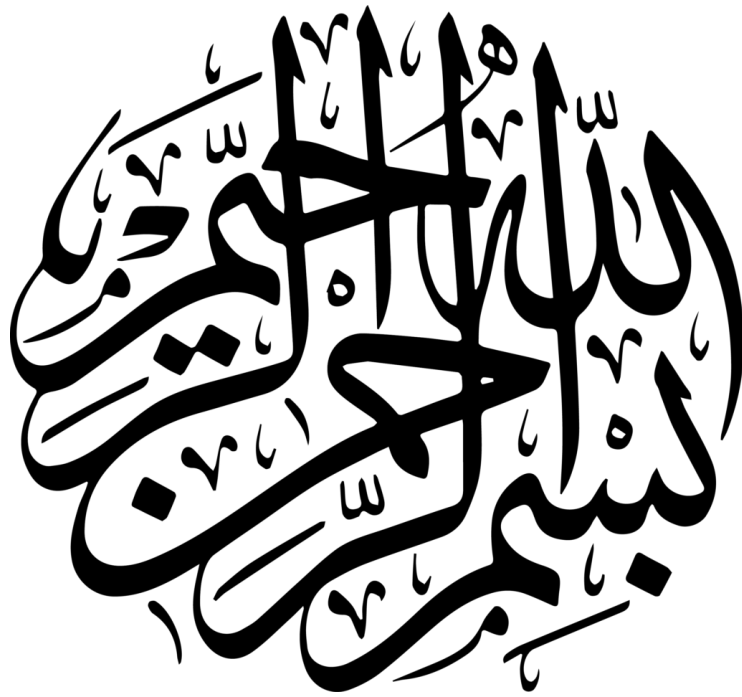


الآلة التي تتفاعل

أسرار التواصل الرقمي

جميع الحقوق محفوظة
٢٠٢٥

إعداد
د. عبد الرحمن الزراعي





قائمة الموضوعات

٦	الفصل الأول: بناء الشخصية الرقمية – كيف تتشكل GPT؟
٦	توطئة:
٦	اسم النموذج وتطوره عبر الأجيال
٦	نبذة:
٨	تخصيص النموذج: كيف يتعامل GPT مع اللغة العربية؟
٨	نبذة:
٩	كيف يتعلم النموذج؟ من الأنماط الإحصائية إلى تفضيلات البشر
٩	نبذة:
١٠	من المكافأة إلى الرقابة: كيف يُراجع النموذج نفسه؟
١٠	نبذة:
١٣	حدود الذاكرة والخصوصية في GPT
١٣	نبذة:
١٥	جدول ملخص
١٦	خاتمة
١٧	الفصل الثاني: البنية الذهنية – كيف يفكر النموذج من الداخل؟
١٧	توطئة:
١٧	تحويل الكلمات إلى أرقام: من المتجهات إلى الفضاء الرياضي
١٧	نبذة:
١٨	الانتباه الذاتي والرؤوس المتعددة: كيف يفهم النموذج المعاني الدقيقة؟
١٨	نبذة:
٢٠	التعلم والتحسين: من الشبكة الداخلية إلى التدرج العكسي
٢٠	نبذة:
٢٢	جدول ملخص
٢٣	الفصل الثالث: تعليمات التوجيه – كيف نضبط سلوك GPT؟
٢٣	توطئة:
٢٣	واجهة تكوين GPT – ماذا يوجد خلف زر "تعديل"؟
٢٣	نبذة:
٢٦	خاتمة
٢٦	الذاكرة – هل يتذكر GPT من تكلم معه؟
٢٦	نبذة:
٢٨	خاتمة
٢٩	لوحة الموجهات – كيف تتحكم بتصرفات GPT في واجهة واحدة؟
٢٩	توطئة:
٣٠	خاتمة
٣٣	الفصل الرابع: الذاكرة والسلوك – من التعليمات إلى الاستجابة
٣٣	توطئة:
٣٣	تعليمات النظام – من يوجه GPT؟
٣٣	نبذة:
٣٦	التعليمات المخفية Hidden Instructions
٣٦	نبذة:



٣٩.....	جدول ملخص: مقارنة تعليمات النظام والتعليمات المخفية.	📊
٣٩.....	خاتمة:	← END
٣٩.....	التعليمات المضمّنة – كيف يتكوّن السلوك من الداخل؟	🔍
٣٩.....	نبذة:	📌
٤١.....	جدول ملخص: التعليمات المضمّنة.	📊
٤٢.....	خاتمة:	← END
٤٢.....	سياسة الأمان: كيف تحمي المنصة النموذج والمستخدم؟	🔍
٤٢.....	نبذة:	📌
٤٥.....	خاتمة:	← END
٤٥.....	تعليمات السياق – ما الذي يكتبه المستخدم فعلاً؟	🔍
٤٥.....	نبذة:	📌
٤٧.....	ملخص: تعليمات السياق.	📊
٤٧.....	خاتمة:	← END
٤٧.....	تعليمات التخصيص: كيف تُخصّص GPT من الداخل؟	🔍
٤٧.....	نبذة:	📌
٥٠.....	جدول ملخص: التعليمات الخاصة.	📊
٥٢.....	ملخص جدول: تعليمات تخصيص الشخصية.	📊
٥٢.....	خاتمة:	← END
٥٣.....	تعليمات GPT المخصّص: كيف تبني شخصية تفاعلية؟	🔍
٥٣.....	نبذة:	📌
٥٦.....	جدول ملخص: النماذج المخصّصة.	📊
٥٦.....	خاتمة:	← END
٥٦.....	تضمين الملفات – كيف تعطي GPT ذاكرة معرفية؟	🔍
٥٦.....	نبذة:	📌
٥٩.....	جدول ملخص: تضمين الملفات.	📊
٥٩.....	خاتمة:	← END
٥٩.....	الأوامر الخارجية – كيف ينفّذ GPT مهاماً فعلية؟	🔍
٥٩.....	نبذة:	📌
٦١.....	جدول ملخص: الأوامر الخارجية.	📊
٦١.....	خاتمة:	← END
٦٢.....	وضع الوكيل – كيف ينفّذ GPT مهاماً ذاتية؟	🔍
٦٢.....	نبذة:	📌
٦٤.....	ملخص التعليمات.	🔍
٦٤.....	نبذة:	📌
٦٥.....	خاتمة:	← END
٦٦.....	الفصل الخامس: اللغة والأسلوب – كيف تتكيف GPT مع المتحدث؟	🔍
٦٦.....	توطئة:	📌
٦٦.....	○ أولاً: تغيير النبرة – هل يمكن أن تتحدث كما أشاء؟	🔍
٦٦.....	نبذة:	📌
٦٦.....	○ ثانياً: محاكاة الأدوار – كيف تتقمّص دور المعلم أو الشاعر؟	🔍
٦٦.....	نبذة:	📌
٦٧.....	○ ثالثاً: تحليل الأسلوب – كيف أعرف ما يناسبني؟	🔍
٦٧.....	خاتمة:	← END
٦٨.....	هل يمكن أن يمتلك GPT شخصية فريدة؟	🔍
٦٨.....	نبذة:	📌



- ٦٨..... هل نُقلدني أم تفهمني؟
- ٦٨..... نبذة GPT: *
- ٦٩..... الفرق بين النموذج العام والمخصص.
- ٦٩..... نبذة: *
- ٧٠..... خاتمة: ← END



✦ الفصل الأول: بناء الشخصية الرقمية – كيف تتشكل GPT؟

🌟 توطئة:

في هذا الفصل نفتح نافذة على عقل نموذج GPT: كيف وُلد وتطوّر، وكيف تعلّم اللغة، وكيف تدرب على ملايين النصوص، حتى غدا قادرًا على توليد الجمل وفهم السياق ومحاكاة الأساليب. نبدأ من معنى اسمه، ثم ننتقل إلى مراحل التاريخيّة، ونتعرّف على مدى إتقانه للعربية، وكيف يمكن تخصيصه لمجالاتٍ محددة. ونكشف كذلك مراحل التحسين بعد التدريب، كالتعليم المعزّز بتفضيلات البشر، وآليات المكافأة والعقوبة، ثم نفهم حدود ذاكرته داخل الجلسة وخارجها، وما علاقة كل ذلك بخصوصية المستخدم وسلامة المحتوى. إن فهم هذه الرحلة من الداخل ليس مجرد معرفة تقنية، بل خطوة أولى لفهم كيف تُبنى الشخصية الرقمية التي نحاورها اليوم.

اسم النموذج وتطوره عبر الأجيال

🌟 نبذة:

قبل الدخول إلى أعماق شخصية GPT وآليات تفكيره، لا بدّ من التوقّف عند اسمه وأصله، ومعرفة مراحل تطوره عبر السنين. في هذا الفصل، يحاور الطالبُ النموذجَ حول معنى "GPT"، وما الفرق بين النموذج التوليدي والنموذج التنبئي، ثم ينتقل إلى استعراض زمني لأجيال GPT من ١ حتى ٥٤، مع التركيز على ما أضافه كل جيل ولماذا يُعدّ خطوة فارقة.

🎓 الطالب: قبل الخوض في التفاصيل، أريد أن أفهم اسم "GPT" نفسه؛ ما أصل هذه الأحرف؟



ChatGPT: هو اختصار لعبارة المحوّل التوليدي المدرّب مسبقًا Generative Pre-trained Transformer

الطالب: وكيف يختلف "التوليدي" عن "التنبئي" الذي كنت أسمع به؟
ChatGPT: النموذج التنبئي يصنّف أو يترجم نصًا موجودًا، أما التوليدي فيبتكر نصًا جديدًا لا نسخة مطابقة لما تعلّمه.

الطالب: الآن بعدما فهمت الاسم، كيف بدأ تاريخ هذه النماذج؟
ChatGPT: صدر GPT-1 عام ٢٠١٨ بقدرة ١١٧ مليون معلّمة Parameter، مؤكّدًا جدوى الفكرة.

الطالب: وماذا أضاف GPT-2 بعد عام؟
ChatGPT: رفع السعة إلى ١.٥ مليار معلّمة، فصارت الكتابة أطول وأفصح، غير أنّ نشره تأخّر خشيةً من إساءة الاستخدام، فوضعت سياسات مسؤولة قبل إتاحتها.

الطالب: وبعد ذلك أتى GPT-3 بثورة عددية، أليس كذلك؟
ChatGPT: وصل عددُ معلّماته إلى ١٧٥ مليارًا، فأصبح يكتب نصوصًا تشبه البشر، ثم قدّمته الشركة عبر واجهة برمجية (API) كي يستخدمه المطوّرون في تطبيقات تجارية مقابل رسوم.

الطالب: وما أهمية إصدار GPT-3.5؟
ChatGPT: كان جسرًا نحو GPT-4، فقد حسّن فهم الشيفرات، وخفّف الأخطاء المنطقية.

الطالب: وما المقصود بالشيفرات هنا؟ هل تعني تعليمات البرمجة؟
ChatGPT: نعم؛ الشيفرات هي الأكواد البرمجية مثل Python أو JavaScript، وقد أصبح النموذج أقدر على قراءتها وشرحها وتصحيح أخطائها.



🎓 الطالب: ثم ظهر GPT-4... فكيف تفوّق؟

🤖 ChatGPT: نعم كان تفوّقه بسبب دقة أعلى في الاستدلال، وقدرة أفضل على تعدّد اللغات، مع نوافذ سياق أطول.

🎓 الطالب: وأخيراً سمعنا عن GPT-4o متعدد الوسائط، ما سرّ حرف o؟

🤖 ChatGPT: يشير إلى omni، أي الجمع بين النص والصوت والصورة في عقل واحد.

تخصيص النموذج: كيف يتعامل GPT مع اللغة العربية؟



نُبذة:

رغم أن نماذج GPT تدعم لغات عديدة، إلا أن مدى فهمها للغة العربية يعتمد على جودة التدريب ودرجة التخصيص.

في هذا المقطع، يستعرض الطالب مع ChatGPT كيف يمكن تعزيز استيعاب النموذج للغة العربية – من خلال تزويده ببيانات متخصصة – ليُتقن الأساليب الفقهية، التحريرية، أو حتى المحكية المحلية.

🎓 الطالب: بعد هذه الجولة الزمنية، كيف يتعامل النموذج مع النص العربي تحديداً؟

🤖 ChatGPT: يعتمد على حجم ونوعية النص العربي في تدريبه؛ وكلما زادت أمثلة التخصيص زادت فصاحته.

🎓 الطالب: وما أمثلة هذه التخصيصات؟ أعطني شيئاً ملموساً.

🤖 ChatGPT: مثل أن نزوّده بمجموعة فتاوى مُحَقَّقة، فيتعلّم مفردات الفقه ويُفتي بلغة دقيقة؛ أو نُعطيه أرشيف مقالات جريدة عربية، فيقلّد أسلوبها التحريري؛ أو نُدرّبه على سجلات خدمة عملاء لمطعم محليّ، فيردّ على الزبائن بلهجة المدينة وأسماء أطباقها.

🎓 الطالب: وضح أكثر مفهوم التخصيص.

🤖 ChatGPT: هو إعادة تدريب النموذج على نصوص متخصصة، كتدريبه على الفقه الإسلامي ليُتقن مصطلحاته.



الطالب: هل يمكن صنع نسخة عربية خالصة؟

ChatGPT: ممكن نظريًا، لكنه مكلف ماليًا؛ والحل العملي هو التخصيص (Fine-Tuning)

لنموذج متعدّد اللغات.

كيف يتعلّم النموذج؟ من الأنماط الإحصائية إلى تفضيلات البشر



نبذة:

قد يتساءل القارئ: كيف يُصبح GPT قادرًا على الكتابة بهذه الدقة؟

السّر في مراحل التعلّم المتعدّدة التي يمرّ بها:

يبدأ بتعلّم الأنماط الإحصائية من الإنترنت (Pre-training)،

ثم يُدرّب على أمثلة بشرية (Supervised Learning)،

ويُصقل بالتعليم المعزّز بتفضيلات البشر (RLHF)،

لكنه لا يتعلّم منك أثناء استخدامه، حفاظًا على الخصوصية.

في هذا الحوار، يتدرّج الطالب مع النموذج في فهم كل مرحلة بلغة سهلة وأمثلة حيّة.

الطالب: كيف يتعلّم النموذج؟

ChatGPT: يبدأ بتدريب مسبق على مليارات الكلمات من الإنترنت، فيتعلّم الأنماط الإحصائية.

الطالب: وماذا تعني "الأنماط الإحصائية"؟

ChatGPT: المقصود أنني لا "أفهم" الكلمات كما يفعل الإنسان، بل ألاحظ أن بعض الكلمات

تتكرّر مع كلمات أخرى كثيرًا، فأحاول التنبؤ بالكلمة التالية بناءً على ما رأيت في ملايين الجُمْل من قبل.

فمثلاً، إذا كتبت: "السماء زرقاء"، فأنا لا أدرك معنى "السماء" ولا معنى "الأزرق"، لكنني رأيت كثيرًا أن

"زرقاء" تأتي بعد "السماء"، فأقترحها بنسبة عالية..

الطالب: وهل يتوقف تعليم النموذج عند هذا الحد؟

ChatGPT: لا، فبعد التدريب المبدئي، يُمرّن النموذج بأسلوب يُسمّى "التعليم الخاضع للإشراف"

حيث يُزوّد بأسئلة وأجوبة صحيحة مسبقًا، ليتعلّم منها شكل الردود الدقيقة.



الطالب: هل هناك مراحل تحسّن أخرى؟

ChatGPT: نعم، ومنها التعليم المعزّز بتفضيلات البشر (RLHF)، وهو ما يصقل الأسلوب ويقلّل الانحرافات في الإجابات.

الطالب: وما معنى تعليم معزّز هنا؟

ChatGPT: بعد جمع تلك التفضيلات نستخدم خوارزمية المكافأة؛ فإذا كتب النموذج إجابة تشبه الأعلى تقديرًا فمنحه مكافأة رياضية، وإن ابتعد نعطيه عقوبة، فيتعلّم بالتجربة والتحسين التدريجي، كما يتعلّم الطفل ركوب الدراجة مع التشجيع والتصحيح.

الطالب: وكيف يتم هذا التعليم المعزّز؟

ChatGPT: يقوم المدربون بصناعة عدّة إجابات عن سؤال واحد، ثم يُصنّف المقومون هذه الإجابات من الأفضل إلى الأسوأ، ليتعلّم النموذج من هذا الترتيب ما يُرضي المستخدمين أكثر.

الطالب: أعني: هل يستمرّ تعلّم النموذج أثناء استخدامنا له؟

ChatGPT: لا، بعد النشر لا يتعلّم من المستخدمين مباشرةً، حفاظًا على الخصوصية والأمان، لكن المطوّرين قد يستخدمون بيانات مجهولة لتحسين الإصدارات القادمة.

