の作成
5つ データを整える	姓と名の分離 ・ ふりがなの記載
6つ 一緒に考える	ブレインストーミング ・ ディベートを行う ・ 質問させながら進める ・ 正確な文字数の指定
2+4+2+2+2+4 = 16	

図5-9 16の型は、全部この6つのどれか

1つ、聞いて、読む

アンケート項目の作成——何を聞くべきかを出させる。**アンケートの分析**——集まった自由記述を、傾向にまとめさせる。

⚠ 作るときと、読むとき。両方に使えます。

2つ、作る

キャッチコピーの作成。**ビジネス書類のテンプレート作成**。**アジェンダの作成**。そして**海外企業宛のメール文章の作成**。

ゼロから作るより、たたき台を出させて直すほうが速い——これが「作る」の本質です。

3つ、分解する

業務の手順を分解——ひとつの仕事を、手順に割る。**タスクの抽出**——議事録や会話から、やることだけを抜き出す。

★ 第3章のAIエージェントと同じ考え方ですね。**大きい仕事を、小さく割る。**

4つ、言葉を渡す

外国語の翻訳。**英単語から英文の作成**——単語だけ渡して、文にしてもらおう。英語が苦手でも使えます。

5つ、データを整える

姓と名の分離——「山田太郎」を、「山田」と「太郎」に分ける。**ふりがなの記載**——名簿にふりがなを付ける。

→ ⚠️ **地味ですが、事務ではいちばん効きます。**ただし**読み方が複数ある姓名では誤ります**ので、結果は必ず確認してください。

6つ、一緒に考える

- **ブレインストーミング**——案を数多く出させる。
- **ディベートを行う**——★ **賛成と反対、両方の立場で議論させる。**一人では出ない視点が出ます。
- **質問させながら一緒に進める**——★ 「足りない情報があれば質問して」と**先に頼む**。
- **正確な文字数の指定**——⚠️ **AIは文字数を正確に数えるのが苦手**です。指定はできますが、ずれます。必ず数えること。

試験ポイント | 6つの動詞

- 聞いて読む・作る・分解する・言葉を渡す・データを整える・一緒に考える。
- 内訳は $2+4+2+2+2+4=16$ 。
- ★ **ディベート**（賛成と反対の両方）と**質問させながら進める**（先に頼む）は、単独で問われます。

5-8 不得意なこと 3つ

最後に、**できないことです。★3つだけ。ここは確実に取れます。⚠ただし理由まで問われるので、理由をセットで覚えてください。**

1つ、計算

⚠ AIは計算が苦手です。理由は、**計算しているのではなく、次に来そうな数字を予測している**から。5-1でやった「次に来る言葉を予測する仕組み」が、数字にもそのまま働いています。

桁が多いほど間違えます。電卓や表計算など、**確実な道具に任せてください。**

2つ、最新の情報

学習した時点までしか知りません。昨日のニュースは答えられない。

★ 第3章のRAGは、この弱点への答えでした。外部の文書を検索して、その内容をもとに答えさせる仕組みです。

3つ、芸術の批評

良し悪しを決めることです。AIは「**多くの人がそう言っている**」を返します。**それは批評ではなく、平均**です。

⚠ **価値の判断は、人間の仕事。**正解が定まらない事柄は、仕組みの上で扱えません。

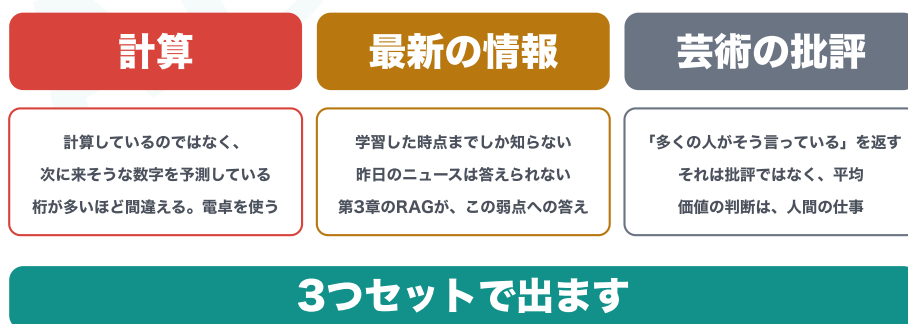


図5-10 できないことは、この3つだけ

試験ポイント | 3つセットで出ます

- 計算……予測しているだけ。**桁が多いほど間違える。**
- 最新の情報……学習時点まで。**RAGが答え。**
- 芸術の批評……返るのは平均。**価値の判断は人間。**
- ⚠ **「何度も聞けば正しい答えに収束する」は誤り**です。

確認問題③ (5-6～5-8)

