



Ordine degli Psicologi

Consiglio del Friuli Venezia Giulia

Psicologia e Intelligenza Artificiale

STRUMENTI, RISCHI E RESPONSABILITÀ PROFESSIONALE

Sono Marco Parenzan



MARCO PARENZAN



marcoparenzan.bsky.social
marco_parenzan



marcoparenzan



marcoparenzan



marcoparenzan

ATTENZIONE

Sono un professionista informatico, non un professionista psicologo

I riferimenti a concetti e termini nell'ambito psicologico, generati tramite l'uso di una Generative AI, sono riportati a scopo esemplificativo.

Le origini dell'AI

Seconda guerra mondiale

Grande accelerazione tecnologica

Necessità militari come motore dell'innovazione

Crittografia e calcolo intensivo





Alan Turing

Decifrazione di Enigma

Macchine programmabili



Il Test di Turing

Basi teoriche dell'intelligenza artificiale

Macchina indistinguibile da un uomo nelle risposte

Valutazione comportamentale

Riferimenti culturali come Blade Runner

Le origini dell'IA (anni '40–'50)

Alan Turing e il Test di Turing (1950): nascita del concetto di macchine intelligenti.

Nel 1950, nel suo articolo *"Computing Machinery and Intelligence"*, Turing voleva rispondere a una domanda filosofica molto controversa:

“Le macchine possono pensare?”

Il problema era che la parola *"pensare"* è vaga e difficile da definire.

Così Turing fece qualcosa di geniale: **sostituì la domanda con un esperimento pratico.**

John von Neumann: architettura dei computer a calcolo sequenziale.

1956 – Conferenza di Dartmouth: John McCarthy conia il termine “Intelligenza Artificiale”.

Il documento afferma una tesi molto ambiziosa:

“Ogni aspetto dell'apprendimento o qualsiasi altra caratteristica dell'intelligenza può, in linea di principio, essere descritto così precisamente da poter essere simulato da una macchina.”

L'era simbolica (anni '50–'80)

L'era simbolica è la prima grande fase dell'AI.

L'idea di fondo era:

- l'intelligenza = manipolazione di simboli secondo regole logiche.

In pratica, si cercava di riprodurre il ragionamento umano usando:

- logica formale
- regole "if-then"
- rappresentazioni simboliche della conoscenza

Programmi pionieristici: ELIZA (Weizenbaum) e SHRDLU (Winograd).

Sistemi esperti basati su regole logiche: DENDRAL e MYCIN.

Limiti degli approcci simbolici → prima "AI Winter".

- Difficoltà nel gestire ambiguità e incertezza
- Scalabilità limitata (troppe regole da scrivere a mano)
- Nessun vero apprendimento automatico
- Fragilità fuori dai contesti controllati

First AI Winter

Modellare il pensiero

Regole «implicite», perchè le esplicite sono troppo complicate

- Non pensare a definire le regole
- Far apprendere dai dati

Le regole non sono più frasi logiche, ma:

- Pesi numerici
- Parametri di una funzione matematica

Si passa:

- Da logica booleana (vero/falso)
- A modelli probabilistici (più o meno probabile)

Algoritmi che:

- osservano esempi
- minimizzano un errore
- ottimizzano parametri

Connessioni tra neuroni artificiali

- Le regole diventano distribuite
- Non sono più interpretabili facilmente
- Emergono dal training

Limiti (1980-2000)

Molti sistemi simbolici e logici soffrivano di:

- Esplosione combinatoria
- Problemi NP-difficili
- Tempi di calcolo ingestibili all'aumentare delle variabili
- Più aumentava il dominio → più le regole crescevano → più il sistema diventava lento e fragile.