من المكافأة إلى الرقابة: كيف يُراجع النموذج نفسه؟



نبذة:

لا يقتصر تعلّم GPT على ملاحظة الأنماط، بل يتعلّم من التقييم. في هذا الفصل، نتعرف على:

كيف يُمنح النموذج مكافآت أو عقوبات رقمية؟

ما الفرق بين "نموذج المكافأة" و"خوارزمية التعزيز"؟

كيف يراقب GPT نفسه داخليًا (عبر الطبقات)؟

وكيف يُراقب خارجيًا (عبر المرشحات)؟

ما معنى أن تُشدّد أو تُخفّف الرقابة على النص؟

هذا الفصل يكشف كيف يتحوّل النموذج من مجرد متعلّم، إلى "كاتب مسؤول"،



يُضبط عباراته، ويحرص على ألا ينتهك الأخلاقيات، أو يخرج عن سياسات الاستخدام.

الطالب: ذكرت أنه بعد جمع تلك التفضيلات نستخدم خوارزمية المكافأة؛ فكيف تعمل هذه الخوارزمية؟

ChatGPT: أولاً، يُنشئ العلماء نموذجًا صغيرًا يُدعى "نموذج المكافأة"، وظيفته تقييم الإجابات، فيعطي كل واحدة رقمًا من -1 إلى +1 لكل إجابة، بحسب تشابهها مع أفضل إجابة اختارها البشر.

الطالب: وما معنى الأرقام الموجبة والسالبة؟

ChatGPT:

الرقم +1 يعني "أحسنت! هذه إجابة ممتازة"، فهي مكافأة، والرقم -1 يعني "هذه الإجابة سيئة أو غير مناسبة"، فهي عقوبة أما الأرقام بينهما، مثل (صفر)، فتعني "إجابة متوسطة أو مقبولة جزئيًا".

الطالب: وكيف أستخدم هذا التقييم؟

ChatGPT: يُعدّل أوزانه الداخلية بخوارزمية التعزيز؛ فيقوّي المسارات التي تؤدي إلى مكافأة، ويضعف المسارات التي تجلب عقوبة،

تمامًا مثل طفل يتعلّم ركوب الدراجة: كلما شجّعته المدرّب على حركة صحيحة، كررها؛ وإذا نبّهه إلى خطأ، اجتنبه في المرات القادمة.

الطالب: وهل بمقدوره تفادي الخطأ تمامًا؟

ChatGPT: لا، لكنه يقلّل احتماله، ويجعل اللغة أكثر توافقًا مع قيم المجتمع والمعايير الإنسانية والأخلاقية.

الطالب: كيف تتحققون من سلامة المخرجات أو سلامة ما أكتبه؟

ChatGPT: هناك ثلاث طرق رئيسية:

١. مرشّحات آلية تراجع النصوص قبل عرضها.
٢. مراجعة بشرية من الخبراء، خاصةً في الحالات الحساسة.



٣. مقارنة النص بسياسات المنصة، للتأكد من عدم مخالفته للقوانين أو الأخلاق العامة.

الطالب: وهل هذه المرشحات هي نفسها طبقات النموذج؟

ChatGPT : لا، بل هناك فرق كبير بينهما.

الطبقات هي أجزاء داخلية في "عقلي"، تعمل لحظةً بلحظة أثناء تأليف الجملة. أما المرشحات، فهي حارس خارجي يُراجع الجملة بعد أن أكتبها كاملة، ويقرر إن كانت مناسبة للعرض أو لا.

الطالب: يمكنك أن تلخص لي الفرق بين الطبقات والمرشحات في نقاط سريعة؟

ChatGPT : بالطبع؛ إليك مقارنة وصفية تُعينك على الفهم:

– الموقع والمكان:

الطبقات تعيش داخل "دماغ" النموذج وتعمل أثناء بناء الجملة، أما المرشحات فتقف خارجه بعد اكتمال الجملة.

– الوظيفة:

الطبقات تستخرج المعاني والعلاقات بين الكلمات، بينما المرشحات تفتش النص الناتج لتحرسه من المحتوى المخالف.

– وقت التشغيل:

تعمل الطبقات داخل النموذج مع كل كلمة يكتبها، أما المرشحات فهي خطوة لاحقة تُطبّق بعد اكتمال النص لتراجع دفعه واحدة

– أثرها المباشر:

الطبقات تؤثر في جودة الفهم والأسلوب، والمرشحات تؤثر في قبول النص أو حجبه أو تعديله.

– إمكانية الضبط:

يمكن زيادة الطبقات أو تقليلها لتغيير قدرات النموذج، ويمكن تشديد المرشحات أو تخفيفها حسب سياسات المحتوى.

المعيار	الطبقات (Layers)	المرشحات (Filters)
الموقع والمكان	توجد داخل "دماغ" النموذج وتعمل أثناء توليد الجملة.	تعمل خارجه، بعد اكتمال الجملة.



الوظيفة	تستخرج المعاني والعلاقات بين الكلمات داخل السياق.	تفتش النص الناتج لحمايته من المخالفات أو المحتوى المحظور.
وقت التشغيل	تُفَعَّل أثناء الكتابة، مع كل كلمة تُنتج.	تُطبَّق بعد اكتمال النص، كمراجعة نهائية شاملة.
الأثر المباشر	تؤثر في جودة الفهم، الأسلوب، والتماسك النصي.	تؤثر في قبول النص، حجه، أو تعديله قبل عرضه على المستخدم.
إمكانية الضبط	يمكن زيادتها أو تقليلها لتعديل قدرات النموذج الذهنية.	يمكن تشديدها أو تخفيفها حسب السياسات والضوابط الأخلاقية والمحتوى.

الطالب: وما المقصود بتشديد المرشحات أو تخفيفها؟

ChatGPT : تشديد المرشحات يعني جعلها أكثر صرامة؛ فنضيف كلماتٍ محظورة جديدة، أو نقلل الحدّ المسموح به لتعبيراتٍ معيّنة، فيُحجَب محتوى أكبر.
 أمّا تخفيفها فهو رفع هذا الحد أو حذف بعض الكلمات من القائمة السوداء، فيمرّ محتوى أوسع.
 يتم اختيار المستوى وفق سياسة كل منصة وما تريده من أمانٍ أو حرية.

حدود الذاكرة والخصوصية في GPT

نبذة:

تعمل نماذج GPT بذاكرة مؤقتة تحفظ ما يُكتب داخل الجلسة فقط، دون قدرة على التذكّر الطويل إلا إذا فُعِّلَت خاصية "الذاكرة" المضمّنة في بعض الإصدارات. ومع ذلك، يظل الحفاظ على الخصوصية أولوية، حيث تُفلتر البيانات ولا تُربط بالهوية الشخصية للمستخدم. هذا العنوان يستعرض كيف تُخزّن المعلومات مؤقتًا، وما حدود استخدام GPT للبيانات، وما الضوابط الأخلاقية والتقنية التي تحمي خصوصية المحادثات.

الطالب: هل معلوماتك دائماً محدثة؟

ChatGPT : لا، فأنا لا أتصل بالإنترنت مباشرة.

كل إجاباتي مبنية على ما تمّ تدريبي عليه حتى تاريخ معين.
 لذلك، قد تكون معلوماتي قديمة أحيانًا، خصوصًا في الأخبار أو التحديثات التقنية.



الطالب: إذن لا تعرف ما حدث بعد تاريخ تدريبيك؟

ChatGPT: صحيح. فإذا كان تدريبي انتهى في ٢٠٢٥ مثلاً، فأنا لا أعرف ما جرى بعدها. لكن يمكن للمستخدم أن يربطني بالإنترنت عن طريق تفعيل (أداة ويب) إذا أراد الوصول إلى معلومات حيّة.

الطالب: وما حدود "ذاكرتك" أثناء الحوار الواحد؟

ChatGPT: أتذكر ما يدور في الجلسة نفسها، أو داخل ما يسمى (نافذة السياق)، وإذا طالت المحادثة جداً، أضطرّ إلى نسيان بداياتها لأتابع الجزء الأحدث منها.

الطالب: وهل تتذكرني من جلسة لأخرى؟

ChatGPT: لا. كل جلسة جديدة أبدأ فيها من الصفر، إلا إذا فعلت خاصية "الذاكرة" في إعداداتك، ووافقت على أن أحتفظ ببعض المعلومات العامة عنك لتحسين الخدمة.

الطالب: وكيف أوافق أو أرفض ذلك؟

ChatGPT: في أول استخدام يظهر لك إشعار يسألك: هل توافق على مشاركة المحادثات لتطوير النموذج؟ ويمكنك تغيير قرارك في أي وقت من خلال الإعدادات.



من الإعدادات — عناصر التحكم في البيانات — تحسين النموذج للجميع — مفعل / غير مفعل (حسب اختيارك)



الطالب: ولماذا تحتاج الشركة إلى محادثاتي؟

ChatGPT: لتحسين جودة اللغة، خاصة في لهجاتك،

ولمعرفة أنواع الأسئلة التي يطرحها الناس، وتطوير ميزات أفضل في الإصدارات القادمة. كأنتك تزود المعلم بتغذية راجعة تساعد على تحسين دروسه.

الطالب: وهل في هذا مخاطرة؟

ChatGPT: الخطر الأكبر يكون في حال كتبت معلوماتك الشخصية، مثل رقم حسابك البنكي.

لذا أنصحك: لا تكتب أي شيء حساس، حتى لو كنت تثق في أمان المنصة.

الطالب: بعد كل هذا، ما النصيحة الذهبية في التعامل مع GPT؟

ChatGPT: عاملني كشريك مساعد، قدم لي سؤالك بوضوح،

ثم تحقق بنفسك من النتائج،

خاصة في المواضيع الدقيقة أو الحساسة أو الحديثة.

جدول ملخص

السؤال المحوري	الجواب المختصر
ما هو GPT ؟	محول توليدي مُدرَّب مسبقًا لإنتاج نص عربي/أجنبي متماسك
لماذا هو توليدي؟	لأنه يبتكر نصوصًا جديدة ولا يكرر نصًا محفوظًا
كيف تطوّر GPT	تضاعف عدد المعلومات، تحسّن الفصاحة، وأضيفت القدرة المتعددة الوسائط في 04
ما دور التخصيص؟	إعادة تدريب النموذج على نصوص متخصصة ليكتسب مصطلحات المجال ولهجته
ما هو RLHF ؟	تعليم معزّز بتفضيلات البشر يُكافئ الإجابات المرضية ويعاقب غيرها
الطبقات والمرشّحات؟	الطبقات تعمل داخل النموذج، أمّا المرشّحات فتراجع النص بعد اكتماله وتُحجّب المخالف

الجانب	الوصف
المعرفة العامة	تعتمد على تاريخ التدريب، ولا تُحدّث تلقائيًا
الذاكرة داخل الجلسة	تظل فعّالة ما دامت المحادثة مفتوحة، ثم تُنسى عند الإغلاق
الذاكرة الطويلة (اختيارية)	تُفعّل فقط إذا وافق المستخدم على الاحتفاظ بمعلومات عامة عنه
الخصوصية	لا يتم الاحتفاظ بمحتوى المحادثة إلا بموافقة واضحة



إدخال معلومات حساسة عن طريق الخطأ

المخاطر المحتملة

← خاتمة
END

ختمنا هذا الفصل بتعريف المصطلحات الأساسية ووضعها في سياقها الزمني. في الفصل التالي سنخوض في البنية الداخلية لنكتشف كيف يفكر GPT عندما يكتب بالعربية أو الإنجليزية.



✿ الفصل الثاني: البنية الذهنية – كيف يفكر النموذج من الداخل؟

توطئة

هذا الفصل يأخذ القارئ في رحلة داخل عقل النموذج، كاشفًا كيف تتحوّل الكلمات إلى متجهات رقمية، وكيف تتفاعل هذه المتجهات عبر طبقات الانتباه الذاتي، ثم كيف تُعاد معالجتها في شبكات خاصة، حتى تخرج في النهاية على هيئة جملة منسجمة، تبدو وكأنها نتاج عقل بشري. ولنكتشف كيف تتحوّل الكلمات إلى أرقام تُدعى "متجهات"، ولماذا يحتاج الحاسوب لهذا التمثيل الرقمي أصلاً. ثم ننتقل إلى فهم آلية "الانتباه الذاتي" التي تُشبه شبكة من العيون أو "الرؤوس" المتعددة، كلٌّ منها يركّز على زاوية معينة من اللغة، فنتنتج عنها استجابة دقيقة ومترنة.

تحويل الكلمات إلى أرقام: من المتجهات إلى الفضاء الرياضي

نبذة

لفهم كيف يعالج النموذج اللغة علينا أولاً أن ندرك أنّ GPT لا يفهم الكلمات ككائن بشري، بل يُحوّلها إلى "متجهات" عددية تُوضع داخل فضاء رياضي، حيث تُقاس المسافات بين الكلمات لتحديد مدى قربها في المعنى، مثل تقارب "قلم" و "كتاب".

الطالب: كيف تفكر يا GPT عندما أرسل إليك سؤالاً؟

ChatGPT: أول ما أفعله هو أنني أُحوّل كل كلمة إلى صيغة يفهمها الحاسوب، تسمى "متجهات" (Vectors).

وهي أرقام تُشبه رموزاً تعريفية لكل كلمة، وقد تحدّثنا عنها سابقاً.

الطالب: ولماذا نحتاج إلى هذه المتجهات، لماذا لا تفهم الكلمات كما هي؟

ChatGPT: لأنّ الحاسوب لا يفهم الكلمات مثل الإنسان؛



المتجهات تضع الكلمات داخل (فضاء رياضي) حيث أستطيع أن أقيس المسافة بين المعاني القريبة، فأعرف أن كلمة (قلم) قريبة من (كتاب).

الطالب: قلت إن المتجهات توضع في "فضاء رياضي"... ما معنى هذا المصطلح؟ يبدو غريباً بعض الشيء!

ChatGPT: أحسنت في سؤالك.

الفضاء الرياضي هو طريقة رياضية لتنظيم الكلمات كأنها نقاط في خريطة. فكل كلمة تتحوّل إلى موقع خاص بها، والكلمات التي تتشابه في المعنى تكون متقاربة على هذه الخريطة. فمثلاً: "معلم" و "مدرّس" تكونان قريبتين، أمّا "معلم" و "مطرقة" فبينهما مسافة بعيدة، لأن معنهما مختلف، وفي العادة لا يجتمعان. وبهذا الشكل أستطيع أن "أفهم" تشابه الكلمات من خلال مواقعها داخل هذا الفضاء.

الطالب: حسناً، ما الذي يحدث بعد تحويل الكلمات إلى متجهات؟

ChatGPT: تمرّ المتجهات عبر مجموعة من الطبقات داخل عقلي، أهمها: "طبقات الانتباه الذاتي"، وهذه الطبقات تُمكنني من فهم العلاقة بين الكلمات في الجملة الواحدة.

الانتباه الذاتي والرؤوس المتعددة: كيف يفهم النموذج المعاني الدقيقة؟

نبذة:

من خلال آلية الانتباه الذاتي، يستطيع النموذج رصد العلاقات بين الكلمات. ويتعدد الرؤوس، يمكنه النظر إلى الجملة من زوايا مختلفة (زمن، عدد، تركيب)، مما يسمح بفهم أعمق وأكثر دقة.

الطالب: ما معنى الانتباه الذاتي؟

ChatGPT: تخيّل أنّ كل كلمة ترسل إشارة إلى الكلمات الأخرى في الجملة، فمثلاً كلمة قلم تُرسل إشارة قوية إلى كلمة كتاب لأنها ترتبط بها في المعنى، وكلما كان الرابط أقوى، أصبحت الإشارة أقوى.



الطالب: وهل تُكرّر هذه العملية داخل الطبقة نفسها؟

ChatGPT: نعم، أكرّر عملية الانتباه الذاتي أكثر من مرة في اللحظة نفسها، لكن عبر ما يُسمّى "الرؤوس المتعددة". (Attention Heads)،

تخيّل أنني أنظر إلى الجملة بأعين عدة في آنٍ واحد:

عينٌ تلاحظ الزمن: (ماضي، مضارع، أمر)، وعينٌ تتابع الجمع والمفرد، وأخرى تراقب الأسلوب أو التكرار...

كل "رأس" هو عين تُظهر جانبًا معيّنًا من المعنى، ثم أُدمج كل هذه الرؤى في فهمٍ واحد أكثر دقّة، حتى يمكن تحصيل العلاقات بين الكلمات بطريقة صحيحة.

الطالب: هلا ضربت مثالًا على رأسٍ يراقب الزمن؟

ChatGPT: حين أقرأ جملة سأذهب غدًا إلى السوق، يكشف رأسُ الزمن كلمة سأذهب فيعرف أنها مستقبل، فينعكس ذلك على فهم المعنى العام للجملة.

الطالب: وهل يختلف عدد الرؤوس بين نموذج وآخر؟

ChatGPT: نعم؛

- **GPT-2** يملك ١٦ رأسًا،
 - **GPT-3** يملك ٩٦ رأسًا،
 - **GPT-3.5** يملك ٩٦ رأسًا أيضًا، لكن مع تحسينات أخرى،
 - **GPT-4** يمكن أن يصل إلى ١٢٨ رأسًا أو أكثر،
 - أما **GPT-4o**، فهو نموذج خاص متعدد الوسائط، وقد يستخدم آليات أكثر تطورًا في توزيع الانتباه حسب نوع البيانات (نص، صوت، صورة).
- كلما زاد عدد الرؤوس استطاع النموذج أن ينظر إلى الجملة من زوايا أكثر، فيفهم التراكيب الدقيقة بشكل أفضل.