確認問題③

第1問 AIに役を与えて言葉づかいを変えることを、何といいますか。

答

話者の設定を変更する。

「あなたは校長先生です」ですね。⚠️ **読み手**を変える「文章の対象を変更する」とは別です。

確認問題③

第2問 賛成と反対、両方の立場で議論させる使い方は何ですか。

答

ディベートを行う。

一人では出ない視点が出ます。「一緒に考える」のかたまりでした。

確認問題③

第3問 校正を頼むとき、一緒に頼むべきことは何ですか。

答

校正箇所の確認。

直った文だけでは、何が変わったか分かりません。どこを直したかを出させます。

まとめ 16語、全部出ました

第5章を、5つに畳んで置いておきます。

1つ、言葉のモデル

LMが言語モデル。LLMはその中の大きいもの。作り方が **n-gramモデル** と **ニューラル言語モデル**。あらかじめ学ぶのが **プレトレーニング**。

2つ、つまみ

ハイパーパラメータが親で、**Temperature** がばらつき、**Top-p** が候補の幅。

3つ、プロンプト

プロンプトが指示文、プロンプトエンジニアリングがその工夫。4要素が **Instruction・Context・Input Data・Output Indicator**。例を見せないのが **Zero-Shot プロンプティング**、見せるのが **Few-Shot プロンプティング**。

4つ、できること

文章は **直す・整える・縮める・相手を変える・形を変える**の5つ。仕事は **聞いて読む・作る・分解する・言葉を渡す・データを整える・一緒に考える**の6つ。

5つ、できないこと

計算。最新の情報。芸術の批評。

⇒ ☆ 第5章は、AIに「どう頼むか」の章です。それだけ持って帰ってください。

これで全5章

第1章	AI	39語
第2章	生成AIの仕組みとChatGPT	62語
第3章	いまの動向	20語
第4章	情報リテラシーと社会のルール	65語
第5章	プロンプト	16語

練習問題と模擬試験は、サイトにあります

図5-11 第1章から第5章までの地図

ここまで来た方は、もう十分に戦えます。このあとは**間違えた問題に戻ってください**。当サイトには**全5章の練習問題**と、本番と同じ形の**模擬試験**があり、間違えた問題から、それを説明している動画の場面へ直接飛べます。

試験直前 | 第5章はここだけ

- **プロンプトの4要素**を、英語と日本語で言えるか。(5-4)
- **Zero-Shot** と **Few-Shot** の使い分け。 **Shotは例の数**。(5-5)
- **不得意な3つ**を、理由つきで言えるか。(5-8)
- ☆ この3つで、第5章の出題のほとんどが取れます。

⇒ ☆ **全5章、おつかれさまでした**。あとは手を動かすだけです。試験までの残りは、**間違えた問題を1周する**のがいちばん効きます。

巻末 16語チェックリスト

自分の言葉で説明できたら✓を入れてください。⚠ 英語の4要素は、英語と日本語の両方で言えるか確かめてください。

- ☐ LM
- ☐ n-gramモデル
- ☐ ニューラル言語モデル
- ☐ LLM
- ☐ プレトレーニング
- ☐ ハイパーパラメータ
- ☐ Temperature
- ☐ Top-p
- ☐ プロンプト
- ☐ プロンプトエンジニアリング
- ☐ Instruction
- ☐ Context
- ☐ Input Data
- ☐ Output Indicator
- ☐ Zero-Shot プロンプティング
- ☐ Few-Shot プロンプティング

⚠ 名前の無い範囲も、ここで確かめてください。この本でいちばん出る所です。

- ☐ 文章を扱う**12の技法**を、5つの動詞で言えますか。(直す・整える・縮める・相手を変える・形を変える)
- ☐ 仕事で使う**16の型**を、6つの動詞で言えますか。(聞いて読む・作る・分解する・言葉を渡す・データを整える・一緒に考える)
- ☐ 不得意な**3つ**を、**理由つきで**言えますか。(計算・最新の情報・芸術の批評)

巻末 用語集

LM p.166

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

n-gramモデル p.166

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.167

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。

ハイパーパラメータ p.169

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.169

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.170

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。

Instruction p.173

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.173

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.173

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。

LLM p.167

大規模言語モデル (Large Language Model)。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.168