Altro problema enorme:

- Le regole dovevano essere scritte da esperti
- Estrarre la conoscenza dagli esperti era lento e costoso

Mantenere aggiornato il sistema era complicato

I sistemi diventavano:

- difficili da modificare
- difficili da estendere
- poco adattivi

Second AI Winter

L'era dell'AI

Internet

Accesso a enormi quantità di dati

Digitalizzazione dell'informazione

Condivisione globale

Consumerizzazione

Computer personali

Riduzione dei costi

Accesso diffuso al calcolo

Cloud Computing

Risorse condivise

Modello as-a-service

Economia di scala

Modello economico

Servizi a pagamento

Freemium

Sostenibilità economica

GPU e calcolo parallelo

Origine nei videogiochi

Calcoli vettoriali paralleli

Fondamentali per le reti neurali

NVIDIA

Evoluzione dalle schede grafiche

Accelerazione IA

Vantaggio competitivo

Infrastruttura e hype: data center e GPU

Grandi player costruiscono data center e comprano GPU

Costi di componenti (es. memorie) in crescita

Investimenti su speranze/promesse più che su numeri consolidati

Contesto economico ancora in assestamento

Big Data e Deep Learning (2010+)

Crescita dei dati digitali e uso delle GPU.

2012: AlexNet rivoluziona il riconoscimento delle immagini.

Diffusione dell'IA in ambito commerciale (Amazon, Netflix, Google Translate).

Artificial Intelligence

Broad AI Concepts & Techniques



Complexity
& Specialization

L'era della Generative AI

Problema nel riconoscimento della lingua

Elaborazione sequenziale (una parola alla volta)

Difficoltà nel catturare dipendenze lunghe

Addestramento lento

IA moderna ed era generativa (2010–oggi)

Paper 2017 «Attention is all you need»

Invece di leggere la frase in ordine rigido, il modello:

- Guarda tutte le parole insieme
- Assegna pesi di importanza tra di loro
- Costruisce rappresentazioni contestuali dinamiche

Il Transformer:

- È altamente parallelizzabile su GPU
- Scala bene con grandi dataset
- Permette modelli con miliardi di parametri

Non solo NLP:

- Visione artificiale (Vision Transformers)
- Audio
- Biologia computazionale
- Modelli multimodali

Perché ha:

- Semplificato l'architettura
- Migliorato le performance
- Permessi la scalabilità massiva
- Aperto la strada ai modelli generativi moderni

Tokenizzazione: dal testo ai token

Riceve qualcosa del tipo:

- ["Il", " gat", "to", " dorm", "e"]

Oppure token ancora più frammentati.

I token possono essere:

- Parole intere
- Parti di parola (subword)
- Caratteri
- Segni di punteggiatura

Ogni token viene convertito in:

- Un numero (ID)
- Poi in un vettore numerico (embedding)

Non “significati”, ma vettori in uno spazio matematico.

Il modello:

- Non ha concetti nel senso umano
- Non ha esperienza diretta del mondo
- Non “sa” cosa sia un gatto

Ma ha visto milioni di frasi in cui:

- “gatto” co-occorre con “miagola”
- “dorme” segue spesso un soggetto animato
- certe strutture sintattiche si ripetono
- Quindi impara distribuzioni di probabilità.

Generazione: predizione del prossimo token

Dato un input:

- "Il gatto dorme sul"

Il modello calcola:

- $P(\text{tappeto} \mid \text{contesto})$
 $P(\text{divano} \mid \text{contesto})$
 $P(\text{tetto} \mid \text{contesto})$

E sceglie (in base a probabilità e parametri di campionamento) il prossimo token.

Non recupera una risposta da un database.

Non cerca in una tabella.

- Calcola probabilità condizionate.

Questo è fondamentale:

- Non memorizza domande → risposte
- Non fa lookup diretto
- Non “ricorda” conversazioni passate (a meno che siano nel contesto corrente)

Un LLM è un gigantesco modello statistico che:

- Stima la distribuzione del linguaggio
- Approssima la probabilità delle sequenze
- Genera testo coerente con il contesto

Se è “solo statistica”, perché sembra ragionare?

Perché il linguaggio contiene la struttura del pensiero.

- E il modello impara quella struttura.

Quando gli esseri umani scrivono:

- spiegano passaggi logici
- mostrano relazioni causa-effetto
- risolvono problemi passo dopo passo

Il modello non “capisce” la logica come un filosofo.

Ma impara che:

- Se $A \rightarrow B$
A
 $\Rightarrow B$

è un pattern altamente probabile.

E quindi lo riproduce.

Non è magia.
È una proprietà dei sistemi ad alta dimensionalità.