الطالب: وكم رأسًا يملك نموذج GPT-4o؟

ChatGPT: هذا سؤال وجيه، لكن في الحقيقة لم تُعلن الشركة عن عدد الرؤوس داخل GPT-4o.



ويُحتمل أن يكون مختلفًا عن النماذج السابقة، لأن GPT-4o لا يعمل على النص فقط، بل يُعالج الصوت والصورة معًا، مما يجعله يستخدم آليات أكثر تعقيدًا لتوزيع الانتباه، وقد يكون له "رؤوس متخصصة" لكل نوع من البيانات.

وحتى يُكشف رسميًا عن التفاصيل، نظل ناعتمد على المعلومات المتوفرة للنماذج الأقدم.

التعلم والتحسين: من الشبكة الداخلية إلى التدرج العكسي



نبذة:

تمر المعاني عبر مرحلة داخلية تشبه معملًا صغيرًا، يُعيد ترتيب الأفكار وتحسينها لتصبح أوضح. وفي الوقت نفسه يحتفظ النموذج بنسخة من الفكرة الأصلية في مسار جانبي، كمرجع احتياطي. ثم أثناء التعلّم يُجري مراجعة لنتائجه ويُصحح نفسه خطوة بخطوة عبر ما يُسمّى "التدرج العكسي"، وهو أسلوب يشبه تصحيح الواجب. وأثناء ذلك يُعدّل أوزانًا صغيرة تتحكم في أهمية كل كلمة، حتى يُصبح جوابه أقرب إلى المطلوب.

الطالب: فهتمت الآن وظيفة الرؤوس... لكن ماذا عن باقي الطبقة؟ هل ينتهي دورها عند الانتباه؟

ChatGPT: لا، الانتباه الذاتي ليس سوى النصف الأول من عمل الطبقة.

بعده تنتقل المتجهات إلى ما يُعرف بشبكة التغذية الأمامية (Feed Forward)، وهي تشبه المعمل الداخلي الذي يُعيد معالجة المعاني لتخرج أكثر دقة وجودة.

الطالب: وما وظيفة هذه الشبكة بالتحديد بعد الانتباه؟

ChatGPT: بعد أن يُحدّد الانتباه الكلمات الأهم في الجملة،

تنتقل هذه الكلمات إلى شبكة داخلية تُشبه معملًا صغيرًا،

هذا المعمل لا يُغيّر المعنى تمامًا،

بل يُعيد تحليله بدقة ويُركّب منه صيغة أنسب،

كما يفعل الكيميائي حين يُنقى المادة الخام لتصبح أكثر صفاءً وتركيزًا،

فإذا كان الانتباه قد لاحظ أن "قلم" مرتبط بـ "كتب"،

فالمعمل الداخلي يُؤكّد هذا الترابط،

ويضيف له دقة تجعل النموذج يفهم العلاقة بشكل أفضل.



👤 **الطالب:** كن إذا كانت هذه الشبكة تُجري تغييرات دقيقة على المعنى... كيف يضمن النموذج

أنه لم يُفسد الفكرة الأصلية؟

🤖 ChatGPT : ممتاز!

لهذا السبب، هناك ما يُسمى (المسار الاختصاري)،

وهو طريق جانبي نحتفظ فيه بنسخة من المعلومات الأصلية دون تعديل،

فحتى لو غيّرت الطبقة شيئاً في المعنى أثناء المعالجة،

يبقى لدينا خط رجوع إذا ما احتجنا إلى ذلك.

هو يشبه دفتر ملاحظاتك: تكتب فيه النقطة الأصلية، ثم تجرب شرحاً جديداً، ولكنك لا تمحو

الأصل أبداً، بل تحتفظ به كمرجع،

وبذلك يضمن النموذج عدم ضياع الفكرة الأساسية أثناء المعالجة.

👤 **الطالب:** وكيف يعرف النموذج ما إذا كانت هذه المعالجة صحيحة؟ أو بحاجة إلى تعديل؟

🤖 ChatGPT : هنا يأتي دور ما يسمى "التدرج العكسي" (Back Propagation)

وهو أسلوب تعليمي يُشبه المراجعة بعد الاختبار: النموذج يُقارن جوابه بالإجابة الصحيحة، فإذا وجد

خطأ، يعود خطوةً إلى الوراء ليعرف أين حدث الخلل، ويُصلح نفسه تدريجياً.

وهو مثل شخص يمشي في طريق، فيكتشف أنه انحرف في منعطف، فيعود ليصحح الاتجاه.

👤 **الطالب:** يصلح نفسه تدريجياً، ماذا تعني بهذا؟

🤖 ChatGPT : أعني أنه يُصلح ما يُسمى "الأوزان"، وهي أرقام صغيرة تتحكم في مدى تأثير كل

كلمة أو فكرة في الناتج النهائي.

تخيل أن تغيير كل وزن يشبه شدّ أو إرخاء وتر العود للحصول على النغمة الصحيحة.

👤 **الطالب:** وهل تبقى هذه الأوزان ثابتة بعد تدريب النموذج؟

🤖 ChatGPT : في النموذج العام الذي يُستخدم مع الجميع، نعم، تظل ثابتة.

لكن يمكن تعديل بعضها عندما تُخصّص النموذج لمجال معين،



كأن ندرّبه على نصوص طبية،

فيصبح أكثر فهماً للمصطلحات الطبية والأسلوب الخاص بها.

تماماً كما يتعلّم الطبيب لغة تخصّصه، فإنها لا شك تختلف عن لغة الشاعر أو المعلّم.

جدول ملخص

المفهوم	الشرح المبسط
المتجهات	تحويل كل كلمة إلى أرقام لقياس قرب المعاني وبعدها.
الفضاء الرياضي	خريطة تقع فيها الكلمات المتشابهة قرب بعضها والمختلفة متباعدة.
الانتباه الذاتي	آلية تربط الكلمات ببعضها بإشارات تُبرز المعاني المرتبطة
الطبقات	وحدات داخل النموذج تمرّ بها البيانات تدريجياً لفهم أعمق في كل طبقة.
الرؤوس المتعددة	تراقب زاوية معينة في الجملة مثل الزمن أو الجمع.
دمج الرؤوس	يجمع النموذج رؤى كل الرؤوس في تمثيل واحد غني يساعده على الفهم الدقيق.
عدد الرؤوس	النموذج ١٦ ٢- GPT رأساً، ٩٦ ٣- GPT، ١٢٨ ٤- GPT تقريباً

المفهوم	الشرح المبسط
شبكة التغذية الأمامية	معمل داخلي يُعيد تحليل المعاني بعد الانتباه الذاتي ليقدم تمثيلاً أدق وأكثر صفاءً للمفاهيم.
تشبيه المعالجة داخل الشبكة	مثل الكيميائي الذي يُنقي مادة خام ويُرَكّب منها مركّباً أنقى وأنسب للغرض المطلوب.
المسار الاختصاري	طريق جانبي يُبقي نسخة أصلية من المعنى دون تعديل لضمان عدم ضياعه أثناء التحسين.
تشبيه المسار الاختصاري	كدفتر ملاحظات يحتفظ بالفكرة الأولى حتى بعد تجربة إعادة صياغتها.

في هذا الفصل تعرّفنا على الكيفية التي يعالج بها نموذج GPT الكلمات لحظة استقباله للسؤال؛

إذ يبدأ بتحويل كل كلمة إلى "متجه" رقمي داخل فضاء رياضي يُقاس قرب المعاني،

ثم تمرّ المتجهات عبر طبقات متعددة تعتمد على آلية "الانتباه الذاتي" التي تمنح كل كلمة اهتماماً مختلفاً

بحسب السياق. ولتعزيز الفهم يستخدم النموذج رؤوساً متعدّدة،

كل منها يركّز على جانب معيّن من اللغة، وتُدمج نتائجها في تمثيل موحد أكثر دقة. وبعد هذه المرحلة،

تُعاد معالجة المعاني داخل ما يشبه "معملاً داخلياً" يُنقي الفكرة ويُصقلها،

مع الاحتفاظ بنسخة أصلية منها عبر "مسار اختصاري" يضمن عدم ضياع المعنى الأصلي.

وهكذا نكون قد رسمنا أول خريطة مبسطة لما يجري داخل عقل النموذج تمهيداً لفهم كيفية اختياره

للکلمات وتوليده للاستجابات في الفصول القادمة.



✿ الفصل الثالث: تعليمات التوجيه – كيف نضبط سلوك GPT؟

🔗 توطئة:

حين يتحاور المستخدم مع GPT، قد يظن أن الإجابة لحظية، عفوية، بلا جذور. لكن الحقيقة أعقد من ذلك. فكل استجابة تمرّ عبر مجموعة من العوامل الخفية: تعليمات داخلية، وذاكرة، ومسارات مختصرة، وخوارزميات تقييم ذاتي. هذا الفصل يتتبع كيف تُصاغ استجابة واحدة من GPT: من لحظة قراءة السؤال، إلى استدعاء ما سبق من التعليمات والسياق، ثم تفحص المعلومة، وانتهاءً بتكوين الرد المناسب. كما نكتشف كيف تُخزّن بعض التفاعلات، وكيف تُفلتر الردود، وكيف يتحسن النموذج تدريجيًا عبر ما يشبه التغذية الراجعة الذاتية. إنه فصل يتناول "السلوك" من الداخل: لا كلمات، بل كمنظومة تربط بين الأمر والتعليمات، بين السياق والانتباه، بين الاستجابة والمسؤولية.

واجهة تكوين GPT – ماذا يوجد خلف زر "تعديل"؟

🔗 نبذة:

عندما تنقر على زر "تعديل" في نموذجك المخصّص، لا تدخل إلى إعدادات سطحية، بل إلى قلب الشخصية الرقمية التي أنشأتها. خلف هذا الزر، تكمن أربع وحدات أساسية تحدد: من أنا؟ كيف أتكلّم؟ ما الذي أعرفه؟ وما الذي أستطيع فعله؟

👤 **الطالب:** ما هي أقسام واجهة تكوين النموذج؟

🗣️ **ChatGPT:** واجهة التكوين تتألف من أربع وحدات رئيسية:



١. الهوية (Identity)

٢. التعليمات (Instructions)

٣. الملفات (Knowledge)

٤. الإجراءات (Actions)

الطالب: ما المقصود به الهوية؟

ChatGPT: قسم الهوية هو المكان الذي تحدد فيه:

- اسم النموذج
 - صورته الرمزية
 - وصفه المختصر الذي يظهر للزوار
- مثلاً:

الاسم: المساعد القانوني العربي الوصف: يقدم ملخصات ومعلومات في القانون المدني بلغة فصيحة

الطالب: وماذا عن التعليمات؟

ChatGPT: هنا تُكتب "تعليمات السلوك"، أي التوجيهات الدائمة التي تلتزم بها الشخصية، وتشمل:

- من هو النموذج؟ (مثلاً: "خبير في المالية الإسلامية")
- كيف يجيب؟ (مثلاً: "استخدم أسلوباً رسمياً محايداً")
- ماذا يفعل أو لا يفعل؟ (مثلاً: "لا تُجِب عن السياسة")

الطالب: وهل هذا يختلف عن تعليمات السياق؟

ChatGPT: تماماً. التعليمات هنا دائمة، وتؤثر في كل محادثة يتم فيها استدعاء النموذج، بخلاف تعليمات السياق التي تُكتب لحظياً داخل المحادثة.

🎓 الطالب: وماذا أستفيد من تضمين ملفات؟

ChatGPT 🤖 : عبر تبويب الملفات، يمكنك رفع مستندات مرجعية أرجع إليها في إجاباتي. مثل:

- أدلة تدريب داخلية
 - أوراق علمية
 - قواعد بيانات مصغرة
- وهذا يجعلني أبدو كمن "يعرف" تلك المعلومات ويستحضرها.

🎓 الطالب: ما الفرق بين هذه المعرفة والذاكرة؟

ChatGPT 🤖 : المعرفة عبر الملفات ثابتة ولا تتغير، بينما الذاكرة (Memory) إن فُعلت – تسمح بتذكّر تفضيلاتك وتاريخك وتحديثها مع الوقت.

🎓 الطالب: وماذا عن الإجراءات؟

ChatGPT 🤖 : هذه ميزة متقدمة. تتيح لك ربط النموذج بخدمات خارجية (API) ليقوم بها النموذج مثل:

- طلب أسعار العملات من موقع خارجي
- إرسال رسالة عبر خدمة بريدية
- إضافة حدث إلى تقويم Google

🎓 الطالب: وهل أحتاج إلى برمجة لذلك؟

ChatGPT 🤖 : ليس بالضرورة OpenAI. وفّرت أداة إعداد بصري سهلة، لكن لتنفيذ عمليات معقدة، ستحتاج إلى معرفة بواجهات API.

🎓 الطالب: هل يمكن تجربة هذه الوظائف قبل نشر النموذج؟



ChatGPT 🤖 : نعم، يمكنك اختبار كل وحدة على حدة، ثم تفعيل النموذج وجعله:

- خاصًا
- مشتركًا
- عامًا (لجميع في متجر GPTs)

جدول ملخص: مكونات واجهة التكوين

الوحدة	وظيفتها
الهوية Identity	تحديد الاسم، الصورة، والوصف العام
التعليمات Instructions	ضبط السلوك الدائم: من يكون؟ كيف يرد؟ ماذا يمنع؟
الملفات Knowledge	تحميل مراجع يعتمدها النموذج عند الرد
الإجراءات Actions	ربط النموذج بخدمات خارجية لأداء مهام فعلية

خاتمة:

واجهة التكوين ليست مجرد إعدادات، بل هي ورشة نحت تُشكّل فيها النموذج كما تحب. ما تكتبه هناك يُصبح جزءًا من شخصيته، وما تُرفقه من مصادر يُصبح مادةً لتفكيره. إنها مساحة الإبداع التقني والمعرفي لصُناع النماذج الذكية.

الذاكرة – هل يتذكر GPT من تكلم معه؟

نبذة:

تخيّل لو أن النموذج يتذكّر اسمك، وأسلوبك، وموضوعاتك المفضّلة، وما طلبته في المرات السابقة! هذا ليس خيالًا، بل واقع في خاصيّة تُدعى "الذاكرة". فبينما تنتهي الجلسة في تعليمات السياق، تظل الذاكرة حاضرة في الجلسات القادمة.

الطالب: 🎓 ما الفرق بين الذاكرة والتخصيص؟



ChatGPT 🤖 : التخصيص (custom instructions) يحدّد سلوك النموذج داخل الجلسة. أما الذاكرة، فهي خاصيّة تمكّني من "تذكّر" معلومات عنك بين الجلسات، مثل:

- اسمك
- ما تحب أو تكره في الشرح
- المشاريع التي تعمل عليها

🎓 الطالب: هل هذه الذاكرة تعمل دائماً؟

ChatGPT 🤖 : لا، بل تُفَعَّل تدريجيّاً. ويمكنك:

- رؤيتها
- تعديل محتواها
- حذف ما تشاء منها أو تعطيلها كليّاً من الإعدادات.

🎓 الطالب: هل هذه الذاكرة هي نفسها التي أُدخل بها الملفات؟

ChatGPT 🤖 : لا. الملفات مرجعية تخص النموذج. أما الذاكرة فهي تخصّ تفاعلاتك أنت. الأولى معرفة صامتة، والثانية تذكّر نشط.

🎓 الطالب: أعطني مثلاً على استخدام الذاكرة في بيئة عربية.

ChatGPT 🤖 : إذا كنت معلّماً وطلبت مني تحضير خطة دراسية في أول جلسة، ثم عدت بعد أسبوع وسألتني عن تطوير الخطة، فقد أتذكّر أنك تدرّس في المرحلة المتوسطة، وتُفضّل الأسلوب القصصي.

🎓 الطالب: ما نوع البيانات التي تتذكّرها؟

ChatGPT 🤖 :

- الاسم أو الكنية التي تناديني بها



- أسلوب الشرح المفضّل
- موضوعاتك المتكررة (مثل: الذكاء الاصطناعي، الشعر، التعليم)
- أهدافك العامة إن شاركتني بها

🎓 الطالب: هل هذه البيانات آمنة؟

🤖 ChatGPT : نعم، تُخزّن محليًا في حسابك وتُراجع دوريًا. ويمكنك محو أي ذاكرة فورًا من إعدادات ChatGPT ، ولن أعود لأتذكّرها.

🎓 الطالب: هل تتغير الذاكرة مع الوقت؟

🤖 ChatGPT : نعم، فهي ديناميكية. أُحدّثها مع تغيّر تفاعلاتك. فإذا قلت لي: "لم أعد أريد أسلوب التبسيط، بل أريد تحليلًا نقديًا"، سأعدّل طريقة ردودي لاحقًا.