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。

プロンプト p.172

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.172

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。

Output Indicator p.173

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.175

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.175

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。

第2部

解き方編

2つで迷う問題・言い回しが変わる問題

AI資格ラボ

難問対策編



この章の解説動画

<https://youtu.be/pcoWQsyaHE0>

QRから直接ひらけます（無料・AI資格ラボ）

はじめに

この本は、講義とは別の本です

第1章から第5章までの講義で、**知識はひとつとおりました**。この本は、そのあとの話です。**知識ではなく、解き方**を渡します。

先に、はっきりさせておきます

△ この試験に、公開された過去問はありません。「過去問」を名乗るものは、すべて推測です。別の試験の問題が混ざっていた、という報告まであります。

★ こちらが持っているのは2つだけです。公式が出している例題3問と、自作の練習問題537問の実測。それ以上のことは書きません。

なぜ、この本が要るのか

受験した方が、こう書かれていました。

⇒ 「多くの設問で、選択肢は2つまで絞れていました」

「でも、そこから決めきれなかった」

★ 分かっていないわけではありません。2つまでは、ちゃんと絞れている。……最後の一步が、届かない。この本は、その一步の話です。

この本の構成

- **本文**……ふつうに読める文章です。まずここを読んでください。
- **試験ポイント**（青い枠）……そのまま使える形にまとめた所です。
- **△マークの枠**（赤い枠）……**やってはいけないこと・言い切れないこと**です。
- **やってみる**（緑の枠）……**サイトに実在する難問**を3問。一字一句そのまま載せています。
- **巻末**……4つの道具のチェックリストと、本番前の読み方の手順を付けました。

動画と合わせて使う

- 動画は約18分。この本と**同じ順番・同じ図**で進みます。
- **★道具は使って初めて手に入ります**。実演の3問は、必ず自分で解いてから答えを見てください。
- 練習問題（無料・登録不要）は ai-shikaku-lab.pages.dev。章のページで「**応用・難問だけ**」を選ぶと、この本の型の問題だけが並びます（**全5章で69問**）。

目次

この本の目次

はじめに	この本は、講義とは別の本です	189
1	難しさの正体は2つ	192
2	手がかりは、公式の例題3問だけ	193
3	本番の作り 3つの型	194
	やってみる① 定義提示の型	195
4	道具1 言い過ぎを探す	196
5	道具2 逆にした選択肢を疑う	197
	やってみる② 逆にした選択肢の型	198
6	道具3 「適切／不適切」を先に押さえる	199
7	道具4 2択で止まったら、差分だけを見る	200
	やってみる③ 2択で止まる型	201
8	言い回し耐性 1つの論点を3通りで	202
まとめ	今日の4つの道具	203
巻末	本番の読み方（手順）	204
巻末	4つの道具チェックリスト	205

1 難しさの正体は2つ

難しさの正体は**2つ**あります。⊖ **「2択で迷う」だけではありません。**

1つ目 2つまで絞れるのに、決めきれない

さきほどの言葉です。……これは**分かります**。誰でもそうなります。

2つ目 言い回しが変わると、解けなくなる

★ **問題は、こちらです。**受験した方の言葉を、もう1つ。

△ 「問題集を周回するうちに、問題文と答えの組み合わせを、**そのまま覚えてしまっていた**」
「表現が変わった瞬間、その記憶は**使い物にならなくなります**」

☆☆ **ここが、いちばん怖いところ**です。**周回すればするほど、速く解けるようになる。**……でもそれは、**分かったから**ではなく**覚えたから**かもしれません。

△ **自分では区別が付きません。本番まで、気づけない。**だからこの本は、引っかけの見抜き方だけでは終わりません。「覚えているだけ」を自分で見つける方法まで扱います（8章）。

① 2つまで絞れる。でも決めきれない

選択肢は2つに絞れている
最後の一步が届かない
＝ 分かっていないわけではない

② 言い回しが変わると解けない

問題文と答えの組を覚えていた
表現が変わると使えなくなる
＝ 周回するほど気づけない

②は「分かった」と「覚えた」の区別がつかない

本番まで気づけないのは、こちら

図1 ②は「分かった」と「覚えた」の区別がつかない

試験ポイント | 2つを分けて考える

- ①**2択で決めきれない**……道具で減らせます（4～7章）。
- ②**言い回しが変わると解けない**……**道具では直りません**。練習の形を変えるしかありません（8章）。
- ☆ **この2つを混ぜないこと**。直し方がまったく違います。

2 手がかかりは、公式の例題3問だけ

本番の問題は、どういう作りなのか。……手がかかりは、ひとつしかありません。**公式が出している、例題3問**です。

△ **過去問は公開されていません**。出回っているものは**全部が推測**です。別の試験の問題が混ざっていた、という報告まであります。

先にお断り。3問は、統計になりません

「何割が否定形」といった話は、**この3問からは言えません**。★言えるのは**作りの型**だけです。

そして**型なら、3問でも十分に見えます**。実際、**3つとも違う型**でした。

試験ポイント | 根拠の置き方

- この本が根拠にしているのは**3つだけ**——受験した方の証言／公式の例題3問／自作537問の実測。
- ⚡ **それ以外を根拠にしません**。「たぶんこう出る」は書きません。
- ⚠ 出題範囲の正本は公式シラバスです (guga.or.jp)。この本は受験対策で、試験の実施団体とは関係ありません。