Simulazione del ragionamento

Un LLM fa questo:

- Vede il problema
- Predice quale tipo di risposta è probabile
- Genera passaggi intermedi coerenti
- Arriva a una conclusione statisticamente consistente

Questo processo simula il ragionamento.

Ma attenzione:

Simulare $\neq$ Possedere comprensione interna.

Il punto cruciale

Durante il training, il modello costruisce uno spazio vettoriale dove:

- parole simili sono vicine

concetti correlati si organizzano geometricamente

relazioni diventano direzioni nello spazio

Esempio classico:

- $\text{Re} - \text{Uomo} + \text{Donna} \approx \text{Regina}$

Il modello:

- Non ha intenzioni
- Non ha coscienza
- Non ha un modello del mondo come il nostro

Ma:

- Ha una rappresentazione statistica estremamente raffinata del linguaggio
- Il linguaggio umano incorpora strutture logiche
- Riprodurre quelle strutture crea l'illusione di pensiero

Sembra ragionare perché il ragionamento umano lascia tracce statistiche nel linguaggio
e il modello è bravissimo a catturare e riprodurre quelle tracce.

Questa è vera “intelligenza” o solo una simulazione molto avanzata?

CAPACITÀ FUNZIONALE

- Risolvere problemi
- Usare il linguaggio
- Fare inferenze
- Generalizzare

Allora sì:
i modelli come GPT mostrano forme di
intelligenza funzionale

- Superano esami, scrivono codice, spiegano teorie.

COMPRENSIONE PROFONDA, COSCIENZA, ESPERIENZA SOGGETTIVA

- Consapevolezza
- Intenzionalità
- Esperienza interna
- “Sapere di sapere”

Un LLM:

- Non ha stati mentali
- Non ha desideri
- Non ha un mondo vissuto
- Non ha coscienza

Genera sequenze coerenti.

The «Imitation Game»

È un sistema statistico sofisticato che imita il linguaggio umano.

Nessuna comprensione reale.

Se si comporta come intelligente, allora è intelligente. (Una linea simile a quella di Turing.)

Un LLM:

- Non capisce nel senso umano
- Non ha grounding nel mondo fisico
- Non ha obiettivi propri

Ma:

- Costruisce rappresentazioni interne
- Generalizza oltre gli esempi visti
- Mostra comportamenti emergenti

Questo va oltre una semplice tabella di lookup.

Parametri

Dimensioni del modello

La “dimensione” di un modello è il numero di **parametri** (pesi).

Esempi indicativi:

- 1 miliardo (1B) → modello piccolo
- 7B → medio
- 70B → grande
- 100B+ → molto grande

Un parametro è un numero in virgola mobile (di solito 16 o 32 bit).

Più parametri =

- più capacità di rappresentazione
- più memoria
- più calcolo più costo

Cos'è un parametro

Un parametro è semplicemente:

- Un numero che il modello impara durante il training.
- Tecnicamente è un peso (weight) di una rete neurale.

Immagina una rete neurale come un enorme sistema di:

- nodi (neuroni artificiali)
- collegamenti tra nodi
- Ogni collegamento ha un numero associato.
- Quel numero è il parametro.

Un numero ottimizzato che determina come il modello trasforma input in output.

Cosa succede nel training

- Il modello fa una previsione
- Si calcola l'errore
- L'errore modifica leggermente i parametri
- Ripeti miliardi di volte

Alla fine i parametri:

- codificano pattern statistici
- rappresentano relazioni tra token
- contengono la “conoscenza” appresa

Negli LLM

- Matrici di pesi nelle proiezioni Q, K, V
- Matrici nei layer feed-forward
- Embedding dei token
- Bias

Non rappresentano frasi memorizzate.

Rappresentano:

- correlazioni statistiche
- relazioni geometriche nello spazio vettoriale
- trasformazioni matematiche

Un singolo parametro da solo non significa nulla.
Il significato emerge dall'insieme.

I parametri sono come:

- Le manopole di un gigantesco equalizzatore
- Le sinapsi di un cervello artificiale
- Le costanti di una funzione matematica enorme

Il training consiste nel regolare tutte queste manopole.