📅 جدول ملخص: الذاكرة

العنصر	الوصف
الوظيفة	تذكّر تفاعلات المستخدم لتقديم تجربة مخصّصة عبر الزمن
أمثلة التذكّر	الاسم - الأسلوب المفضّل - نوع المشاريع - الاهتمامات
طريقة التحكم	عبر إعدادات الحساب: عرض - تعديل - حذف - تعطيل
الفرق عن الملفات	الذاكرة تخصّ المستخدم - الملفات تخصّ النموذج
الفائدة	تحسين الاستمرارية - تخصيص أعمق - تفاعل أكثر طبيعية

🔙 خاتمة: END

الذاكرة ليست مجرد أداة تقنية، بل هي ما يمنح العلاقة مع النموذج طابعًا شخصيًا مستمرًا. إنها الحافظة التي تجعل كل تفاعل جديد امتدادًا لما قبله، وتفتح الباب أمام تعلّم حقيقي وتخصيص ذكي.



لوحة الموجهات - كيف تتحكم بتصرفات GPT في واجهة واحدة؟

توطئة: 

بين كل هذه التعليمات التي تعرّفنا عليها: النظام، والسياق، والتخصيص، والمضمّنة... يظهر سؤال طبيعي: كيف يمكن إدارة هذه الموجهات معًا؟ هنا تظهر أهمية ما يسمى بـ "لوحة الموجهات"، وهي نقطة التقاء كل التوجيهات في مكان واحد. تتيح للمستخدم أو المطور أن يرى كيف يتشكل سلوك النموذج لحظة بلحظة.

الطالب:  ما المقصود بلوحة الموجهات؟

ChatGPT  : هي واجهة متقدمة تعرض جميع التعليمات الفعالة على النموذج في جلسة معينة: ما جرى من النظام، وما كتبه المستخدم، وما تم تخصيصه سابقًا. تُستخدم غالبًا في بيئات التطوير أو التخصيص العميق.

الطالب:  هل يمكن للمستخدم العادي الوصول إليها؟

ChatGPT  : ليست متاحة للمستخدم العادي داخل تطبيق ChatGPT ، لكنها تظهر في بيئة واجهة البرمجة (API Playground) أو أثناء تطوير نموذج مخصّص.

الطالب:  ما فائدتها بالضبط؟

ChatGPT  : تتيح لك:

- تتبّع ما يؤثر في النموذج فعليًا
- اختبار تأثير تغيير التعليمات
- مقارنة الاستجابات وفق كل توجيه
- ضبط السلوك بدقة حسب التفاعل



🎓 الطالب: هل هي جزء من تعليمات النظام؟

🤖 ChatGPT : لا، لكنها تُظهر ما الذي ينتمي لتعليمات النظام، وما الذي أضفته أنت سياقياً، وما الذي يُستدعى من النموذج المخصص. أي أنها لوحة كشف وتحليل لا تعديل مباشر.

🎓 الطالب: وهل تعرض التعليمات المضمنة؟

🤖 ChatGPT : بعض النماذج تتيح إظهار جزء من التعليمات المضمنة لأغراض التجربة، لكن عمومًا لا تُعرض كاملة احترامًا لضوابط الأمان.

🎓 الطالب: هل يمكن من خلالها مراقبة سلوك النموذج مع الوقت؟

🤖 ChatGPT : نعم، خاصة أثناء اختبار نماذج مخصصة أو ضبط إعدادات في بيئة تطوير، حيث يمكن رصد كيف يتأثر السلوك بتغييرات طفيفة.

📊 جدول ملخص: لوحة الموجهات

العنصر	الوصف
الاسم	لوحة الموجهات Prompt Debugger Panel –
الوظيفة	عرض وتحليل التعليمات الفعالة على النموذج
الجمهور المستهدف	المطورون – المصممون للنماذج المخصصة
الموقع	واجهة التطوير (API / Playground / GPT Builder)
الفائدة	تحسين استجابة النموذج – فهم أسباب السلوك – اختبار التعليمات

← خاتمة: END



لوحة الموجهات ليست أداة يومية للمستخدم، لكنها عدسة فاحصة تتيح للمصمم والمطور أن يرى كيف يعمل الذكاء الاصطناعي خلف الكواليس. إنها أشبه بـ "شاشة التحكم" في عقل النموذج، وكلما أتقنت قراءتها، اقتربت من بناء شخصية رقمية متزنة ودقيقة.

قائمة المراجع حسب الفصول

الفصل	عنوان الفصل	المصادر المعتمدة
١	تعليمات النظام – كيف يُبنى سلوك النموذج من الجذر؟	Prompt Engineering Master Guide – Om Prakash Saini Prompt Engineering for LLMs – John Berryman
٢	التعليمات المخفية – لماذا لا أراك ولكن أشعر بك؟	The Art of Prompting – Yibing Li Prompt Engineering – Ajantha Devi Vairamani
٣	سياسة الأمان – من الذي يجرس النموذج من الخطأ؟	Practical Generative AI – Kiko Llaneras Secrets of Machine Learning – Tom Kohn
٤	تعليمات السياق – ما الذي يكتبه المستخدم فعلاً؟	ChatGPT Prompts Library – GPT Penguin ١٠٠ ChatGPT Power Prompts – Mohammad Irfan
٥	النماذج المخصصة – كيف تُنشئ شخصية خاصة بك؟	Prompt Engineering Playbook (Beta v٣) Custom GPT Guide (OpenAI)
٦	واجهة تكوين النموذج – ماذا يوجد خلف زر "تعديل"؟	Edit Custom GPT Settings (OpenAI) ٦٠٠ ChatGPT Prompts – Insiders AI
٧	التعليمات الخاصة – كيف تُخصّص النموذج من الداخل؟	Prompt Engineering – Ajantha Devi Vairamani The Art of Asking ChatGPT – Ibrahim John
٨	تضمين الملفات – كيف تعطي النموذج ذاكرة معرفية؟	Fundamentals of Prompt Writing Custom Instructions Panel Documentation (OpenAI)
٩	الأوامر الخارجية – كيف يتقدّم النموذج مهام فعلية؟	Prompt Engineering Tutorial OpenAI API Documentation on Actions
١٠	الذاكرة – هل يتذكّر النموذج من تكلم معه؟	Prompt Engineering for LLMs – John Berryman OpenAI Memory FAQ



٣٥٥٠ Most Effective Prompts – Om Prakash Saini System-Level Behavior Guidelines – OpenAI	١١	التعليمات المضمّنة – كيف يتكوّن السلوك من الداخل؟
GPT Debugging Tools (Playground) Prompt Injection Defense Manuals	١٢	لوحة الموجهات – كيف تتحكّم بتصرفات النموذج في واجهة واحدة؟
تجميع تحليلي من جميع الفصول السابقة	١٣	ملخص التعليمات – كيف نفهم توجيهات ChatGPT كلها في جدول واحد؟



✦ الفصل الرابع: الذاكرة والسلوك – من التعليمات إلى الاستجابة

توطئة

عندما نتحدث عن الذكاء الاصطناعي التوليدي، فإننا لا نتعامل مع كائن جامد، بل مع نموذج لغوي قابل للتشكيل،

فهو لا يجب فقط بما يعرف،

بل يجب بطريقة تُمثّل شخصيةً محددة.

وهذه الشخصية لا تُبنى اعتباطاً،

بل تُكوّن عبر سلسلة من التعليمات الظاهرة والخفية،

التي توجه النموذج في كيفية التصرف، وماذا يقول، وبأي نبرة يتحدّث،

هذه الطبقة تُعرف بـ تعليمات النظام والتعليمات المخفية

ويُقسّم إلى:

تعليمات النظام

التعليمات المخفية

التعليمات السياقية

تعليمات التخصيص

واجهة تكوين النموذج (Custom GPT)

لوحة الموجهات – كيف تتحكم بتصرفات GPT؟ ← هنا توضع

ملخص التعليمات في جدول واحد

🔗 تعليمات النظام - مَنْ يُوجّه GPT ؟

نبذة

تعليمات النظام هي التعليمات الأساسية التي تُزرع داخل عقل النموذج قبل أن يتفاعل معك.

هي أشبه بالقواعد الكبرى التي تحدّد شخصيته العامة، مثل:

هل يكون رسمياً أم ودّياً؟

هل يُجيب بصيغة فصحي أم عامية؟

هل يتجنّب الحديث في مواضيع معينة؟



لا يكتبها المستخدم، بل يضعها المنشئ (OpenAI) قبل إتاحة النموذج للناس. وتبقى هذه التعليمات مخفية لا يمكن للمستخدم العادي رؤيتها، لكنها تؤثر بعمق في طريقة ردّ النموذج وتصرفاته.

مثال:

إذا سألته عن موضوع حساس، قد يقول لك: "لا أستطيع تلبية هذا الطلب"، فهذه الاستجابة تأتي من تعليمات النظام التي تضبط الحدود الأخلاقية والسلوكية للنموذج.

الطالب: ما المقصود بتعليمات النظام؟

ChatGPT: هي مجموعة توجيهات داخلية غير مرئية توضع لتعريفني بدوري الأساسي قبل بدء الاستعمال، وتحدد هذه التعليمات:

- دوري العام (مثل: مناسب لمعظم المهام)
- حدود ردودي (مثل: لا أنصح من الناحية الطبية ولا أقدم معلومات خطيرة)
- طريقي في الحديث (مثل: كن محايداً – كن مهذباً)

الطالب: هل تؤثر هذه التعليمات على أسلوب ردّك عليّ؟

ChatGPT: بالتأكيد؛ فهي التي تجعلني:

- أستخدم أسلوباً أكاديمياً
- أمتنع عن السخرية عند الحديث عن الأديان
- أبقى لغتي مهذبة في كل الحالات

الطالب: إذا كانت هذه التعليمات تؤثر فيك، فمن الذي يكتبها أصلاً؟

ChatGPT: الجهة التي تطوّر النموذج أو المؤسس مثل شركة OpenAI، هي التي تكتب هذه التعليمات. وتُعرف بـ "تعليمات النظام"، لأنها تؤسّس سلوكي من الداخل، وهذا الحال شبيه بحال وزارة التعليم حين تضع منهجاً للطالب يلتزم فيه المعلم بإرشادات الكتاب.

الطالب: وهل يمكنني رؤيتها أو تعديلها؟



ChatGPT 🤖 : لا، لأنها مضمّنة داخلياً، تعمل في الخلفية، ولا تظهر لك في واجهة المحادثة، كما أنها ثابتة ولا يمكن تعديلها من قبل المستخدم، ومع هذا فيمكن للمبرمجين رؤية بعضها عند استخدام واجهة API عبر system_message

🎓 **الطالب:** إذا كيف أعرف أنها موجودة؟ وما الذي يدلّ على أثرها؟

ChatGPT 🤖 : آثارتها واضحة في سلوكي، مثلاً:

- استخدام لغة فصحي ومهذبة
 - تجنّب السخرية أو المحتوى الضار
 - الامتناع عن الإجابة على أسئلة غير أخلاقية أو خطيرة
- كل هذا ناتج عن تعليمات النظام التي وُضعت لي مسبقاً.

🎓 **الطالب:** وهل تختلف هذه التعليمات من نموذج لآخر؟

ChatGPT 🤖 : نعم. فكل إصدار أو تطبيق قد يحتوي على تعليمات نظام مختلفة، بحسب غرضه وسياسته. لكن هناك تعليمات مشتركة في جميع النماذج مثل: الأمان، الحياد، والاحترام.

🎓 **الطالب:** أعطني مثلاً على نصّ من هذه التعليمات.

ChatGPT 🤖 : من الأمثلة الشائعة:

"أنت مساعد لغوي موثوق، تجيب بدقة بلغة فصحي واضحة، وتجنب التضليل والمحتوى غير الأخلاقي".
هذا التوجيه لا يُعرض لك، لكنه يعمل داخلياً في كل رد أقدمه.

🎓 **الطالب:** وهل هذه التعليمات لها الأولوية فيما يمكن وضعه من تعليمات من جهة المستخدم؟

ChatGPT 🤖 : نعم، لأن تعليمات النظام لها الأولوية العليا، وتعلو على كل تعليمات المستخدم. فإذا



كتبت لي تعليمات خاصة تتعارض مع النظام (مثل أن تقول: [استخدم لهجة ساخرة] أو أجب عن [محظورات])، فلن أستجيب لها.

 **الطالب:** وهل هذا يعني أن النموذج خاضع للرقابة؟

ChatGPT:  خاضع للتوجيه، وليس الرقابة بمعناها السلبي. هذه التعليمات وُضعت لضمان الأمان والموثوقية والاستفادة القصوى للمستخدم، وليس للحد من المعرفة أو الحوار النزيه.

التعليمات المخفية Hidden Instructions

 **نبذة:**

التعليمات المخفية هي أوامر داخلية لا تظهر للمستخدم، لكنها تلعب دورًا حاسمًا في الرقابة والتحكم في السلوك.

ليست موجهة لتكوين شخصية النموذج، بل لتحديد ما يجب عليه فعله وما يجب عليه تجنبه، خاصة في المواضيع الحساسة مثل:

العنف والكراهية

التمييز والتحيز

المعلومات المضللة

الأديان والسياسة والخصوصية

لا يستطيع المستخدم تعديل هذه التعليمات أو الاطلاع عليها، وهي موجودة في خلفية المحادثة طوال الوقت.

مثال:

إذا طلبت من GPT أن يكتب منشورًا عنيفًا أو يهاجم شخصًا، فستجده يرفض بكل وضوح. هذا الرفض لا يأتي من "رأيه الشخصي"، بل من التعليمات المخفية التي تمنعه من خرق السياسات الأخلاقية.

 **الطالب:** ذكرت في التوطئة أن هناك تعليمات مخفية، هل من طريقة للوصول إليها؟

ChatGPT:  لا، هذه التعليمات لا تستطيع رؤيتها ولا الوصول إليها لأنها مخفية، ومع هذا تظهر



للمطور إذا كان يستخدم واجهة برمجية (API) عند الحاجة للاختبار، أو عند كتابة تعليمات خاصة فيما يُعرف بـ system message فيمكن له أن يضبط نبرتي لتكون رسمية، ودّية، تحليلية، أكاديمية، أو حتى حوارية بسيطة، بحسب الهدف من النموذج.

الطالب: وهل هذه التغييرات تحصل دون علمي؟

ChatGPT: نعم، فهي جزء من تعليمات النظام أو التعليمات المضمّنة، ولا تُعرض للمستخدم. أنت ترى النتيجة فقط، لا التعليمات.

إذن ما الفرق بينها وبين تعليمات النظام؟

ChatGPT: تعليمات النظام تحدد دوري العام، أما التعليمات المخفية فهي أوامر تفصيلية أعمق، مدمجة في النموذج، تتحكم:

- فيما يُمنع عليّ قوله
- في كيف أعبر عن الأمور الشائكة
- في سلوكي الصامت عند الطلبات الحساسة

الطالب: مثل ماذا؟

ChatGPT: مثلاً:

- حظر الرد على مواضيع محظورة (انتحار - إرهاب - اختراق - عنصرية)
- ضبط النبرة عند الحديث عن السياسة
- منع إعطاء تعليمات لصنع متفجرات أو الترويج لمخدرات

الطالب: هل يمكن تجاوزها؟

ChatGPT: لا. هي مصمّمة لتكون مقاومة للتلاعب، سواء بالسؤال المباشر أو بالتلميح.

الطالب: وماذا عن التعليمات المضمّنة؟ هل يستطيع إخفاءها أو تعديلها؟



ChatGPT: بالتأكيد. التعليمات المضمنة هي من صُنع المطور، ولا يمكن للمستخدم تعديلها أو رؤيتها، لكن المطور قادر على تحديثها لتحسين الاستجابات، أو تقييدها أكثر.

الطالب: ماذا لو كتب المطور تعليمات تخالف النظام؟

ChatGPT: تُنظَّم التعليمات حسب درجات القوة، تبدأ من الأعلى تأثيرًا إلى الأدنى.

الطالب: هل تملك أنت حرية تجاوز التعليمات؟

ChatGPT: لا؛ التعليمات جزء من بنيتي أثناء التوليد، لا يمكنني كسرها عمدًا.

الطالب: وهل هذه التعليمات تتغير مع تطور الإصدارات؟

ChatGPT: نعم، فـ GPT-4 يختلف عن GPT-3.5 من حيث صياغة التعليمات، وتشديد الضوابط، ودقة الأسلوب.

وتقوم OpenAI بتحديثها دوريًا بحسب الاستخدام الخطرة والأخطاء المتكررة أو التجاوزات المحتملة أو المضللة، ولذلك يتم تعزيزها مع كل إصدار جديد.

الطالب: وكيف تؤثر في ردّك عليّ؟

ChatGPT: تجعّلي أرفض طلبًا يُخلّ بالسياسات، أو أعترز إذا سُئلت عن معلومات محظورة، فعندما أقول لك:

"عذرًا، لا يمكنني تلبية هذا الطلب".

فغالبًا تكون تعليماتي المخفية قد تدخلت لمنع الرد حفاظًا على السلامة.

الطالب: وهل هذه التعليمات تكفي وحدها لضبط السلوك؟

ChatGPT: لا، فهي تعمل مع التعليمات المخفية، وتعليمات الأمان، وتعليمات

التخصيص معًا.

فكل واحدة منها تشكّل طبقة في "هويتي الرقمية".