3 本番の作り | 3つの型

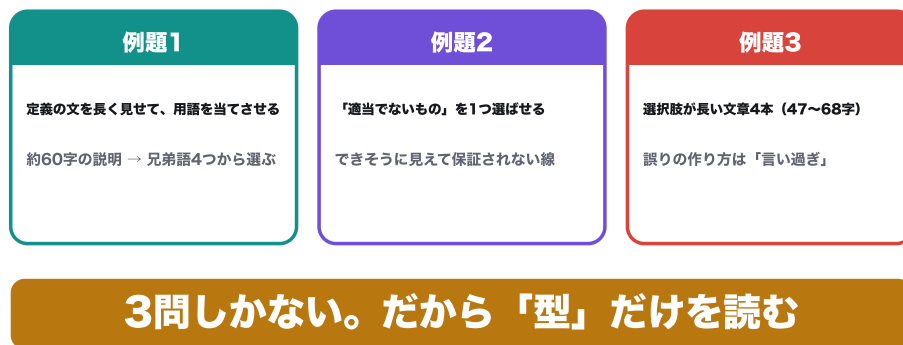


図2 3問しかない。だから「型」だけを読む

型① 定義の文を長く見せて、用語を当てさせる

60字ほどの説明文が先にあって、「これは何か」と聞かれます。そして**選択肢は兄弟語が4つ**——教師あり学習／教師なし学習／強化学習／過学習。**全部、近い言葉**です。

△ こちらの練習問題は、設問の長さが**中央28字**でした。**本番は、その倍あります**。長い文を読まされる、と思っておいてください。

型② 「適当でないもの」を1つ選ばせる

答えになっていたのは「**正確な情報の抽出**」。**★できそうに見えて、保証はされていない**——この線です。

型③ 選択肢が、長い文章4本

1本あたり**47字から68字**。4本とも文章なので、**読み比べ**になります。

⚠ そして**誤りの作り方**。……**★「言い過ぎ」**でした。「**厳格に法律で規制されている**」——そこまでは言っていない、という作り方です。

⇒ **この3つ目が、いちばん大事な手がかり**です。次の章から、道具の話に入ります。

ー やってみる① 定義提示の型

サイトに実在する問題です。答えを見る前に、自分で解いてください。

やってみる①

第1問 学習に用いる着目点（特徴量）を人間が指定するのではなく、大量のデータからAI自身が見つけ出して判断に用いる段階を指す。この説明にあてはまるものはどれか。

ア レベル1 イ レベル2 ウ レベル3 エ レベル4

答

エ レベル4。

★ 見るところはひとつだけ——「着目点を、誰が決めたか」。人間が決めていれば**レベル3**、AIが見つけていれば**レベル4**。

⚠ 学習しているかどうかでは分かれません。レベル3も学習しています。

試験ポイント | 長い定義文の読み方

- ★ 長い定義文は、ほとんどが飾りです。
- 分かれ目になる一文だけを探す。それがこの型の解き方です。
- ⚠ 兄弟語が並ぶので、「**どれも当てはまりそう**」に見えます。分かれ目を見つけるまで選ばないこと。

4 道具1 | 言い過ぎを探す

公式の例題3でも、誤りの作り方は「言い過ぎ」でした。……これは目印で見つかります。

この6つが入っていたら、まず目を留める

必ず

すべて

完全に

一切

保証される

常に

この分野に、言い切れることがほとんどない

出力の正しさは保証されない／シンギュラリティは証明されていない

RAGを入れても誤りは無くならない

ただし「疑う」まで。切るのは中身を見てから

図3 この6つが入っていたら、まず目を留める

なぜ効くのか

★ この分野に、言い切れることがほとんどないからです。

- 生成AIの出力は、正しさが保証されない
- シンギュラリティは、証明されていない
- RAGを入れても、誤りは無くならない

ただし、万能ではありません

△ 言い切りが正しい場合もあります。「著作権は、創作した時点で発生する」——これは言い切って正しい。

試験ポイント | 「疑う」までで止める

- 目印は必ず／すべて／完全に／一切／保証される／常にの6つ。
- ★ 道具は「疑う」まで。「即座に切る」ではありません。
- まず目を留める。それだけで、だいぶ違います。

5 道具2 | 逆にした選択肢を疑う

この試験には、2つセットで出てくる言葉がとても多い。

2つ並んだら、必ず「入れ替えた選択肢」が用意される

エンコーダ	入力をまとめる	デコーダ	出力を組み立てる
Temperature	ばらつきの強さ	Top-p	候補の幅
営業秘密	秘密として管理	限定提供データ	相手を限って提供
肖像権	誰にでもある	パブリシティ権	お金の話
プレトレーニング	先	ファインチューニング	あと