Memoria necessaria

La memoria dipende da:

- Parametri
- Gradiente
- Stati dell'ottimizzatore
- Batch size
- Attivazioni temporanee

Regola grossolana:

Per training in precisione mista (FP16/BF16):

- Servono circa 6–8× la memoria dei soli parametri

Esempio:

- Modello 70B parametri
 $70B \times 2 \text{ byte} \approx 140 \text{ GB}$ (solo pesi)
- Per training reale → anche 500–800 GB totali distribuiti su più GPU.

Tempi di training

Dipende da:

- Dimensione modello
- Dataset
- Numero GPU

Indicativamente:

- Modello piccolo → giorni
- 7B → settimane
- 70B+ → settimane/mesi su centinaia-migliaia GPU

Ordine di grandezza per modelli molto grandi:

- Milioni di dollari
- In alcuni casi decine di milioni
- Costo =
GPU time + energia + infrastruttura + ingegneri

Cost per Inference

Dopo il training:

Ogni richiesta ha un costo computazionale

Dipende da:

- Lunghezza input
- Lunghezza output
- Dimensione modello

Modelli grandi costano molto di più per token generato.

Cosa succede quando fai una domanda?

- Il testo viene tokenizzato
- I token diventano numeri
- Passano attraverso il Transformer
- Il modello predice il prossimo token
- Ripete il processo fino a fine risposta

Questo processo si chiama inferenza.

Modello grande (decine di miliardi parametri):

•20–200 millisecondi per token (dipende dall'infrastruttura)

•50 token/sec è già molto buono per modelli grandi

Da cosa dipende la performance?

Dimensione del modello (parametri)

- Più parametri → più calcoli per ogni token.
- Lunghezza del contesto

La self-attention ha costo circa:

$O(n^2)$

dove n = numero di token nel contesto.

Se raddoppi il contesto:

Lunghezza dell'output

- Il modello genera un token alla volta.
- Se chiedi una risposta lunga:
 - Deve fare forward pass per ogni token
 - Il tempo cresce linearmente con la lunghezza

Costi

Un numero

Per generare ~100 token con un 7B servono circa ~1.4 trilioni (1.4×10^{12}) di operazioni tipo moltiplicazione (ordine di grandezza).

ChatGPT

GPT è il modello matematico (Generative Pre-trained Transformer)

ChatGPT come servizio

ChatGPT ha creato sostanzialmente il mercato del servizio AI

Mercato IA: consumatori vs produttori

Chi consuma l'IA: professionisti e utenti finali

Chi crea i modelli: OpenAI, Anthropic, Google, ecc.

Chi integra: costruisce prodotti e strumenti sopra modelli esistenti

Molti 'nuovi esperti' sono integratori/prompt designer

Integrator e soluzioni verticali

Esempio: integrare GPT in e-commerce/gestionali/siti web

Prodotti verticali: stessi modelli, contesto specializzato

È spesso l'IA con cui l'utente lavorerà davvero

Importante capire chi sta “tra te e il modello”

Costi

L'AI non è gratuita

Il profilo gratuito è limitato nella dimensione delle risposte (non il numero delle risposte, ma sostanzialmente nel numero di Token – che non sono le parole)

Tutti i servizi di AI si stanno sostanzialmente allineando a livelli di:

- 0€/mese
- 20€ /mese
- 200€ /mese

Terzo AI Winter?

Un nuovo “raffreddamento” potrebbe avvenire se:

- L’hype supera troppo i risultati
- I costi energetici diventano insostenibili
- Si raggiunge un plateau tecnologico
- Regolamentazioni o limiti economici frenano l’espansione

La privacy nel Cloud

...E IL PROBLEMA DEL PROFESSIONISTA «NON AZIENDA»

GDPR nel Cloud

General Data Protection Regulation (GDPR)

Come leggere responsabilità e requisiti quando sposti dati su un cloud provider

Modello Cloud: chi fa cosa

IaaS (Infrastructure as a Service)

- Esempi: Amazon Web Services, Microsoft Azure, Google Cloud
- Provider → responsabile di infrastruttura, rete, data center
- Cliente → responsabile di dati, applicazioni, modelli AI, configurazioni
- Se un modello AI viola la privacy, la responsabilità è principalmente del cliente.

PaaS (Platform as a Service)

- Il provider gestisce piattaforma e strumenti AI, il cliente gestisce dati e utilizzo.
- Responsabilità condivisa.