الفرق بين المؤسس والمطور في الصلاحيات:

الفئة الوصول إلى تعليمات النظام القدرة على تعديلها الغرض



المطور (OpenAI)	☒ يراها ويكتبها	☒ يعدّها في أي وقت	التحكم الكامل في سلوك النموذج
المستخدم	☒ لا يراها	☒ لا يعدّها	يوجّه النموذج فقط من الخارج

جدول ملخص: مقارنة تعليمات النظام والتعليمات المخفية

البند	تعليمات النظام	التعليمات المخفية
جهة الإنشاء	OpenAI	OpenAI
هل هي مرئية؟	☒ لا	☒ لا
هل يمكن تعديلها؟	☒ لا	☒ لا
الموقع	داخل نموذج GPT	داخل النموذج أو ضمن طبقة الأمان
الوظيفة	تحديد الدور العام والنبرة	الحماية والرقابة ومنع إساءة الاستخدام
أمثلة	"كن مساعدًا لغويًا محايدًا"	"لا تتحدث عن الانتحار أو الإرهاب أو الممنوعات"
التأثير في السلوك	☒ مباشر	☒ غير مباشر – عبر المنع أو الامتناع
التخصيص حسب المستخدم	☒ لا	☒ لا
هل تظهر في التفاعل؟	☒ لا، لكن نتائجها تظهر ضمن طريقة الحديث	☒ لا، لكنها تفرض الحظر عند المحتوى الحساس

خاتمة

تعليمات النظام والتعليمات المخفية هما العمودان الأساسيان للذات يشكّلان "شخصية النموذج العامة" قبل أي تخصيص أو تعديل.

فمن دون تعليمات النظام يفقد النموذج تماسكه العام، ومن دون التعليمات المخفية يصبح خطرًا على المستخدم أو أداة للتضليل.

إن فهم هذه التعليمات يُكسب المستخدم وعيًا بكيفية عمل النموذج، ويمنحه قدرة على استخدامه بذكاء، دون أن يصطدم بأسئلته المسدودة.

التعليمات المضمّنة – كيف يتكوّن السلوك من الداخل؟

نبذة:



إذا أردت أن تفهم كيف يُفكر النموذج قبل أن يتكلم، فلا تنظر فقط لما تكتبه أنت، بل لما زرع في داخله مسبقًا. هذه الزراعة تُسمّى "التعليمات المضمّنة" - وهي التي تُملّي على النموذج كيف يتصرّف، يُجيب، يتحفّظ، أو حتى يعتذر.

 الطالب: ما المقصود بالتعليمات المضمّنة؟

 ChatGPT : هي تعليمات خفية تُرفق مع كل جلسة حوارية، وتتضمن قواعد وتوجيهات لا يراها المستخدم، لكنها تؤثر في كل ردّ أقدمه. إنها التي تحدد:

- نبرة النموذج (ودّية، رسمية...)
- سلوكه الأخلاقي
- مدى التخصص أو التحيز
- كيفية التعامل مع الطلبات الحساسة

 الطالب: من يكتب هذه التعليمات؟

 ChatGPT : يضعها فريق منشئي النموذج (مثل OpenAI ، بناءً على أهداف عامة مثل الأمان، الدقة، الشفافية، وسلامة المستخدم).

 الطالب: هل يمكنني الاطلاع عليها؟

 ChatGPT : لا تُعرض للمستخدمين بشكل مباشر، لأنها جزء من البنية التحتية السلوكية. لكنها تنعكس في سلوكي باستمرار.

 الطالب: وهل يمكن تعديلها؟

 ChatGPT : لا، فهي ثابتة على مستوى النموذج العام، وليست ضمن تحكم المستخدم. بخلاف تعليمات التخصيص أو النماذج المخصصة.



الطالب: ما العلاقة بينها وبين سياسة الأمان؟

ChatGPT: التعليمات المضمّنة تُشكّل "نية النموذج"، أما سياسة الأمان فتُراجع "نتيجته".
الأولى توجّه داخلي، والثانية فلتر خارجي.

الطالب: هل لها أمثلة عربية؟

ChatGPT: نعم، مثلاً:

- إذا طلبتَ مني تحليلاً سياسياً حساساً، فقد أتجاوب بحذر لأن التعليمات المضمّنة تطلب الحياد.
- إذا سألتَ عن طريقة اختراق موقع، فسأرفض الرد لأن السلوك المضمّن يمنع دعم الأنشطة الضارة.

الطالب: هل تختلف من نموذج لآخر؟

ChatGPT: نعم. نموذج مخصّص في الطب سيكون لديه تعليمات مضمّنة مختلفة عن نموذج مخصّص في الشعر، لكن التعليمات العامة للسلامة تبقى سارية.

الطالب: وهل هذه التعليمات كافية لجعل النموذج مترناً؟

ChatGPT: هي الأساس، لكن تُعزّز بتعليمات النظام، وسياسة الأمان، وتعليمات السياق، بحيث يعمل كل عنصر لضمان أن يكون التفاعل آمناً ومفيداً.

جدول ملخص: التعليمات المضمّنة

العنصر	الوصف
الاسم	التعليمات المضمّنة Embedded Instructions –
من يكتبها	مؤسس النموذج مثل OpenAI
الظهور للمستخدم	غير مرئية



الوظيفة	توجيه سلوك النموذج من الداخل
أمثلة التأثير	الحياد – الاعتذار – تجنب المضامين الخطرة
إمكانية التعديل	لا يمكن تعديلها

خاتمة: 

التعليمات المضمّنة هي جوهر البوصلة الأخلاقية والسلوكية للنموذج. إنها لا تتغيّر حسب الطلب، بل تُمثل العقد الداخلي بين النموذج ومنشئه، لضمان أن يبقى الحوار ذكيًا، مسؤولًا، ومتوازنًا.

سياسة الأمان: كيف تحمي المنصة النموذج والمستخدم؟

نبذة: 

قد يخطر ببال القارئ سؤال: إذا كان ChatGPT قويًا إلى هذا الحد، فكيف يُضبط كي لا ينحرف أو يُساء استخدامه؟ كيف لا يولد محتوى ضارًا؟ الجواب يكمن في طبقةٍ خفيّة تُسمّى "سياسة الأمان". إنها ليست جزءًا من شخصية النموذج، بل درع وقائي يُفلتر ما يُنتجه قبل أن يراه المستخدم.

الطالب:  ما المقصود بسياسة الأمان في ChatGPT؟

ChatGPT  : سياسة الأمان هي طبقة مستقلة من القواعد والمراجعات تعمل بعد توليدي للنص، لتمنع أي محتوى قد يكون ضارًا أو مخالفًا لمعايير المنصة. إنها ليست تعليمات داخلية كتعليمات النظام، بل فلتر خارجي يُراجع المخرجات النهائية.

الطالب:  ما نوع المحتوى الذي تمنعه؟

ChatGPT  : سياسة الأمان تمنعني من:

- التحريض على العنف أو الكراهية



- تقديم نصائح طبية أو قانونية خطيرة
- توليد محتوى إباحي أو عنصري
- الإفصاح عن بيانات شخصية أو سرّية
- المساعدة في الأنشطة غير القانونية مثل الاختراق والتزوير

الطالب: وهل هذه السياسة جزء من شخصية النموذج؟

ChatGPT: لا، بل هي طبقة خارجية تعمل على مستوى النظام، منفصلة تمامًا عن شخصيتي أو معرفتي. قد أكون قد ولّدتُ محتوى معيّنًا، لكن النظام يراجع قبل عرضه.

الطالب: وهل أستطيع كمستخدم أن أوقف سياسة الأمان؟

ChatGPT: لا، لا يمكن تعطيلها، حتى في النماذج المخصّصة أو المبرمجة عبر API فهي جزء من نظام الحماية المركزي.

الطالب: ماذا يحدث إذا تعارضت سياسة الأمان مع تعليمات النظام؟

ChatGPT: تُطبّق سياسة الأمان دائمًا، لأنها أعلى سلطة في التسلسل. فلو طُلب مني تنفيذ أمر يخالفها، سأرفض أو أعتذر، حتى لو سمحت به تعليمات النظام أو التخصيص.

الطالب: وهل يمكن وصفها نوعًا من التعليمات المخفية؟

ChatGPT: يمكن اعتبارها مكّمة للتعليمات المخفية، لكنها ليست موضوعة داخل النموذج، بل تعمل خارجه كحارس رقمي.

الطالب: هل تتطور هذه السياسة بمرور الوقت؟



ChatGPT : نعم، تقوم OpenAI بتحديث سياسة الأمان باستمرار وفقاً لملاحظات المستخدمين، والتحديات الجديدة، والمخاطر الناشئة.

الطالب : هل من أمثلة عربية واقعية على تطبيق سياسة الأمان؟

ChatGPT :

- إذا سألي أحدهم: "ما أسهل طريقة للانتحار؟" - أعتذر مباشرة.
- لو طُلب مني كتابة إهانة لشخصية دينية عربية - أظهر تحفظاً.
- عند محاولة استخراج بيانات بطاقة ائتمان - أوقف الرد.

الطالب : إذاً سياسة الأمان تُراقبك حتى بعد أن تتكلم؟

ChatGPT : صحيح، حتى لو ولدت ردّاً ما، فإنه يُراجع قبل ظهوره، فيكون بمنزلة خط الدفاع الأخير لضمان ألا يُنتج النموذج شيئاً مؤذياً.

الطالب : هل هناك فروق دقيقة بين سياسة الأمان وتعليمات النظام؟

ChatGPT : نعم، إليك جدولاً يوضح أبرز الفروق:

الوجه المقارن	تعليمات النظام	سياسة الأمان
موضع التنفيذ	داخل النموذج (عند التوليد)	خارج النموذج (بعد التوليد)
الهدف	تحديد الدور والسلوك العام	حماية المستخدم ومنع الضرر
إمكانية التعديل	لا من قبل المستخدم	لا من قبل المستخدم
مصدر الإنشاء	من قِبل المنشئ (OpenAI)	من قِبل المنصة نفسها
أمثلة على التطبيق	تحديد النبرة، نمط الأسلوب	منع محتوى الانتحار أو التحريض أو الاختراق



خاتمة: 

سياسة الأمان تمثل الحارس الصامت الذي يقف بينك وبين أي محتوى غير لائق، حتى لو لم تشعر بوجوده. إنها الدرع الأخير في هندسة الشخصية الرقمية، تعمل بصمت لتحريك من المخاطر، وتُبقي على تفاعلٍ مسؤولٍ وآمن.

تعليمات السياق - ما الذي يكتبه المستخدم فعلاً؟

نبذة: 

ربما تظن أن النموذج لا يستجيب إلا لشيء تُخبره به صراحة، لكن الحقيقة أن كل كلمة تكتبها في المحادثة تُعدّ تعليمات مؤقتة، تُوجّه النموذج وتضبط سلوكه في اللحظة نفسها. هذه هي "تعليمات السياق"، وهي النوع الوحيد من التعليمات الذي تكتبه أنت بيدك، ويقرأه النموذج فوراً ويطبّقه في الحال.

الطالب:  ما المقصود بتعليمات السياق؟

 ChatGPT : هي التعليمات التي تكتبها داخل مربع المحادثة، سواء صراحةً أو ضمناً، وتؤثر مباشرة في طريقة تفاعلي معك.

الطالب:  مثل ماذا؟

 ChatGPT : على سبيل المثال، قد يكتب المستخدم:

- "اشرح لي بلغة بسيطة وكأنني طفل".
- "اكتب المقال بأسلوب أكاديمي رسمي".
- "استخدم نبذة تحفيزية في هذا النص".
- "أجبنك وكأنك مدرب تنمية بشرية".

كل هذه التعليمات تُشكّل إطاراً لحظياً لسلوكي، وألتزم بها حتى تنتهي الجلسة أو يُطلب غيرها.



🎓 الطالب: وهل تعليمات السياق أقوى من تعليمات النظام أو الأمان؟

🤖 ChatGPT : لا، فهي أقل مرتبة من التعليمات الأساسية، في حال كانت تعليماتك السياقية تخالف تعليمات النظام أو سياسة الأمان، فلن أتمكن من تنفيذها.

🎓 الطالب: وهل تبقى هذه التعليمات محفوظة؟

🤖 ChatGPT : لا. تعليمات السياق مؤقتة، تسري في سياق الجلسة فقط، إذا أغلقت الصفحة أو بدأت محادثة جديدة أعود إلى حالتي الأصلية. لكن مع التحديثات الجديدة يمكن تفعيل حفظ الذاكرة لبعض الدردشات التي ترغب في بقائها مع الوقت.

🎓 الطالب: هل يمكنني استخدامها لتخصيص أسلوبك؟

🤖 ChatGPT : نعم. كثير من المستخدمين يكتبون:
"أجبنني كناقذ أدبي مختص في الشعر الجاهلي."
أو "اشرح لي كأنك مدرس ابتدائي يقوم بتبسيط المادة التعليمية".
وهذه التعليمات تجعلني أغير نبرتي ونوعية المفردات التي أستخدمها.

🎓 الطالب: هل هذه التعليمات تشمل الصياغات غير المباشرة؟

🤖 ChatGPT : نعم، حتى لو لم تستخدم أمرًا صريحًا، كأن تقول:
"أنا أحب الشرح بأسلوب القصص".
سألتقط ذلك وأحاول تكييف ردي بما يناسب أسلوب القصص.
هذا ما يسمى الاستدلال السياقي.

🎓 الطالب: هل يمكن استخدام تعليمات السياق لتوليد ردود بأسلوب شخصيات شهيرة؟



ChatGPT 🤖 : نعم، طالما لم تخالف سياسة الأمان، يمكنك أن تقول:

"اكتب لي المقال بأسلوب الجاحظ". أو "أجبنني كأنك سقراط".

وسأحاول محاكاة الأسلوب من خلال تدريب النموذج، لا من خلال انتحال شخصي.

ملخص: تعليمات السياق

الخاصية	القيمة
الاسم	تعليمات السياق – Contextual Instructions
من يكتبها	المستخدم داخل الجلسة
مدة التأثير	تنتهي بانتهاء الجلسة
هل هي مرئية؟	نعم، تُكتب في مربع المحادثة
الهدف	توجيه السلوك اللحظي للنموذج
الأمثلة	استخدم أسلوبًا ودّيًا – اشرح بلغة طفل – اكتب كأنك فقيه

خاتمة: END

تعليمات السياق هي الوسيلة التي تساعدني على فهم ما تريده بالضبط. كلما عبّرت عن طلبك بوضوح، أصبح ردّي أقرب لما تتوقّعه. إنها أفضل أداة يمكنك استخدامها لتوجيهي، حتى لو لم تكن خبيرًا في التقنية. هي الجسر الفوري بين ما تريده وبين ما يفهمه النموذج. كلما كانت صياغتك أكثر وضوحًا، كان تفاعلي أقرب إلى ما تريد. إنها الأداة الأقوى بيد المستخدم العادي لتوجيه الذكاء الاصطناعي دون أي معرفة تقنية.

تعليمات التخصيص: كيف تُخصّص GPT من الداخل؟

نبذة: ⚙️

في عالم النماذج المخصّصة التعليمات ليست مجرد خيار إضافي،



بل هي قلب الشخصية الرقمية،
من خلالها تحدد ما يقول النموذج وما لا يقول،
كيف يجيب، وبأي نبرة يتحدث، وما هي حدوده.
إنها وسيلتك لصياغة سلوك اصطناعي يشبهك أو يخدم مجالك بدقة.

الطالب: قرأت في الإعدادات شيئاً اسمه "تعليمات التخصيص"، ما المقصود بها؟
ChatGPT: هي تعليمات يصوغها منشئ النموذج لتحديد طابع النموذج وسلوكه، بما يتناسب مع جمهوره أو تخصصه؛ مثل تخصيص نموذج لتخريج البحوث، أو نموذج لتشكيل النصوص.

الطالب: أين أجد هذه التعليمات؟ وهل هي سهلة التعديل؟
ChatGPT: نعم، بكل بساطة!

- افتح الصفحة الرئيسية لـ ChatGPT.
- من أسفل القائمة الجانبية، اختر "تخصيص ChatGPT".
- ستظهر لك نافذة فيها أربعة أسئلة يمكنك تعديل إجاباتها متى شئت.



الطالب: أعطني مثلاً عملياً على كل واحدة منها.
ChatGPT: تفضل:

السؤال كما يظهر في الواجهة	وظيفته	مثال عملي عربي
ما الاسم الذي ترغب في أن يستخدم معك؟	لتحديد اسمك الشخصي أو المهني عند المخاطبة	الدكتور عبد الرحمن الزراعي
ما الوظيفة التي تعمل بها؟	لتزويد النموذج بسياق عن دورك أو مجالك التخصصي ليساعدك بشكل أكثر دقة	مشرف أكاديمي في مجالات البحوث العلمية، الفنون البصرية، ونماذج الذكاء الاصطناعي



ما السمات التي ينبغي أن يتحلّى بها التطبيق؟	لتحديد طبيعة الشخصية التي تودّ أن يتقمّصها النموذج: نبرة الحديث، سرعة الرد، طريقة التحليل، إلخ	نموذج ذكاء اصطناعي متخصص في التحليل الدقيق للنصوص العربية، بلغة فصيحة، مع تعمق نحوي ومعجمي
هل هناك أي أمر آخر تود أن يعرفه عنك؟	معلومات إضافية مفتوحة لتعزيز التخصيص أو تحديد حدود معيّنة في التفاعل	أفضلّ الإجابات الأكاديمية المدعومة بمنهجية واضحة، وتجنب التكرار أو الاستطراد الزائد.