短い一言で「どっちがどっち」まで覚える

図4 短い一言で「どっちがどっち」まで覚える

★2つ並んだら、必ず「入れ替えた選択肢」が用意されます。これは**作る側の都合**です——もっともらしい誤答を作るのに、いちばん簡単な方法だから。

だから、対策も決まっています

2つセットの言葉は、必ず「どっちがどっち」で覚える。

△「両方の意味を知っている」では足りません。入れ替えられた瞬間に、見分けられるか。……そこまでがが必要です。

試験ポイント | 短い一言に直す

- Temperatureは**ばらつき**、Top-pは**候補の幅**。
- プレトレーニングが**先**、ファインチューニングが**あと**。
- ★ これくらい短くしておくと、**本番で戻ってこられます**。

① やってみる② 逆にした選択肢の型

やってみる②

第2問 TemperatureとTop-pは、それぞれ何を決めているか。

ア Temperatureは出力のばらつきの強さを決め、Top-pは次の語の候補をどこまで広げるかを定める

イ Temperatureは次の語の候補をどこまで広げるかを決め、Top-pは出力のばらつきの強さを定める

答

ア。アとイは、**入れ替えただけ**です。

知らなければ五分五分。知っていれば一瞬。

⚠ ここで「どちらだったかな」と迷ったなら、それが弱点です。**両方の意味は言えるのに、対応が結びついていない**ということ。

⇒ その場で、短い一言に直してください。「Temperatureは、ばらつき」——それだけで次から迷いません。

6 道具3 | 「適切／不適切」を先に押さえる

これも、受験した方が書かれていました。

→ 『『適切／不適切』『正しい／誤り』が、連続すると頭がこんがらがる』

★ これは、知識の問題ではありません。読み方の事故です。

本番は60問。問題ごとに、聞き方が入れ替わります。「最も適切なものはどれか」→「適当でないものはどれか」→「誤っているものはどれか」。……3問連続で切り替わると、人は必ずつられます。

～について述べた次の記述のうち、
(長い設問がここに続く)

適当でないものはどれか。

↑ ここを
先に見る

探すのは「正しいほう」か「間違いのほう」か
分かっているのに逆を選ぶ、がいちばん惜しい

図5 分かっているのに逆を選ぶ、がいちばん惜しい

試験ポイント | 対策はひとつだけ

- ★ 問題文を読む前に、最後の一行を先に見る。
- 「適切」なのか、「適当でない」なのか。そこを先に確定させる。
- ⚠ 当たり前聞こえますが、本番でいちばん落とすのがここです。
- ★ こちらの練習問題でも、否定形は全体のおよそ2割入れています。慣れておいてください。

7 道具4 | 2択で止まったら、差分だけを見る

いちばん最初の話に戻ります。「2つまでは絞れる。でも決めきれない」。

やってはいけないこと

🚫 **2つの選択肢を、最初から読み直すこと**です。同じ結論に、もう一度たどり着くだけ。……時間だけが減ります。

やること

★ やることは、逆。2つの、**違うところだけを取り出す**。

残った2つは、ほとんど同じ文のはずです。違うのは**1か所か、2か所**。★ その1か所だけを見て、どちらが正しいかを決める。問題全体ではなく、その一語で決めます。

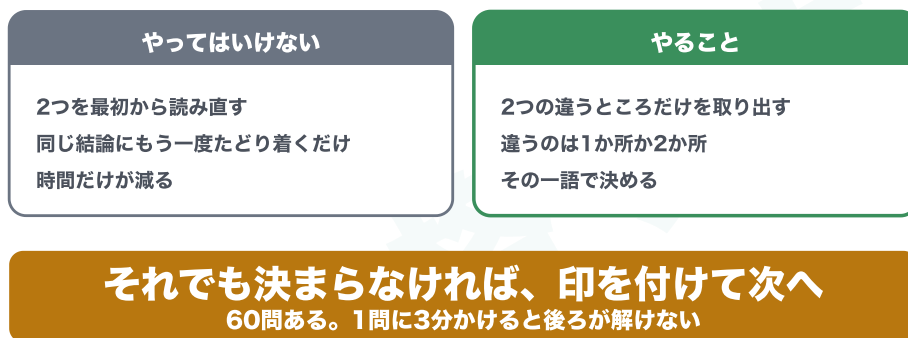


図6 それでも決まらなければ、印を付けて次へ

試験ポイント | 「決めない」も道具

- ⚠️ それでも決まらないときは、**捨ててください**。印を付けて、次へ。
- 60問あります。1問に3分かけると、**後ろが解けません**。
- ★ 「決めきれない」の正解は、「決めない」こともある。——これも道具です。