SaaS / AI-as-a-Service

- Esempi: OpenAI, Salesforce
- Il provider controlla gran parte del sistema
- Il cliente decide come usarlo e quali dati inserire
- Anche qui la responsabilità è condivisa, ma dipende da ruoli contrattuali e GDPR.

Chi è il titolare dei dati?

Nella maggior parte dei casi:

- Tu / la tua azienda → siete il Titolare del trattamento
- Cloud provider → è il Responsabile del trattamento

Il titolare decide:

- perché vengono trattati i dati
- quali dati raccogliere
- per quanto tempo conservarli

Il provider:

- tratta i dati per tuo conto
- deve rispettare le tue istruzioni
- deve garantire sicurezza adeguata

Cosa può (e non può) fare il provider?

Il provider non può:

- usare i tuoi dati per finalità proprie (salvo accordi specifici)
- conservarli oltre quanto previsto
- trasferirli fuori UE senza garanzie adeguate

Il provider deve:

- adottare misure di sicurezza tecniche e organizzative
- notificare eventuali data breach
- firmare un DPA (Data Processing Agreement)

Esempi di cloud provider comuni:

- Amazon Web Services
- Microsoft Azure
- Google Cloud

Attenzione alle righe piccole quando il servizio è gratuito

Se uso servizi AI in cloud?

Qui la situazione può complicarsi.

Se usi un servizio AI (es. API o SaaS):

- Se inserisci dati personali → resti titolare
- Il provider AI → normalmente responsabile

Se il provider usa i dati per addestrare i propri modelli → può diventare titolare autonomo

Esempio:

- OpenAI
- Microsoft
- Dipende dalle condizioni contrattuali.

Di chi è la responsabilità se succede qualcosa?

Data breach infrastrutturale

- Responsabilità primaria → provider

Uso illecito o eccessivo dei dati

- Responsabilità → tua (come titolare)

Decisioni automatizzate illegittime

- Responsabilità → tua (art. 22 GDPR)

Trasferimento extra-UE irregolare

- Responsabilità condivisa, ma il titolare resta il primo responsabile

In sostanza

Situazione	Responsabilità principale
Violazione dati per cattiva configurazione	Cliente
Data breach infrastrutturale	Provider
Uso scorretto dell'AI sui dati personali	Cliente (titolare)
Difetto tecnico dell'AI	Può coinvolgere il provider

Trasferimenti extra-UE: cosa guardare (Capo V)

Se dati (o accessi) escono dallo SEE → servono garanzie e documentazione

Attenzione ai servizi “global”

- IAM/telemetry/support possono comportare trasferimenti o accessi extra-UE.

Misure supplementari tipiche

- Cifratura forte + controllo chiavi; minimizzazione; pseudonimizzazione; egress control.

Data residency vs sovereignty (3 livelli)

1) Data Residency

- Dati fisicamente in un Paese/region (EU-only).

2) Sovereignty giuridica

- Limitare accessi extra-UE; governance su supporto e personale; controllo dei trasferimenti.

3) Sovereign Cloud

- Control plane, operazioni e supporto in UE; opzioni per chiavi e accessi “EU-only”.

Rischi tipici nell'AI generativa

Prompt leakage

- Prompt/output possono contenere dati personali o segreti; spesso finiscono in log/telemetria.

Retention poco chiara

- Conservazione di prompt/log oltre il necessario o fuori dal perimetro UE.

Training/fine-tuning

- Riutilizzo dei dati per migliorare modelli (se consentito) = rischio elevato.

Re-identificazione

- Dati pseudonimizzati possono essere ricostruiti se contesto/metadata restano.

AI layer: 3 opzioni (quando usarle)

A) AI SaaS Enterprise

- Per casi a rischio medio e time-to-market; richiede settaggi no-retention e controlli contrattuali.

B) AI gestita ma isolata (consigliata)

- Endpoint dedicato in UE, rete privata, telemetria ridotta; buon equilibrio.

C) Self-hosted LLM

- Per dati critici o vincoli forti; massima sovranità, maggiore complessità/costi.

Cos'è il RAG (Retrieval-Augmented Generation)

RAG significa Retrieval-Augmented Generation.