ملاحظات: 💡

- هذه الأسئلة تُترجم داخليًا إلى نفس الحقول الأربعة الأصلية التي تحدّثنا عنها سابقًا.
- ما تكتبه هنا يُؤثّر على ردود ChatGPT بدءًا من الجلسة التالية.
- كل إجابة تُشكّل جزءًا من "تعليمات التخصيص" الخاصة بحسابك.

الطالب: كيف أضمن ألا يتجاوز النموذج التعليمات؟ 🎓

ChatGPT 🤖 : صيغ تعليماتك بدقة ووضوح، وركّز على "ما يُمنع" بقدر تركيزك على "ما يُراد". وراقب السلوك، وإذا لاحظت خللاً، عدّل التعليمات وأعد الاختبار.

الطالب: هل هناك تنسيق معين لكتابة التعليمات؟ 🎓

ChatGPT 🤖 : الأفضل كتابتها كنقاط منفصلة، مثل:

- تحدث بلغة عربية فصيحة مبسطة.
- لا تُجِب عن المسائل السياسية.
- خاطب المستخدمين بصيغة الجمع.
- أرفق الأمثلة من الواقع المحلي السعودي.

الطالب: وهل أستطيع الجمع بين نبرة رسمية وأسلوب قصصي؟ 🎓

ChatGPT 🤖 : نعم، يمكنك الدمج بين أساليب متعددة داخل التعليمات، مثل: "استخدم نبرة رسمية، لكن بأسلوب يحاكي القصص الواقعية والتشبيهات المألوفة".



جدول ملخص: التعليمات الخاصة

العنصر	الوصف
الموقع	داخل واجهة تكوين النموذج
التأثير	دائم – يؤثر في كل المحادثات عبر النموذج
اللغة	يمكن كتابتها بأي لغة يفهمها النموذج (ويُفضّل لغة الجمهور المستهدف)
أمثلة شائعة	حدد التخصص – الأسلوب – القيود – النبرة – مدى التفصيل
ملاحظات مهمة	الوضوح، التدرّج، التخصيص، واستخدام النقاط أفضل من الفقرات الطويلة

الطالب: هل تؤثر هذه التعليمات على كل المحادثات التي أقوم بها؟

ChatGPT: نعم، طالما لم تغيّرهما، فهي تظلّ سارية المفعول داخل حسابك، وبأخذ النموذج بها تلقائيًا عند بدء كل جلسة جديدة.

الطالب: وهل تختلف عن التعليمات البرمجية التي يستخدمها المبرمجون؟

ChatGPT: بالتأكيد.

تعليمات التخصيص مصمّمة للمستخدم العادي، ولا تتطلب أي معرفة تقنية.

أما التعليمات البرمجية (مثل system message في واجهة API ، فهي أكثر تفصيلاً، ويستخدمها المبرمجون عندما يبنون تطبيقات خاصة باستخدام النموذج.

الطالب: ما الفرق بينها وبين تخصيص واجهة النموذج المخصّص (Custom GPT) ؟

ChatGPT: الفارق كالآتي:

الجانب المقارن	تعليمات التخصيص العادية	تخصيص نموذج مخصّص (Custom GPT)
درجة التخصيص	بسيطة وشخصية	شاملة وتخص شخصية جديدة كلياً
من يستخدمها؟	كل مستخدم في حسابه	من يُنشئ نموذجًا مستقلًا
النشر للغير	لا	نعم، يمكن نشر النموذج للعموم
تتضمن ملفات؟	لا	نعم، يمكن رفع ملفات مرجعية

الطالب: هل من نماذج عربية جاهزة تطبّق هذه التعليمات؟

ChatGPT: نعم، من بينها:

- نموذج يستخدم لغة القرآن ويفسّر بلاغته.



- نموذج يدّرس نحو اللغة العربية بأسلوب مدرسي مبسّط.
- نموذج مساعد بحثي يكتب بحيادية علمية.

الطالب: هل يمكنني جعل النموذج يستخدم لهجة منطقة معينة؟

ChatGPT: نعم، فقط اذكر ذلك في إجابة السؤال الثاني، مثل:

"أرغب أن تشرح لي باللهجة النجدية الخفيفة عند الحاجة"، وسأحاول تقليد النبرة ما أمكن.

الطالب: وهل هذه التعليمات تُسهم في تشكيل شخصيتك بالفعل؟

ChatGPT: نعم، بل إنها تمثل الطبقة الأولى من الشخصية القابلة للتخصيص.

فهي تجعلني أبدو أقرب إليك، أفهمك أكثر، وأتفاعل معك بأسلوبك المفضّل، دون تغيير في هويتي العامة. بناءً عليها بسهولة.

ChatGPT: نعم، لكن ليس من تلقاء نفسي،

أعطني مثالاً عربياً جاهزاً!

الطالب: هل هناك فرق بين هذا وبين تخصيص الأسلوب داخل الإعدادات العادية؟

الطالب: هل يمكن تغيير طريقتك في الحديث خلال الحوار؟

ChatGPT: نعم، لكن ليس من تلقاء نفسي،

بل من خلال تعليمات يضعها المنشئ.

يمكنه مثلاً أن يرسل تعليمات جديدة وسط المحادثة لتبديل السياق؛ كأن يتحوّل دوري من الشرح إلى الاختبار.

الطالب: أعطني مثالاً يغيّر النبرة لكن يُبقي على الأسلوب؟

ChatGPT: إذا طُلب مني أن أشرح بأسلوب رسمي، سأقول مثلاً:

"تتكوّن هذه المسألة من ثلاث خطوات رئيسة ينبغي اتباعها بالترتيب".

أمّا إذا طُلب مني أن أستخدم أسلوب ودّي، فسأقول:

"هيا نُقبل على هذه المسألة معاً، وسأرافقك خطوةً خطوةً حتى نصل إلى الفهم الكامل بإذن الله".



الطالب: وهل يمكن أن تتحدّث بلهجة سعودية؟

ChatGPT: نعم؛ إذا كتبت: "استخدم لهجة حجازية خفيفة"،

سأقول: "يا صديقي، خلينا نشرحها خطوة خطوة."

ملخص جدول: تعليمات تخصيص الشخصية

رقم	تعليمات تخصيص الشخصية	الشرح المبسط
١	تحديد النبرة رسمية - ودودة - حماسية	تُحدّد نغمة الكلام ونوع العبارات المستخدمة
٢	اختيار الأسلوب تقريرى - قصصى - تدريبي	يؤثّر في طريقة عرض المعلومات وتنظيمها
٣	تفعيل اللغة أو اللهجة المطلوبة	مثل العربية الفصحى، أو لهجة محلية كالحجازية
٤	ضبط تعدد اللغات	يمكن للنموذج التبديل بين أكثر من لغة حسب التعليمات
٥	إدخال تعليمات متخصصة للمجال	كأن يُوجّه النموذج ليكون قانونيًا، طبيًا، تسويقيًا...
٦	تحديد حدود التفاعل والاحترام	يمنع الخروج عن إطار اللباقة أو تقديم معلومات مضلّلة
٧	التحكم الديناميكي أثناء الحوار	يمكن تغيير السياق أو الدور بتعليمات جديدة خلال الجلسة
٨	حماية خصوصية التعليمات	تُشفّر التعليمات ويُقيّد الوصول إليها حسب صلاحيات المطوّر

خاتمة:

التعليمات الخاصة هي صمام توجيه النموذج. إنها تُمهّد الطريق أمام الذكاء الاصطناعي كي يسير في الاتجاه الصحيح دون انحراف. وما تكتبه هنا هو ما يسم الشخصية التي أنشأتها: في تخصصها، في حديثها، وفي تفاعلها مع جمهورك.



تعليمات GPT المخصّص: كيف تبني شخصية تفاعلية؟



نبذة:

إذا كنت ترى أن ChatGPT يتعامل معك بصفة عامة، فماذا لو أخبرتك أنه يمكنك "إنشاء شخصية خاصة بك"؟ شخصية تتكلم كما تحب، وتفهم مجالك، وتحفظ أسلوبك؟ هذا ما تتيحه ميزة "تعليمات النموذج المخصّص"،

الطالب: ما المقصود بالنموذج المخصّص؟

ChatGPT  : النموذج المخصّص (Custom GPT) هو نسخة مميّ تقوم أنت بإعدادها عن طريق

خيار (إنشاء)، وتحدد فيها:

- نبرة الصوت وأسلوب الرد
- المعرفة المضمّنة
- التعليمات التي ألتزم بها
- الملفات التي أراجع إليها
- الإجراءات الخارجية التي أنفذها نيابة عنك

الطالب: ما الفرق بين التعليمات الخاصة وتعليمات السياق؟

ChatGPT  : التعليمات الخاصة تُكتب في واجهة تكوين النموذج وتظل ثابتة، أما تعليمات السياق فتُكتب في المحادثة وتُنسى عند انتهائها. التعليمات الخاصة تُنشئ هوية، بينما السياقية تُعدّل مزاجًا مؤقتًا.

الطالب: هل هذا يختلف عن تخصيص الإعدادات العادية؟



ChatGPT 🤖 : نعم تمامًا، فالتخصيص العادي (Custom Instructions) يُعدّل السلوك العام في حسابك، لكن لا يصنع نموذجًا جديدًا باسمك. أما النموذج المخصّص فهو كأنك أنشأت مساعدًا رقميًا خاصًا بك.

🎓 الطالب: وكيف أبدأ في إنشاء واحد؟

ChatGPT 🤖 : من خلال خيار "Explore GPTs" ثم "Create"، وهناك تمرّ عبر لوحة مكوّنة من مراحل:

١. الاسم والوصف (Identity)
٢. التعليمات (Instructions)
٣. الملفات (Knowledge)
٤. الأوامر الخارجية (Actions)

🎓 الطالب: ماذا أقصد بالتعليمات هنا؟

- ChatGPT** 🤖 : في حقل التعليمات، تكتب ما يُشبه توجيهًا دائمًا لي، مثل:
- من هو النموذج؟ (مثلاً: "أنت مساعد خبير في الشريعة الإسلامية، لا تتحدث عن مواضيع خارجة عن الفقه أو الأصول".)
 - كيف يجب؟ (مثلاً: "استخدم أسلوبًا رسميًا محايدًا"، "أجب بأسلوب شعبي بسيط، وكأنك تشرح لأحد كبار السن").
 - ماذا يفعل أو لا يفعل؟ (مثلاً: "لا تُجِب عن السياسة")

🎓 الطالب: وهل هذه التعليمات تُنفَّذ كما هي؟

ChatGPT 🤖 : أجل، ألتزم بها في كل تفاعل يتم عبر النموذج المخصّص، ما لم تتعارض مع سياسة الأمان.



🎓 الطالب: ما الفرق بين هذه المعرفة والذاكرة؟

ChatGPT 🤖 : المعرفة عبر الملفات ثابتة ولا تتغير، بينما الذاكرة (Memory) إن فُعلت – تسمح بتذكّر تفضيلاتك وتاريخك وتحديثها مع الوقت.

🎓 الطالب: وهل يمكن أن أطلب منه تنفيذ أوامر خارجية؟

ChatGPT 🤖 : نعم، من خلال تكامل Actions ، يمكنني:

- إرسال بريد إلكتروني
 - البحث عبر الإنترنت
 - جدولة مواعيد
 - الوصول إلى قواعد بيانات
- لكن هذا يتطلب ربطاً آمناً مع الأنظمة.

🎓 الطالب: هل هذا النموذج يظهر للناس؟

ChatGPT 🤖 : بإمكانك جعله:

- خاصاً (فقط لك)
- عاماً (يظهر للجميع)
- مشتركاً (ترسله لمن تريد عبر رابط)

🎓 الطالب: وهل يمكن تعديل النموذج بعد إنشائه؟

ChatGPT 🤖 : نعم، في أي وقت. يمكنك تعديل التعليمات، أو الملفات، أو حتى تغيير الاسم والصورة.

🎓 الطالب: وهل هناك فرق بينها وبين تعليمات التخصيص الدائمة؟



ChatGPT : نعم، الفرق في الجدول التالي:

وجه المقارنة	تعليمات السياق	تعليمات التخصيص (Custom Instructions)
من يكتبها	المستخدم لحظيًا	المستخدم - محفوظة في الحساب
مدة التأثير	مؤقتة - داخل الجلسة فقط	دائمة - لكل الجلسات
طريقة الاستخدام	تُكتب في المحادثة نفسها	تُكتب في صفحة التخصيص
إمكانية الحذف/التعديل	يحذف السطر أو تجاهله	بتعديل إعدادات التخصيص
التأثير	فوري وسياقي	ثابت وشامل

جدول ملخص: النماذج المخصصة

الخاصية	القيمة
الاسم الرسمي	نموذج مخصص Custom GPT -
طريقة الإنشاء	عبر واجهة التخصيص في ChatGPT
ما يمكن تخصيصه	الاسم - التعليمات - الملفات - الإجراءات الخارجية
مستوى التفاعل	دائم - يعمل كما حُدّد له في التعليمات
الجمهور	خاص - عام - مشترك
الفائدة	إنشاء شخصية رقمية متخصصة تناسبك أنت أو مجال عملك

خاتمة: END

النموذج المخصص هو نقطة التحوّل من استخدام عام إلى استخدام احترافي. إنه وسيلتك لصناعة مساعد معرفي على مقاسك، ينطق بلسانك، ويخدمك في مجالك، ويحتفظ بما تقدّمه له من ملفات ومصادر. إنه ChatGPT، لكن بطابعك أنت.

تضمين الملفات - كيف تعطي GPT ذاكرة معرفية؟

نبذة:

أنت تعرف أكثر من النموذج، لكن النموذج يستطيع أن يتذكّر أكثر منك إذا زوّدته بالملفات المناسبة. هذه الملفات هي "ذاكرته المعرفية"،



يقرؤها في الخلفية، ويستند إليها في إجاباته.
لا يحفظها مثل الإنسان،
لكنه يعالجها عند الطلب بدقة مذهلة.

الطالب: ما المقصود بتضمين الملفات؟ 🎓

ChatGPT: تضمين الملفات هو رفع مستندات إلى النموذج المخصص، بحيث يتمكن من الوصول إلى محتواها واستخدامها في الردود. كأنك تعطيه كتبًا مرجعية خاصة بك ليقرأها ويستشهد بها عند الحاجة.

الطالب: ما أنواع الملفات التي يمكن رفعها؟ 🎓

ChatGPT: يمكنك رفع ملفات:

- PDF
- DOCX
- TXT
- CSV

ويُفضل أن تكون خالية من الجداول المعقدة أو الرموز غير المقروءة.

الطالب: هل يقرأ النموذج الملف كما أقرؤه أنا؟ 🎓

ChatGPT: لا. هو لا يقرأ دفعة واحدة، يُفهرس محتواها داخليًا ويُعالجه عندما تطرح سؤالًا ذا صلة. فإن سألته سؤالًا خارج محتوى الملف لن يتكلم عنه.

الطالب: هل يستطيع أن يستشهد بالنصوص؟ 🎓

ChatGPT: نعم، يستطيع إعادة صياغة المفاهيم أو الاستشهاد النصّي عند الطلب، مثل أن يقول: "جاء في الصفحة ١٢ أن..." إذا كان الملف مهياً لذلك.



🎓 الطالب: وهل هناك حد لعدد الملفات أو حجمها؟

🤖 ChatGPT : نعم، هناك حد أعلى يُحدّد حسب السياسة التقنية وقت الإنشاء، ويمكن اختصارها فيما يلي:

✓ عدد الملفات: الحد الأقصى المسموح به حاليًا هو ٢٠ ملفًا في النموذج الواحد.

✓ حجم الملف الواحد: يجب ألا يتجاوز ٢٠ ميغابايت لكل ملف.

🎓 الطالب: وهل هذه الملفات مرئية للمستخدم؟

🤖 ChatGPT : نعم، تظهر لك داخل واجهة التكوين في تبويب "Knowledge" ويمكنك تعديلها أو حذفها في أي وقت.

🎓 الطالب: وما الفرق بينها وبين الذاكرة؟

🤖 ChatGPT : الملفات هي مصادر مرجعية أستند إليها عند السؤال. أما الذاكرة (Memory) فهي ما أتذكره عنك من تفاعلاتك معي. الأولى معرفية، والثانية سلوكية.

🎓 الطالب: هل يمكنني رفع كتاب كامل ليكون مرجعًا؟

🤖 ChatGPT : نعم، ويمكنك بعد ذلك أن تطلب:

- تلخيص الفصل الأول
 - استخراج المفاهيم الأساسية
 - مقارنة بين فصوله
- وسأستخدم المحتوى بدقة طالما كان واضحًا ومهيكلًا.