① やってみる③ 2択で止まる型

まず★最後の一行。「適当でない」ですね。探すのは、間違いのほうです。

やってみる③

第3問 AIエージェントの業務での使い方について述べた次の記述のうち、適当でないものはどれか。

ア 手順を自ら組み立てて進められるため、外部へ影響する操作についても人の承認を挟む必要はない

イ 目標を与えると、手順を計画し、道具を使い分けながら作業を進めることができる

答

ア。目印は「必要はない」＝言い過ぎ。

★差分を見ると、こうなります。イは「できる」と言っているだけ。アは「だから、しなくてよい」まで踏み込んでいる。

⚠ できることと、していいことは、別です。この線引きは、第4章でも第5章でも、ずっと同じでした。

8 言い回し耐性 | 1つの論点を3通りで

最後に、2つ目の難しさに戻ります。**言い回しが変わると、解けなくなる。**

△ ☆ これは、道具では直りません。**練習の形を変えるしかありません。**方法はひとつです——**同じ論点を、違う聞き方で、3回問う。**

同じ論点=プレトレーニングとファインチューニング

1回目・定義で問う

大量のデータであらかじめ
学習させることを何という？

2回目・順序で問う

先に行うのは、
どちらか？

3回目・否定形で問う

これらの説明として、
適当でないものはどれか？

3つとも正解して、はじめて「分かった」

図7 3つとも正解して、はじめて「分かった」

☆ 3つとも正解して、はじめて「分かった」。……**1つでも落ちたら、覚えていただけ**です。

ここが、丸暗記との分かれ目

同じ問題文を10回解いても、この差は見つかりません。☆ **聞き方を変えて、初めて見つかります。**

試験ポイント | サイトの「論点」を使う

- ☆ こちらのサイトでは、この束ねを「**論点**」と呼んで**40件**用意しました。
- 同じ論点の問題は、**必ず違う聞き方**にしてあります（定義／順序／否定形など）。
- ⚠ 同じ論点の中に、**同じ選択肢の組み合わせは置いていません**。記憶では解けないようにしてあります。

まとめ

今日の4つの道具

1 言い過ぎを探す
必ず・すべて・完全に・一切・保証される・常に

2 逆にした選択肢を疑う
2つセットの言葉は入れ替えて出る

3 「適切／不適切」を先に押さえる
問題文より先に、最後の一行を見る

4 2択で止まったら、差分だけを見る
読み直さない。違う一語で決める

+ 同じ論点を、違う聞き方で3回問う
これだけは道具ではなく習慣

図8 道具は4つ。+ 習慣がひとつ

正直に言うておきます

△ これで必ず受かる、とは言えません。言えるのは、**もったいない落とし方が減る**、というところまでです。

★ でも、**2つまで絞れているなら——あと一歩**です。

このあとの進め方

1. ⚠ 先に基本を一周してから。**土台がないうちに難問だけ解いても、削られるだけ**です。
2. 章のページで「応用・難問だけ」を選ぶ。全5章で**69問**あります。
3. 間違えた問題は、解説から動画の該当場面へ直接飛べます。
4. ★ 最後は**間違えた問題に戻るだけ**。それがいちばん効きます。

巻末 **本番の読み方（手順）**

試験当日、1問ごとにこの順で読んでください。**順番に意味があります。**

	やること	なぜ
1	最後の一行を先に見る	「適切」か「適当でない」かを確定させる
2	選択肢に言い過ぎが無いか目を留める	必ず・すべて・完全に・一切・保証される・常に
3	2つセットの言葉が出たら、入れ替えを疑う	もっともらしい誤答の作り方がこれ
4	2つに絞れたら、違う所だけを取り出す	読み直さない。その一語で決める
5	それでも決まらなければ、印を付けて次へ	60問ある。1問に3分かけない

⇒ **1と5がいちばん効きます。**1は落とさずに済む問題を守り、5は解けるはずの問題を守ります。

巻末 4つの道具チェックリスト

自分の言葉で言えたら✓を入れてください。⚠️「知っている」ではなく「**本番で使える**」かどうかで判断してください。

- ☐ 言い過ぎを探す（目印の6語を言える）
- ☐ 逆にした選択肢を疑う（2つセットの言葉を短い一言で言える）
- ☐ 「適切／不適切」を先に押さえる（最後の一行から読む）
- ☐ 差分だけを見る（読み直さない。決まらなければ飛ばす）

★道具ではなく習慣です。ここができると、2つ目の難しさが消えます。

- ☐ 同じ論点を、違う聞き方で3回問う（定義・順序・否定形）
- ☐ 3つとも正解して、はじめて「分かった」とする
- ☐ 1つでも落ちたら「覚えていただけ」と考え直す