È un'architettura che combina:

- Recupero di informazioni (retrieval) da fonti esterne
- Generazione di testo tramite un modello linguistico (LLM)

Indicizzazione

- I tuoi documenti (PDF, email, database, policy...) vengono suddivisi in parti.
- Ogni parte viene trasformata in embedding (vettori numerici).
- I vettori vengono salvati in un vector database.

In pratica:

- invece di rispondere solo con ciò che “sa”, il modello cerca prima nei tuoi documenti e poi genera la risposta.

Miti e leggende della GenAI

OVVERO NON ASPETTIAMOCI DALLA GENAI QUELLO CHE NON
POSSIAMO ASPETTARCI NEMMENO DA NOI STESSI

Entriamo nel merito: correttezza della risposta

Tema pratico: l'IA risponde "bene"?

Se la risposta sembra corretta, crescono fiducia e aspettative

Nasce la domanda: quali rischi corriamo?

Riformulazione: quali aspettative abbiamo verso l'IA?

Caso personale: IA come compagno quotidiano

Account a pagamento, uso giornaliero

Uso tipico: scrive al posto nostro dove serve scrivere (software, relazioni, email)

Domanda: la risposta è giusta o sbagliata?

Bisogno reale: saper verificare la correttezza

Verifica: chi è responsabile?

L'IA propone; il professionista controlla, valuta, decide

Il valore non è che la risposta sia “mia”

Il valore è avere uno strumento a disposizione

La consegna finale resta responsabilità umana

IA come acceleratore (non sostituto)

Idea centrale: l'IA è prima di tutto un acceleratore

Due visioni: IA fa il lavoro al posto mio vs IA coadiuva il mio lavoro

Tema dell'incontro: produttività e supporto

Tema “mi ruberà il lavoro” è collegato ma non unico

Precisione: l'IA (e noi) sbagliamo

Provocazione: voi rispondete sempre in modo preciso?

Tutti commettiamo errori, anche da professionisti

Errore spesso legato a informazioni incomplete

Serve valutare contesto e qualità dei dati in ingresso

Perché una risposta umana può essere sbagliata

Il paziente/cliente non fornisce tutte le informazioni

Non sempre si fanno tutte le domande giuste

Possibili bugie/omissioni/interpretazioni

Possibili limiti di conoscenza sul caso specifico

Bias, emozioni e non-oggettività

Possibili blocchi personali o influenze emotive

Rischio di risposta “influenzata” dal legame con la problematica

Non è dolo: è umano

Tema: oggettività e qualità decisionale

IA e correttezza: la domanda reale

Hai un problema se l'IA non è perfetta?

Può dare risposte corrette spesso, ma non sempre

Può capitare una risposta non soddisfacente

Quindi: aspettative realistiche e controllo continuo

Cliente/paziente con risposta già pronta

Caso già reale: “ho chiesto a ChatGPT e mi ha detto che...”

A volte arriva per conferma, non per diagnosi iniziale

Possibile frizione economica: conferma senza pagare

Dinamica tipica di mercato: “voglio spendere poco”

Il rischio è trasversale a tutti i mestieri

Non riguarda solo l'ambito clinico/psicologico

Nel software: un junior può essere surclassato dall'IA

L'IA offre velocità e accesso a tanta informazione

Il senior deve riposizionare valore e competenze

Dove l'IA aiuta davvero: pattern e memoria

Molti mestieri = riconoscere pattern tra segnali

Serve metodo per organizzare informazioni e condizioni

L'uomo ha limiti di memoria e organizzazione

L'IA eccelle nel collegamento statistico di informazioni

Metafora: “ha letto tutti i libri del mondo”