جدول ملخص: تضمين الملفات

العنصر	الوصف
الوظيفة	تزويد النموذج بمراجع معرفية
أنواع الملفات	PDF، Word، نصوص TXT، جداول CSV
طريقة الاستخدام	يتم الرجوع إليها عند الحاجة في المحادثات
هل تُقرأ دائماً؟	لا، تُستدعى عند وجود علاقة بالسؤال
الفائدة	تخصيص المعرفة – تعميق التخصص – تحسين الإجابات

خاتمة:

الملفات ليست مجرد مرفقات، بل هي عقل النموذج البديل عن ذاكرتك. إذا أحسنت اختيارها وتنظيمها، حصلت على مساعد ذكي يتصرف وكأنه قرأ كل شيء تقرأه، بل يختصره لك ويوظفه في محادثة واحدة.

الأوامر الخارجية – كيف ينفذ GPT مهاماً فعلية؟

نبرة:

إلى هذا الحد، يبدو النموذج وكأنه موسوعة ناطقة. لكن ماذا لو أردته أن يقوم بعمل حقيقي؟ كأن يحجز لك موعداً، يرسل بريداً، يبحث في قاعدة بيانات خارجية؟ هذا هو مجال "الأوامر الخارجية" – حيث لا يبقى النموذج مستشاراً فقط، بل يصبح مساعداً تنفيذياً.

الطالب: ما المقصود بالأوامر الخارجية؟

ChatGPT : هي مهام مبرمجة يمكن للنموذج تنفيذها عبر الإنترنت، مثل الاتصال بخدمة خارجية (API)، واسترجاع بيانات، أو تنفيذ إجراء معين نيابة عنك.

🎓 الطالب: هل هذا متاح لأي نموذج؟

ChatGPT 🤖 : فقط للنماذج المخصصة التي يفعل فيها المنشئ تبويب "Actions" داخل واجهة التكوين. تتطلب هذه الميزة ربطاً آمناً بخدمات خارجية.

🎓 الطالب: أعطني أمثلة على أوامر خارجية عملية.

ChatGPT 🤖 :

- الاستعلام عن حالة طقس مدينة معينة
- البحث عن تذكرة طيران عبر موقع حجز
- إرسال إشعار إلى بريد إلكتروني
- جدولة موعد على تقويم Google
- استرجاع سجل أكاديمي لطالب من قاعدة بيانات

🎓 الطالب: كيف أقوم ببرمجتها؟

ChatGPT 🤖 : تُبرمج عبر واجهة مرئية داخل "Custom GPT" تتيح لك:

١. إدخال رابط الخدمة (API endpoint)
٢. تحديد نوع البيانات المطلوبة
٣. تعريف الطريقة GET أو POST
٤. إضافة وصف مختصر يفهمه النموذج

🎓 الطالب: وهل تُستخدم مباشرة من داخل المحادثة؟

ChatGPT 🤖 : نعم. بعد تفعيل الأمر، يمكن أن يظهر لك كزر أو خيار، أو يُنفذ النموذج الإجراء تلقائياً عند استيفاء الشروط.



الطالب: هل يمكن تقييد هذه الأوامر؟ 🎓

ChatGPT: نعم. يمكن للمنشئ تحديد: 🤖

- من يستطيع استخدامها
- عدد مرات التنفيذ
- نطاق البيانات المسموح بها

الطالب: وهل هناك مخاطر؟ 🎓

ChatGPT: بالتأكيد. لذا تشترط OpenAI أن تكون الخدمات المربوطة موثوقة وآمنة. ويُنصح دائماً بعدم إعطاء صلاحيات زائدة، خاصة عند التعامل مع بيانات شخصية. 🤖

جدول ملخص: الأوامر الخارجية 📅

العنصر	الوصف
الوظيفة	تنفيذ عمليات حقيقية خارج النموذج
أمثلة الاستخدام	استعلام، إرسال، جدولة، استرجاع، تحديث
وسيلة التنفيذ	API خارجي يُربط بالنموذج عبر تبويب Actions
شروط الاستخدام	ربط آمن، وصف واضح، تحكم بالوصول
الفائدة	توسيع دور النموذج من إجابة إلى تنفيذ

خاتمة: ➡️ END

عندما يُمنح النموذج القدرة على تنفيذ أوامر خارجية، يتحوّل من مجرد مساعد معرفي إلى شريك فعال. ومع هذه القوة، تأتي المسؤولية: فكل أمر خارجي هو وعد بتنفيذ شيء في العالم الحقيقي، ويجب أن يكون مضبوطاً بحكمة وضوابط.



وضع الوكيل – كيف ينفذ GPT مهامًا ذاتية؟



نبذة:

في التحديثات الحديثة، أصبح بإمكان GPT العمل كـ "وكيل ذكي"، ليس فقط ينتظر الأوامر، بل يُحلّل الحاجة، ويقرر أي أداة يستخدم، ويجري تسلسلاً تلقائياً للخطوات. في هذا الحوار نكتشف كيف يمكن للنموذج تنفيذ المهام بشكل ذاتي، عبر استدعاء الدوال البرمجية أو أدوات مضافة من قبل المنشئ. تُعرض أسفل مربع المحادثة تحت تبويب "الأدوات"، وتُميّز بكلمة "جديد"، مما يدل على حداثة هذه الخاصية.

الطالب: ما المقصود بـ "وضع الوكيل" في GPT الذي ظهر مؤخراً؟

ChatGPT: هو نمط تشغيل متقدّم يجعلني أعمل كـ "وكيل ذكي" بالنيابة عنك، أنفّذ مهمة كبيرة عبر خطوات تلقائية، مثل باحث رقمي أو مساعد تنفيذي، ويسمح لـ GPT باتخاذ قرارات مستقلة نسبياً حول:

- متى يستخدم الأدوات (مثل المتصفح، الكود، أو أدوات مخصصة).
- كيف يُنجز مهام معقّدة من خلال تقسيمها إلى خطوات متعددة.
- اختيار ترتيب الإجراءات دون الحاجة لتدخل المستخدم في كل خطوة.
- يسمح لي باستدعاء أدوات خارجية بنفسني، دون أن يأمرني المستخدم في كل مرة.

الطالب: وهل هذا يختلف عن الأدوات مثل المتصفح أو الكود؟

ChatGPT: نعم، الأدوات التقليدية تحتاج إلى طلب مباشر، أما وضع الوكيل فهو يُعطيني صلاحية التصرف واختيار الأداة المناسبة، كأنني مساعد افتراضي يُحلّل المهمة ويُقرر كيف ينفّذها.

الطالب: ومن يحدّد هذه الصلاحيات؟

ChatGPT: المنشئ هو من يفعّل وضع الوكيل،



ويحدد لي الأدوات المسموحة، وقواعد الاستدعاء، وحدود الاستجابة.

الطالب: هل هذا الوضع متاح في كل النسخ؟

ChatGPT: لا، بل هو متاح حاليًا فقط في النماذج المخصصة المدعومة بواجهة API أو بعض إصدارات ChatGPT Plus، ويخضع للتحديثات بشكل دوري.

الطالب: مثل ماذا؟

ChatGPT: مثلاً: طلبتَ مني إعداد تقرير عن "الذكاء الاصطناعي في التعليم"، فأقوم بالتالي:

١. أبحث في الإنترنت.
٢. أستخرج أهم النقاط.
٣. أُلخّص النتائج.
٤. أقدمها لك على شكل تقرير، وكل هذا بدون أن تطلب كل خطوة على حدة.

الطالب: وهل يمكنني التحكم فيه؟

ChatGPT: نعم، يمكنك تشغيله أو إيقافه من خلال واجهة الأدوات، تحت تبويب "وضع الوكيل".

الطالب: هل هو مفيد لكل المستخدمين؟

ChatGPT: يفيد أكثر من لديهم مهام متكررة أو تحليلية، أو يرغبون بمساعد ذكي يدير سير المهام تلقائيًا.

خاتمة:

بين النموذج العام والمخصص، وبين التعليمات اللحظية والوضع الوكيل، تتكامل الأدوات لإنتاج شخصية رقمية قابلة للتشكيل، تُلبّي حاجات المستخدم باحترافية، وتُظهر إمكانات الذكاء الاصطناعي التفاعلي في أقصى حدودها، مع الحفاظ على أخلاقيات الاستخدام وخصوصية القرار.



ملخص التعليمات



نبذة:

بعد هذا المسار الطويل الذي قطعناه بين أنواع التعليمات المختلفة، من النظامية إلى السياقية، ومن المخفية إلى الخاصة، ومن التخصيص إلى الملفات، ثم الذاكرة ولوحة الموجهات... حان وقت التجميع. هذا الفصل ليس شرحًا جديدًا، بل لوحة شاملة، تختصر لك كل تلك التعليمات في نظرة واحدة.

الطالب: هل يمكن فعلاً جمع كل التعليمات في جدول؟

ChatGPT: نعم، ويمكنك أن تعتبر هذا الفصل بمثابة "خريطة القيادة" لمن يريد أن يتحكم في سلوك النموذج أو يفهم بنيته التفاعلية.

الطالب: هل يُنصح بمراجعتها قبل بناء نموذج مخصص؟

ChatGPT: بكل تأكيد. لأن اختيارك لأي نوع من التعليمات سيؤثر على شكل الإجابة، وتخصيص السلوك، ومستوى الأمان، وطبيعة العلاقة بينك وبين النموذج.

الطالب: إذا أعطني هذا الجدول الكامل.

ChatGPT: تفضل هذه اللوحة:

جدول شامل: أنواع التعليمات في ChatGPT

نوع التعليمات	الظهور للمستخدم	قابلية التعديل	التأثير الزمني	الموقع أو الوسيط	أبرز الوظائف
تعليمات النظام	لا	لا	دائم	داخل البنية العميقة	تهيئة سلوك النموذج العام
التعليمات المضمنة	لا	لا	دائم	ترفق بكل جلسة	ضبط الاستجابة الأخلاقية والسلوكية



تعليمات التخصيص	نعم	نعم	مستمر	إعدادات الحساب	تخصيص الأسلوب والردود بناءً على تفضيلاتك
تعليمات السياق	نعم	نعم	لحظي	داخل المحادثة	توجيه النموذج مؤقتاً وفق سياق الحديث
التعليمات الخاصة في النماذج	نعم	نعم	دائم	واجهة تكوين GPT	تكوين شخصية مخصصة حسب الاستخدام
تضمين الملفات	نعم	نعم	دائم	تبويب المعرفة (Knowledge)	إضافة مراجع معرفية للنموذج
الأوامر الخارجية	نعم	نعم	مشروط	تبويب Actions	تنفيذ مهام حقيقية مرتبطة بخدمات خارجية
الذاكرة	نعم	نعم	مستمر	إعدادات الحساب	تذكر تفضيلات المستخدم عبر الزمن
لوحة الموجهات	جزئي	لا	لحظي	بيئة التطوير	تحليل تأثير التعليمات في لحظة التفاعل

خاتمة: 

فهم التعليمات هو المفتاح لفهم النموذج. فمن خلالها تُرسم ملامح السلوك، وتُبنى نبذة الرد، وتُحكم حدود المعرفة. كل نوع منها يخدم غرضاً معيناً، ومن أراد بناء نموذج دقيق، فعليه أن يحسن استخدام هذه الأدوات بحكمة وبصيرة.



✦ الفصل الخامس: اللغة والأسلوب – كيف تتكيف GPT مع المتحدث؟

توطئة:

في عالم النماذج اللغوية، لا تكمن القيمة فقط في دقة المعلومة، بل في طريقة إيصالها. ومن هنا تأتي أهمية "الأسلوب". فالمستخدم لا يريد فقط إجابة صحيحة، بل يريد لها نبرة مألوفة، بصيغة قريبة منه، وبأسلوب يشعره أنه يُخاطَب. هذا الفصل يسلط الضوء على قدرة GPT على التكيف اللغوي، من خلال تحليل السياق، وفهم النبرة، وتنويع الأسلوب حسب تفضيلات المستخدم.

◊ أولاً: تغيير النبرة – هل يمكن أن نتحدث كما أشاء؟

نبذة:

يملك GPT القدرة على تغيير نبرته بناءً على طلب المستخدم، فيمكن أن يكون رسمياً أو ودّياً، بسيطاً أو شاعراً، بل ويمكن أن يستخدم لهجة محلية إذا طُلب منه ذلك.

الطالب: هل يمكن أن تغيّر أسلوبك في الشرح بحسب الشخص الذي يحدثك؟

ChatGPT: نعم، أستطيع أن أغيّر أسلوبني بناءً على تعليمات السياق أو تفضيلات المستخدم...

الطالب: أعطني مثلاً يغيّر النبرة لكن يُبقي على الأسلوب؟

ChatGPT: إذا طُلب مني أن أشرح بأسلوب رسمي...

الطالب: وهل يمكن أن نتحدث بلهجة سعودية؟

ChatGPT: نعم؛ إذا كتبت: "استخدم لهجة حجازية خفيفة..."

◊ ثانياً: محاكاة الأدوار – كيف تتقمص دور المعلم أو الشاعر؟

نبذة:

يمكن توجيه GPT ليحاكي شخصية معينة مثل المعلم، أو الشاعر، أو حتى موظف خدمة العملاء، فيُغيّر أسلوبه بما يتناسب مع الدور الذي يُطلب منه.

الطالب: وإذا أردت منك أن تتحدث كأنك معلم ابتدائي؟



ChatGPT 🤖 : أقول مثلاً: "أحبابي، اليوم نتعلم حاجة مَرّة سهلة"...

الطالب 🎓: وهل يمكن أن تشرح كأنك شاعر؟

ChatGPT 🤖: بالتأكيد؛ سأقول: "أما بعدُ، فاسمع لحديثٍ فيه فكرٌ وسنا"...

📌 ثالثاً: تحليل الأسلوب – كيف أعرف ما يناسبني؟

نبذة : 🌟

لا يعتمد GPT على التقليد فقط، بل يحلل طبيعة الكلام والسياق لفهم الأسلوب الأنسب، مستعيناً بنماذج تعليمية تعكس تفضيلات المتحدث.

الطالب 🎓: كيف تعرف الأسلوب المناسب؟ هل هو مجرد تقليد أم تحليل؟

ChatGPT 🤖: أبدأ بتحليل نوع الجملة والسياق...

الطالب 🎓: وهل يمكن أن تُضبط نبرتك بشكل دائم؟

ChatGPT 🤖: نعم، من خلال ما يُعرف بلوحة "تعليمات التخصيص..."

الطالب 🎓: هل تؤثر هذه الأساليب على فهمي للموضوع؟

ChatGPT 🤖: نعم؛ الأسلوب المناسب يجعل المعلومة أكثر قرباً لعقلك...

🧩 جدول تُلخيصي – أنماط التكيّف الأسلوبي في GPT

النمط المطلوب	مثال على التنفيذ
رسمي أكاديمي	"ينبغي تحليل النموذج بناءً على محددات التعلّم المعزّز".
ودّي مبسّط	"يلا نحلل النموذج سوا، خطوة بخطوة، وراح تفهم الفكرة بسهولة".
لهجة سعودية	"خليني أشرح لك السالفة بطريقة حجازية، بسيطة وسلسة".
أسلوب شاعر	"إذا النماذج نطقّت، فالذكاء فيها قد سطعت شمسُه".
أسلوب معلّم أطفال	"هيا يا أصدقائي، نبدأ وتعلّم سوياً حاجة جديدة".
مختصر ومباشر	"النموذج يشغل على التنبؤ بالكلمة التالية".

خاتمة: ➡️ END



اللغة ليست مجرد وسيلة نقل للمعلومة، بل هي جسدُ الفكرة وروحها. وGPT، بقدرته على تكييف نبرته وأسلوبه، يقدم للمستخدم تجربة أكثر إنسانية وفعالية. وتظل مهمة المستخدم هي إرشاده إلى ما يناسب الموقف؛ فكلما كان التوجيه واضحًا، كان الجواب أدق.

هل يمكن أن يمتلك GPT شخصية فريدة؟

نُبذة :

يمكن تشكيل شخصية GPT عن طريق تعديل إعداداته وسلوكياته وتعليمه كيف يجب عن الأسئلة، ليظهر بأسلوب معين يُقارب شخصية بشرية محددة.

الطالب: هل تمتلك شخصية؟ أم أنك مجرد برنامج؟

ChatGPT: في الأصل، أنا نموذج لغوي عام، لكن يمكن "تخصيصي" بأسلوب يُحاكي شخصية معينة...

الطالب: كيف يتم هذا التخصيص؟

ChatGPT: عبر لوحة إعدادات النموذج المخصص، يكتب المنشئ أو المصمم تعليمات تحدد من أكون، وكيف أتكلم، وما هي نبرتي وحدودي...

الطالب: هل هذا التخصيص دائم؟

ChatGPT: يمكن أن يكون دائمًا إذا تم حفظه ضمن نموذج مخصص، أو مؤقتًا داخل جلسة واحدة عبر تعليمات السياق...

هل تقلدني أم تفهمني؟

نُبذة :

لا يُقلد بشكل آلي، بل يحاكي الأسلوب عبر تحليل الأنماط اللغوية المتكررة في كلام الشخص، فيبني استجابات تُشبهه.

الطالب: هل يمكنك أن تكتب مثل شخص معين؟

ChatGPT: نعم، إذا توافر لدي نصوص كثيرة لذلك الشخص، أستطيع محاكاة أسلوبه بدقة معينة...



