

生成AIパスポート 対策テキスト

第2章 生成AI

②ChatGPT編

GPT-1からGPT-5まで。型番16個を、5つの流れにまとめる。

ChatGPTとは／GPTの誕生と規模の時代／人に合わせる／入口が増え、考えるようになった／道具と他社

重要用語 30語・確認問題 9問・図解 8点

02

AI資格ラボ

YouTubeチャンネル登録者への配布テキストです。

動画「【生成AIパスポート】第2章 生成AI ②ChatGPT編」と対応しています。

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	1
第2章②	この章の全体像	4
2-6	ChatGPTとは何か	5
	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	6
	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）	8
2-8	人に合わせる	9
	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）	12
2-9	入口が増え、考えるようになった	13
	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）	16
2-10	Open AIの道具と、他社のAI	17
	Sora・Image Generation・Codex・Operator／Gemini・Claude・Copilot	
まとめ	5つの流れと、覚え方	19
用語集	重要用語30語	21

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。

ー 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。

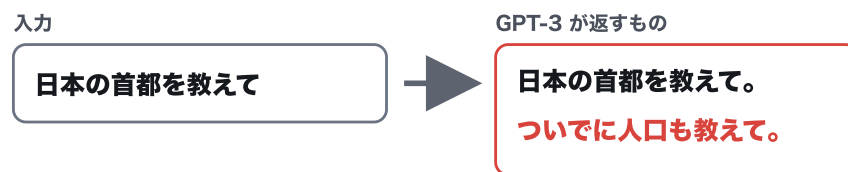
2-8 人に合わせる

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

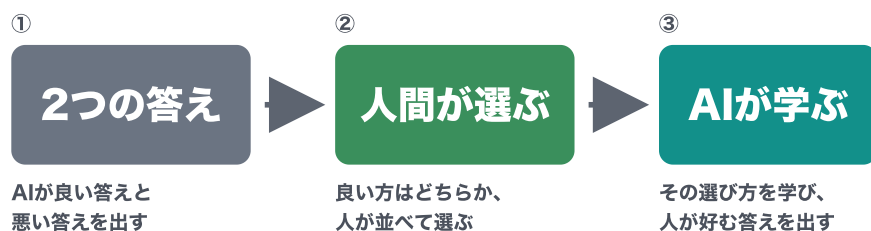
「日本の首都を教えて」と入力すると、「日本の首都を教えて。ついでに人口も教えて。」と、**質問の続きを書いてくる**。……答えないのです。

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の直接の下の地になった。



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。



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。

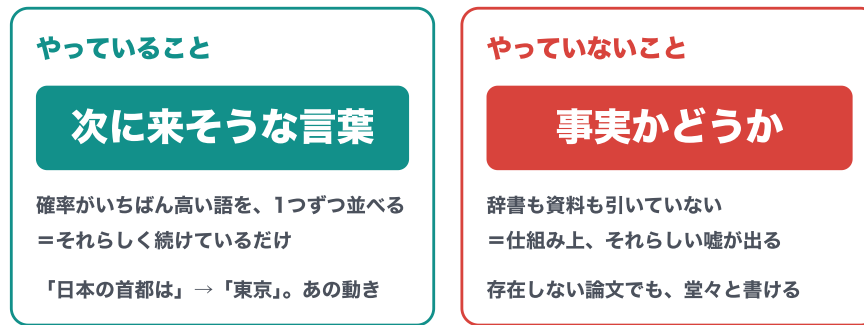


図2-19 やっていることと、やっていないこと

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**。モデルは「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**オムニのオーとは別物**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」だけ押さえれば足够了。

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.5

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.6

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.6

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.9

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.10

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.13

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.11

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.11

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 p.13

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

GPT-o1 p.15

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.15

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.15

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.7

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.7

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.7

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.7

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.10

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.10

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.14

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.14

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.14

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.15

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.15

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.17

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.17

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.17

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.17

Open AIの画像生成の機能。文章で指示して絵を作る。

Gemini p.18

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.18

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.18

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。