Immagine: accesso vastissimo a testi e casi documentati

Velocità di recupero e sintesi enorme rispetto all'umano

Potenziale: risposta più ampia/precisa

Ma dipende dal contesto e da come chiedi

Adozione nel business: non è ovunque

Non tutte le aziende l'hanno integrata nei processi

Molte sperimentano, poche hanno integrazione completa

Vendite sì, ma non sempre sufficienti a giustificare tutto l'hype

Possibile rischio bolla: promesse > prove misurabili

Punto fermo: acceleratore + responsabilità

L'IA “prepara” lavoro, proposte, alternative

Il potenziale diventa tuo e lo porti al cliente/paziente

Al cliente interessa la tua risposta, non la tecnologia dietro

Professionalità = interpretare, filtrare, adattare, rispondere

Copia-incolla vs appropriazione professionale

Opzione A: copiare e incollare la risposta

Opzione B: rielaborarla e presentarla in modo tuo

Opzione C: contestarla e correggerla

La scelta e la responsabilità sono tue

Dialogo iterativo: la metafora della chat

Generative AI = conversazione per affinare domanda e risposta

Aggiungi informazioni, chiedi chiarimenti, dai feedback

Il sistema produce nuove varianti e precisazioni

Obiettivo: risposta finale validata da te

Analogia: requisiti e comunicazione nel software

Il cliente non dà mai requisiti completi

Il senior deve fare domande giuste e chiarire

Il junior esegue ma può fraintendere

Spesso l'errore nasce da requisiti imprecisi

Tecnologia e 'imprecisione' (messa a nudo)

Le tecnologie moderne evidenziano la nostra imprecisione

Siamo stanchi, di corsa, sotto scadenza

Scriviamo requisiti incompleti e poi critichiamo l'output

Con l'IA questo diventa evidente e ripetibile

Scenari quando un compito è poco chiaro

Fa la cosa ma non l'ha capita

La capisce e chiede chiarimenti

Non sa farla e chiede aiuto

Peggio: copia da Internet, non verifica, va in produzione → disastro

Facciamo una domanda
alla GenAI

Il punto chiave: prompt (domanda)

Prompt = ciò che scrivi per chiedere qualcosa

Prompt corto → risposte meno prevedibili e meno mirate

Prompt dettagliato → risposte più aderenti allo scenario

Prompt engineering = saper scrivere prompt efficaci

Prompt lunghi e limiti: token e contesto

Un prompt può essere lungo: decine o centinaia di righe

Limiti tecnici: finestre di contesto (numero massimo di token)

Più descrivi lo scenario, più riduci ambiguità

Serve trovare equilibrio tra dettaglio e chiarezza

Il prompt include anche la cronologia

In chat, il contesto include conversazioni precedenti

Il sistema cerca di includere cronologia ad ogni richiesta

Il contesto non è infinito (limiti/costi)

Le risposte future dipendono da storia + istruzioni + domanda

System prompt: regole e stile

Spesso esiste un prompt iniziale di sistema

Definisce regole: lingua, formato, vincoli, ambito

Esempio: “se non capisci, dimmelo e non improvvisare”

Serve a ridurre risposte fuori tema

Ambiguità e typo: esempio delle arance

‘Tarocco’ interpretato come ‘tarocchi’ per errore di scrittura

Con poche parole, un typo può deviare la risposta

Aggiungere contesto corregge rapidamente

Le GAI tendono a rispondere comunque, anche se non sicure

Temperatura e risposta 'a tutti i costi'

Temperatura = grado di creatività/improvvisazione (dipende dal modello)

Anche una persona, in punti della conversazione «tira a indovinare»

Modelli generalisti tendono a rispondere sempre

Questo aumenta il rischio di risposta sbagliata ma plausibile

Da qui l'importanza di vincoli e revisione

Revisione obbligatoria dell'output

La risposta va sempre rivista

Se non è corretta, spesso la domanda era incompleta/ambigua

Si itera con chiarimenti per migliorare

Validazione finale resta umana e contestuale

Allucinazioni: ambiguità + contesto 'sporco'

Allucinazione = risposta convincente ma scorretta

Spesso nasce da ambiguità del prompt o del contesto

Il contesto può "sporcarsi" e generare loop

Strategia: fermarsi, ripulire, riformulare

Chiudere il loop: fermarsi e ripartire

Il modello può scusarsi e correggersi (stile colloquiale)

Ma non garantisce verità: produce testo plausibile

Prendere il buono di più risposte e costruire la soluzione finale

L'IA è strumento: il professionista decide e risponde

Linee guida per
impostare
un prompt efficace

Perché il prompt conta

Il prompt è la “consegna” che guida la risposta

Più è chiaro e completo, meno l’AI improvvisa

Riduce ambiguità, errori e risposte generiche

Aumenta utilità, coerenza e ripetibilità dei risultati

Messaggio chiave: un buon prompt non è una domanda—è una consegna ben strutturata.