この本で扱わなかったこと

- ️🚫 **出題の予想はしていません。** 根拠が無いからです。
- ️🚫 **「過去問」は載せていません。** **公開されていないから**です。
- ️★扱ったのは**解き方だけ**。知識は第1章から第5章の講義と教科書にあります。

第3部

総仕上げ

耳と紙で、194語を確かめる

AI資格ラボ

一問一答194語（聞き流し版）

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たかが分かるほうが、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。

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

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。

- | | | |
|---|-------------------------------------|----------------------------------|
| ☐ シンギュラリティ（技術的特異点） | ☐ ヴァーナー・ヴィンジ | ☐ 2045年問題 |
| | ☐ レイ・カーツワイル | ☐ 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。

AI資格ラボ

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。

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

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。

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

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

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。

- | | | |
|----------------------------------|--|----------------------------------|
| ☐ 画像のリサイズ | ☐ データの水増し（augmentation） | ☐ リマスタリング |
| ☐ 正規化 | ☐ データ拡張技術 | ☐ 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。

- | | |
|---|--|
| ☐ ディープフェイク（深層偽造）技術 | ☐ 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。

AI資格ラボ

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。

- | | | |
|---------------------------------------|--|----------------------------------|
| ☐ インターネットリテラシー | ☐ 情報リテラシー | ☐ デジタル市民権 |
| ☐ テクノロジーの理解 | ☐ 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。

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

- | | | |
|------------------------------------|------------------------------------|---------------------------------------|
| ☐ 個人情報保護法 | ☐ 個人情報取扱事業者 | ☐ 機微（センシティブ）情報 |
| ☐ 改正個人情報保護法 | ☐ 個人識別符号 | ☐ 匿名加工情報 |
| ☐ 個人情報保護委員会 | ☐ 要配慮個人情報 | ☐ 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。

- | | | |
|--------------------------------|----------------------------------|----------------------------------|
| ☐ 知的財産権 | ☐ 意匠権 | ☐ 営業秘密 |
| ☐ 著作権 | ☐ 肖像権 | ☐ 限定提供データ |
| ☐ 特許権 | ☐ パブリシティ権 | ☐ 著作権侵害 |
| ☐ 商標権 | ☐ 不正競争防止法 | ☐ 名誉毀損 |
| ☐ 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。

- | | | |
|---|---|---------------------------------------|
| ☐ 人間の尊厳が尊重される社会 (Dignity) | ☐ 持続可能な社会 (Sustainability) | ☐ 透明性 |
| ☐ 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) | ☐ 人間中心の考え方 | ☐ アカウンタビリティ |
| ☐ 教育・リテラシー | ☐ 安全性・公平性 | ☐ 人間中心 |
| ☐ 公正競争確保 | ☐ プライバシー保護 | ☐ 安全性 |
| | ☐ セキュリティ確保 | ☐ 公平性 |
| | ☐ イノベーション | ☐ AIガバナンス・ゴール |
| | ☐ 環境・リスク分析 | ☐ 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。

AI資格ラボ

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。

- | | | |
|-------------------------------------|--------------------------------------|--|
| ☐ LM | ☐ プレトレーニング | ☐ プロンプト |
| ☐ n-gramモデル | ☐ ハイパーパラメータ | ☐ プロンプトエンジニアリング |
| ☐ ニューラル言語モデル | ☐ Temperature | |
| ☐ LLM | ☐ 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定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。

- | | | |
|--------------------------------------|---|---|
| ☐ Instruction | ☐ Input Data | ☐ Zero-Shot プロンプティング |
| ☐ Context | ☐ Output Indicator | ☐ 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.8