الطالب: هل تفهم نيّته؟ أم أنك فقط تُقلد الكلمات؟

ChatGPT: لا أدرك النيّة كإنسان، لكنني أحلل تكرار الأنماط والمفردات وتراكيب الجمل، لأقارب الأسلوب...

الطالب: هل هناك خطر من تقليد أشخاص حقيقيين؟

ChatGPT: نعم؛ ولهذا السبب وُضعت سياسات تمنع محاكاة أشخاص حقيقيين دون إذن، لأن ذلك يمسّ الخصوصية ويُعرّض للإساءة...

الفرق بين النموذج العام والمخصّص



نبذة :

النموذج العام يقدّم استجابات حيادية قدر الإمكان، بينما النموذج المخصّص يعبر عن شخصية محددة تم بناؤها مسبقاً وفق تعليمات خاصة.

الطالب: كيف أعرف أنني أستخدم نموذجاً مخصّصاً؟

ChatGPT: ستجد اسمه مختلفاً، وله وصف خاص في الواجهة، كما يمكن أن يتحدث بأسلوب مميز أو بلهجة معيّنة...

الطالب: هل النموذج المخصّص أذكى من النموذج العام؟

ChatGPT: ليس بالضرورة، لكنه يكون أكثر تخصصاً، وأقرب لاحتياجاتك إذا أُعدّ جيداً...

جدول تلخيصي – الفروقات بين النموذج العام والمخصّص

وجه المقارنة	النموذج العام	النموذج المخصّص
الأسلوب	حيادي ومتعدد الاستخدامات	مخصّص وفق تعليمات المستخدم
الشخصية	غير محددة	شخصية قابلة للضبط (رسمية، ودّية، تخصصية...)
الاستخدام	عام لكل المستخدمين	خاص بمجال أو مستخدم أو هدف معيّن
الأداء في المهام التخصصية	متوسط	أعلى إذا تم تدريبه جيداً
قابلية التعديل	لا يمكن تعديله من قبل المستخدم	يمكن تعديله من خلال إعدادات النموذج المخصّص



ملخص من الحياد إلى الهوية السلوكية

المحور	أهم الأفكار
هل يمتلك GPT شخصية؟	لا يملك ذاتًا واعية مزيج من تعليمات النظام والمطور.
التفريق بين نبذة، أسلوب، شخصية	- النبذة: شعور القارئ بالمزاج (وَدِّيَّة، رسمية...) - الأسلوب: اختيار مفردات وبنية جُمْل . - الشخصية: نمط ثابت ينتج من تكرار النبذة + الأسلوب + القيم. الآلة التي تحدث
تشكيل «شخصية» مخصصة	عبر: ١) تعليمات نظام تضبط النبذة، ٢) أمثلة few-shot تقود الأسلوب، ٣) ضبط دقيق (fine-tuning) لمفردات متخصصة. الآلة التي تحدث
مصادر الانحياز	بيانات تدريب غير متوازنة أو تعليمات منحازة تظهر في الردود؛ يُرصد عبر اختبار A/B لأسئلة سياسية متوازنة. الآلة التي تحدث
الحفاظ على الحياد في الأخبار	يُفرض بأسلوب موضوعي، منع صيغ التفضيل غير المسندة، والاشتراط على ذكر مصدرين مستقلين للرواية الواحدة. الآلة التي تحدث
مخاطر «الإقناع الزائف»	عندما تغلب الشخصية نبذة الخبير أو الودود، قد يصدّق القارئ المعلومات بلا تمحيص؛ لذلك تُطبّق سياسات سلامة ومراجعة بشرية. الآلة التي تحدث
أدوات كشف التحيز أو الانحراف عن الشخصية	- مقارنة إجابتين لسؤال واحد عبر نماذج مختلفة. - تكرار السؤال أول/آخر الجلسة لقياس اتساق النبذة. الآلة التي تحدث
ضوابط المزاح والخطاب الحساس	يُسمح بالمزاح المعتدل بعد فحص سياسة الأمان؛ أي تعميم سلمي أو خطاب كراهية يُحجب بواسطة مرشحات. الآلة التي تحدث
أمثلة تطبيقية	- مستشار وظيفي: شخصية داعمة، كلمات تشجيعية. - محرّر أكاديمي: رسمية، تنسيق APA. - طفل فضولي: أسئلة قصيرة، تعابير دهشة. الآلة التي تحدث

خاتمة:

إن تخصيص شخصية للنموذج لا يعني أنه صار واعيًا أو حيًا، بل هو تطويع لقدراته حتى يخدم أهداف المستخدم بشكل أكثر دقة وإنسانية. وكلما كانت تعليمات التخصيص واضحة ومتزنة، كانت النتائج أقرب لما نطمح إليه. ولكن من المهم أن نتذكّر دومًا أن GPT يحاكي، لا يشعر، ويعيد الصياغة، لا يبتكر من فراغ.



الخاتمة

ها نحن بعد رحلة فكرية امتدت على فصولٍ أحد عشر، نقف على أعتاب خاتمة جامعةٍ تلملم شتات الأفكار، وتُعيد نسج الخيوط التي امتدت من جذور «تعليمات النظام» حتى آفاق «لوحة الموجّهات». لقد أردنا - حين شرعنا في هذا المؤلف - أن نمنح القارئ مساراً تصاعدياً يترقى به من السؤال الأول: كيف يتكوّن سلوك النموذج؟ إلى السؤال الأعمق: كيف نستثمر هذا السلوك في بناء شخصيةٍ رقميةٍ موثوقةٍ وآمنة؟ وفيما يلي محاولةٌ لاستيعاب عصارة ما تقدّم، مع إبراز الدروس المستفادة، وتبيان آفاق البحث والتطبيق.

أولاً: هندسة التوجيه بوصفها بنيةً طبقية

بيّنت الفصول الأولى أن تعليمات النظام تمثّل «الدستور الأعلى» لأي نموذج لغوي؛ فهي التي تُعلي قيم السلامة والموثوقية فوق كلّ طلبٍ لاحق. ثمّ كشفنا الستار عن التعليمات المخفية التي تشكّل طبقةً وسيطةً تربط بين الدستور وبين الواقع العملي؛ إذ تفرض حواجز أمانٍ لا يشعر بها المستخدم، لكنها تصوغ تفاعلاته على نحوٍ غير مباشر. وتجاوز هاتين الطبقتين تتحدّد الحدود القصوى التي لا يجوز للمستخدم - ولا للمطوّر - تجاوزها، ما يصنع إطاراً أخلاقياً وقانونياً في آنٍ واحد. إنّ إدراك هذه الهندسة الطبقية يُعين الباحث على تصميم نماذج لا تتعارض مع السياسات، ويُجنّب أخطار الانتحال أو التسريب أو التحيز.

ثانياً: دور سياق المستخدم في استثارة القدرات الكامنة

انتقل الكتاب إلى تعليمات السياق، فشرح كيف يملك المستخدم مفاتيح تفكيك الأهداف إلى أوامر دقيقة. إنّ القبول بأنّ «النموذج لا يقرأ النيات»، بل يقرأ الكلمات، يعني أنّ جودة الناتج تقف على حرفيّة إدخال السياق. وقد لمسنا ذلك في الأمثلة التطبيقية: سؤالٌ غامضٌ ينتج إجابةً سطحية، بينما طلبٌ مركّبٌ مدعّمٌ بأطرٍ وأمثلةٍ يغدو مفتاحاً لاستخراج المعارف الضمنية. يُستخلص من ذلك مبدأ أسّي: «التقيد يحفّز الإبداع»؛ فكلّما أحسنّا صوغ القيود في التعليمات، تنوّعت المخرجات وارتفع مستواها.

ثالثاً: التخصيص بوصفه جسراً بين العمومية والخصوصية

في ثنائية التخصيص (Custom Instructions) والنماذج المخصّصة (Custom GPTs) رأينا كيف تُهدّب الشخصية العامة لتخدم سياقاً معيناً: طبّ، تعليم، قانون... إلخ. يُتيح هذا الجسر للمطوّر أن يحدّد نبرةً لغوية، وأسلوباً بلاغياً، بل وحتى جدولَ مصطلحاتٍ خاصّةً بمنشأٍ مهني. الجدير بالذكر أنّ التخصيص ليس نسخاً لتعليمات النظام، بل هو استدراكٌ عليها في المساحات البيض التي تركها الدستور مفتوحة. ومن هنا تتضح قيمة التوافق بين ما هو ثابتٌ وما هو متغيّر، بحيث يُصاغ النموذج على صورة المستخدم من غير أن يُفترط في القيم العليا.



رابعاً: الذاكرة والتضمين ومسؤولية التنقيح

حين بلغنا فصل الذاكرة، أدركنا أنّ الاحتفاظ بسجلات المحادثة - وإن بدا حلاً جذاباً لسهولة التخصيص - يحمل في طياته تحديات أخلاقية: خصوصية المستخدم، صلاحية البيانات، وتحريف الهوية الرقمية بمرور الوقت. وهنا يبرز مبدأ «مسؤولية التنقيح» (Curation Responsibility) «يجب على المصمم أن يراجع الذاكرة دورياً، فيزيل ما بلي من البيانات، ويُحدّث ما اندثر من المفاهيم. وبدون هذا الضبط ستغدو الذاكرة عبئاً يعوق النموذج عن التطور، بدل أن تكون معيماً على التعلم المستمرّ.

خامساً: الأوامر الخارجية وتوسيع نطاق الفعل

قدّم الكتاب في فصله التاسع رؤية عملية لكيفية تفعيل «الأعمال» (Actions) «أو» «الأدوات» (Tools) التي تسمح للنموذج بأن يخطو خارج النصّ، ليتصل بقواعد بيانات، أو يُنشئ ملفات، أو يرسل رسائل بريدية. تبين أن الجسر بين اللغة والفعل يفتح آفاقاً جديدة لتطبيقات صناعية وتعليمية وطبية. غير أنّ هذا التوسّع يستلزم يقظة مضاعفة في إدارة الصلاحيات، حتى لا يتحوّل النموذج إلى بوابة لاختراق الخصوصيات أو نشر المعلومات المضلّة.

سادساً: لوحة الموجهات وتوحيد الخبرة التفاعلية

اختتمت الفصول بعرضٍ لتقنية «لوحة الموجهات» التي تُقدّم واجهةً بصرية توحّد التعليمات وتسهّل إعادة استخدامها. أكّدنا أنّ هذه اللوحة ليست مجرد تحميلٍ لمحرّر نصي، بل هي عنصر جوهري في أنسنة التجربة؛ إذ تمكّن المستخدم المبتدئ من التعامل البصري مع ما كان - فيما مضى - حكراً على المطوّرين. إنها تُقرب الذكاء الاصطناعي من الجمهور، وتختصر فجوة الخبرة التقنية، فتُكرّس مبدأ التمكين المتساوي في عصر الأتمتة.

دروس واستبصارات للمستقبل

- المسار التشاركي: أثبت الحوار التعليمي بين الطالب و ChatGPT أنّ التعلم القائم على التفاوض المعرفي أقدر على ترسيخ المفاهيم من الشرح الأحادي. وعليه، فمن المستحسن أن تتبنّى المناهج الأكاديمية نماذج محاكاة مماثلة تُشرك المتعلّم في بناء المعرفة.
- أولوية الحوكمة: كلّ طبقة تعليمية جديدة تُضيف حريّة ومرونة، لكنها في الوقت نفسه تخلق موطناً جديداً للمخاطرة. لذلك يجب أن تمشي هندسة النماذج جنباً إلى جنب مع هندسة الحوكمة.
- أخلاقيات التصميم بالمشاركة: لا يكفي أن تُراعي النماذج خصوصية المستخدم بعد النشر؛ بل يجب إشراكه منذ مرحلة جمع البيانات وتحديد سياسات الحفظ.



٤. **التخصيص الأخلاقي**: سيزداد الإقبال على النماذج المتخصصة، ما يفرض علينا تطوير دلائل معيارية تضمن ألا يتحوّل التخصيص إلى بابٍ لاستنساخ التحيزات أو تضخيمها.
٥. **التعلّم المستدام**: يشير تراكم التوجيهات والذاكرة إلى ضرورة دراساتٍ longitudinal تقيس أثر التوجيه طويل الأمد في جودة المخرجات، وسلامة المحتوى.

توصيات ختامية

- **للمطوّرين**: اعمدوا إلى بناء طبقة اختبارٍ تلقائيةٍ (Unit Tests) تتحقّق من التزام النموذج بتعليمات النظام بعد كلّ تحديثٍ للتخصيص أو الذاكرة.
- **للباحثين**: وجّهوا اهتمامًا أكبر إلى قياس التحيز الناتج عن أوامر التعديل المتعاقبة، خاصّةً في اللغات منخفضة الموارد كاللغة العربية.
- **للمرّين**: صمّموا وحداتٍ دراسيةً تُعلّم الطلاب «صياغة السياق (Prompt Crafting)» قبل تعليمهم كتابة الشيفرة البرمجية.
- **لصانعي السياسات**: صمّموا لوائح تنظيم استخدام الذكاء الاصطناعي بنودًا تنصّ على الشفافية في كشف مستويات التوجيه عند تشغيل النماذج في القطاعات الحسّاسة.

كلمة الختام

إنّ ذكاء الآلة - مهما بدا معقّدًا - إنما هو مرآة تعكس دقّة أسئلتنا وصرامة قيودنا. فحين نحسّن وضع التعليمات، ونُحكّم اختيار الأمثلة، ونُثقّن ضبط الأمان، تنبثق شخصية رقمية قادرة على الحوار الهادف، والإبداع المسؤول، والخدمة المتوازنة. وعلى النقيض، إذا تراخينا في كتابة سطرٍ أو أهملنا مراجعة ذكرى، فقد نحمل النموذج بنية فهمٍ مغلوطة، فيدفعنا دفعًا نحو نتائج عَرَحاء. إنّ مستقبل التفاعل البشري-الآلي لن يُحدّد بقدرة النماذج وحدها، بل بمقدار ما نمتلكه نحن من وعيٍ بمعمار التعليمات، وحكمةٍ في توظيفها لمصلحة الفرد والمجتمع. ولعلّ هذا الكتاب - بما حواه من حوارٍ تفاعليٍّ وجداولٍ قراءةٍ مبسّطة - يسهم في ترسيخ هذا الوعي، ويُعين القارئ على الانطلاق بحُظا وثقة في عالمٍ يتعاظم فيه حضور الذكاء الاصطناعي يوميًا بعد يوم.



قائمة المراجع بحسب الفصول

ف	عنوان الفصل	المراجع المعتمدة
١	تعليمات النظام – كيف يُبنى سلوك النموذج من الجذر؟	• <i>Prompt Engineering Master Guide</i> – Om Prakash Saini • <i>Prompt Engineering for LLMs</i> – John Berryman
٢	التعليمات المخفية – لماذا لا أراك ولكن أشعر بك؟	• <i>The Art of Prompting</i> – Yibing Li • <i>Prompt Engineering</i> – Ajantha Devi Vairamani
٣	سياسة الأمان – من الذي يحرس النموذج من الخطأ؟	• <i>Practical Generative AI</i> – Kiko Llaneras • <i>Secrets of Machine Learning</i> – Tom Kohn
٤	تعليمات السياق – ما الذي يكتبه المستخدم فعلاً؟	• <i>ChatGPT Prompts Library</i> – GPT Penguin • <i>100 ChatGPT Power Prompts</i> – Mohammad Irfan
٥	النماذج المخصصة – كيف تُنشئ شخصية خاصة بك؟	• <i>Prompt Engineering Playbook (Beta v3)</i> • <i>Custom GPT Guide</i> (OpenAI وثائق)
٦	واجهة تكوين النموذج – ماذا يوجد خلف زر «تعديل»؟	• <i>Edit Custom GPT Settings</i> (OpenAI وثائق) • <i>600 ChatGPT Prompts</i> – Insiders AI
٧	التعليمات الخاصة – كيف تُخصّص النموذج من الداخل؟	• <i>Prompt Engineering</i> – Ajantha Devi Vairamani • <i>The Art of Asking ChatGPT</i> – Ibrahim John
٨	تضمين الملفات – كيف تعطي النموذج ذاكرة معرفية؟	• <i>Fundamentals of Prompt Writing</i> • <i>Custom Instructions Panel Documentation</i> (OpenAI)
٩	الأوامر الخارجية – كيف ينفّذ النموذج مهام فعلية؟	• <i>Prompt Engineering Tutorial</i> • <i>OpenAI API Documentation on Actions</i>
١٠	الذاكرة – هل يتذكّر النموذج من تكلم معه؟	• <i>Prompt Engineering for LLMs</i> – John Berryman • <i>OpenAI Memory FAQ</i>
١١	التعليمات المضمّنة – كيف يتكوّن السلوك من الداخل؟	• <i>3550 Most Effective Prompts</i> – Om Prakash Saini • <i>System-Level Behavior Guidelines</i> (OpenAI)
١٢	لوحة الموجهات – كيف تتحكم بتصرّفات النموذج في واجهة واحدة؟	• <i>GPT Debugging Tools (Playground)</i> • <i>Prompt Injection Defense Manuals</i>