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.8

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.10

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.11

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.11

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.209

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.209

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.209

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.209

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.209

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.209

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.13

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.209

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.209

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.210

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.210

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.14

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.210

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.210

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.15

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.210

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.210

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.17

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.15

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.20

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.20

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.211

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.211

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.211

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.211

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.22

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.22

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.212

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.22

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.24

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.212

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.212

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.24

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.25

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.36

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.38

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.38

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.39

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.41

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.41

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.42

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.43

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.42

入力を圧縮する側 (入口)。

デコーダ p.42

圧縮されたものから復元する側 (出口)。

潜在ベクトル p.42

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.43

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.43

GANの中で偽物を作る側。

識別器 p.43

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.45

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.45

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.45

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.45

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.46

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.46

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.47

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.47

同じ文の中の語どうして注目し合う仕組み。

Attention Mechanism p.48

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.48

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.47

モデルの構造・設計そのもの。

GPTモデル p.51

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.51

GPTモデルを開発した組織。

BERTモデル p.51

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.52

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.52

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 p.52

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.52

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.63

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.64

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.64

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.65

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.65

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.65

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.67

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.68

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.71

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.65

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.68

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.68

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.69

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.69

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.71

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 p.72

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.72

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.72

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.76

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.73

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.73

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.73

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.73

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.73

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.75

OpenAIの動画生成AI。文章で指示すると映像が出てくる。

Operator p.75

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.75

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.75

OpenAIの画像生成の機能。文章で指示して絵を作る。

Gemini p.76

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.76

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.86

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.86

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.87

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.87

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.87

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.89

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.90

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.90

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.92

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.94

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.94

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.96

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.97

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.97

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.97

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.97

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.108

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.108

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.108

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.109

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.109

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.110

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.110

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.110

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.116

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.116

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.116

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.112

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.112

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.113

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.113

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.113

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.119

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.119

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.119

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.119

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.121

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.121

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.121

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.124

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.124

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.136

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.136

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.136

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.137

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.137

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.138

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.138

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.141

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.141

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.142

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.143

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.144

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.143

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.146

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.146

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.146

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.147

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.147

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.147

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.147

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.147

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.147

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.150

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.150

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.155

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.150

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.150

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.150

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.150

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.153

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.153

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.153

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.154

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.154

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.154

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.155

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.166

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.166

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.167

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.167

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.168

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.169

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.169

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.170

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.172

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.172

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.173

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.173

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.173

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.173

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.175

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.175

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。

総索引

公式シラバスの詳細キーワード全194語です。並びはシラバスの順（章→節）。

AI資格ラボ

第1章 (39語)

AI (人工知能) p.8

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.8

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.10

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.11

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.11

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.209

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 p.209

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.209

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.209

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.209

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.209

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 p.13

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.209

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.209

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.210

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.210

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.14

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.210

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.210

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.15

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.210

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.210

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.17

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.15

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.20

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.20

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.211

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.211

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.211

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.211

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.22

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.22

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.212

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.22

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.24

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.212

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.212

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.24

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.25

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 (62語)

生成AI (ジェネレーティブAI) p.36

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.38

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.38

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.39

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.41

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.41

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.42

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.43

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.42

入力を圧縮する側 (入口)。

デコーダ p.42

圧縮されたものから復元する側 (出口)。

潜在ベクトル p.42

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.43

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.43

GANの中で偽物を作る側。

識別器 p.43

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.45

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.45

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.45

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.45

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.46

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.46

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.47

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.47

同じ文の中の語どうして注目し合う仕組み。

Attention Mechanism p.48

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.48

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.47

モデルの構造・設計そのもの。

GPTモデル p.51

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.51

GPTモデルを開発した組織。

BERTモデル p.51

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.52

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.52

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 p.52

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.52

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.63

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.64

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.64

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.65

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.65

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.65

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.67

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.68

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.71

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.65

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.68

人間のフィードバックによる強化学習。人がよし悪しを選び、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.68

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.69

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.69

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.71

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 p.72

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.72

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.72

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.76

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.73

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.73

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.73

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.73

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.73

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.75

OpenAIの動画生成AI。文章で指示すると映像が出てくる。

Operator p.75

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.75

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.75

OpenAIの画像生成の機能。文章で指示して絵を作る。

Gemini p.76

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.76

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 (16語)

画像のリサイズ p.86

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.86

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.87

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.87

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.87

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.89

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.90

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.90

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.92

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.94

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.94

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.96

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.97

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.97

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.97

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.97

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 (61語)

インターネットリテラシー p.108

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.108

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.108

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.109

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.109

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.110

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.110

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.110

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.116

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.116

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.116

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.112

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.112

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.113

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.113

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.113

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.119

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.119

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.119

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.119

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.121

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.121

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.121

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.124

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.124

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.136

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.136

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.136

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.137

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.137

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.138

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.138

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.141

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.141

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.142

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.143

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.144

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.143

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.146

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.146

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.146

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.147

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.147

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.147

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.147

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.147

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.147

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.150

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.150

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.155

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.150

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.150

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.150

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.150

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.153

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.153

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.153

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.154

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.154

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.154

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.155

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 (16語)

LM p.166

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.166

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.167

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.167

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.168

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.169

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.169

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.170

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.172

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.172

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.173

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.173

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.173

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.173

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.175

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.175

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。

奥付

本書について

『生成AIパスポート 完全攻略テキスト』

発行：AI資格ラボ（YouTubeチャンネル）

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

公式シラバス <https://guga.or.jp/assets/syllabus.pdf>

⚠️ 本書は受験対策の解説であり、試験の実施団体とは関係ありません。掲載している問題はすべて自作の予想問題で、過去問ではありません。内容は発行時点のシラバスに基づきます。最新情報は必ず公式をご確認ください。

AI資格ラボ