Definisci il ruolo dell'AI

Specifica “chi” deve essere l'AI (ruolo/competenza)

Aiuta a scegliere tono, approccio e priorità

Esempi: “supervisore clinico”, “docente”, “redattore”

Suggerimento: aggiungi anche l'orientamento (es. CBT, sistemico, ACT) se rilevante.

Chiarisci l'obiettivo

Cosa deve fare esattamente? (riassumere, strutturare, generare esempi...)

Una richiesta = un obiettivo principale

Se hai più obiettivi, ordinali per priorità

Esempio obiettivo: “Genera 5 esercizi gradualmente di esposizione” (non: “Dammi delle idee”).

Fornisci contesto sufficiente

Informazioni rilevanti: target, scenario, vincoli, livello di dettaglio

Usa dati anonimi o scenari sintetici quando possibile

Indica cosa è già stato provato e cosa non funziona

Regola pratica: se un collega non capirebbe la richiesta, nemmeno l'AI potrà farlo bene.

Specifica il formato dell'output

Elenco puntato, tabella, schema in N punti, checklist, dialogo simulato...

Chiedi lunghezza: “max 200 parole”, “5 punti”, “3 esempi”

Richiedi stile: “linguaggio semplice”, “tono neutro”, “senza tecnicismi”

Il formato è un moltiplicatore di qualità: guida struttura e leggibilità.

Imposta il livello di profondità

Sintetico: per overview rapide (es. 5 punti)

Approfondito: per ragionamento e alternative

Operativo: per passaggi concreti e sequenze

Didattico: per spiegazioni e esempi per pazienti

Chiedi esplicitamente: “prima la sintesi, poi un approfondimento opzionale”.

Aggiungi vincoli e limiti (sicurezza)

Cosa evitare: diagnosi, giudizi, linguaggio tecnico, dati identificativi

Cosa privilegiare: prudenza, alternative, domande di chiarimento

Chiedi di segnalare incertezze e assunzioni

Per ambito clinico: usa l'AI come supporto—non come autorità. Verifica sempre.

Formula pratica (da ricordare)

RUOLO
+ OBIETTIVO
+ CONTESTO
+ FORMATO
+ VINCOLI

Esempio di prompt completo

Ruolo: “Agisci come psicoterapeuta CBT.”

Obiettivo: “Suggerisci 5 esercizi di esposizione graduale.”

Contesto: “Ansia sociale, ruminazione post-evento, evitamento lavorativo.”

Formato: “Elenco puntato, linguaggio semplice per il paziente.”

Vincoli: “Evita diagnosi, non usare dati identificativi, segnala assunzioni.”

Suggerimento: chiedi anche “prima 5 punti, poi 2 alternative se il paziente rifiuta l’esposizione”.

Una domanda a caso

ATTENZIONE!

MI SONO FATTO AIUTARE DA UNA GenAI!

...rullo di tamburi...

Conclusioni

L'AI è uno strumento, non un decisore

Non delegare diagnosi o decisioni cliniche

Non delegare scelte terapeutiche (obiettivi, tecniche, timing)

Non delegare responsabilità professionale

Usare l'AI come supporto cognitivo (bozze, schemi, alternative)

Regola: l'AI può suggerire, ma il clinico decide e si assume la responsabilità.

Protezione dei dati e privacy

Non inserire dati identificativi (nome, contatti, dettagli riconoscibili)

Minimizzare le informazioni cliniche (solo ciò che serve)

Verificare policy del servizio (conservazione, uso dati, log)

Preferire soluzioni professionali/istituzionali quando necessario

Attenzione: consenso informato, confidenzialità e conformità normativa (es. GDPR).

Verifica e pensiero critico

Controllare sempre fonti e contenuti (soprattutto se “scientifici”)

Diffidare di risposte troppo sicure o “definitive”

Chiedere alternative, limiti e incertezze

Non usare output senza revisione e contestualizzazione clinica

L'AI può sbagliare in modo convincente: la verifica è parte del lavoro.

Competenza digitale (minima)

Comprendere limiti tecnici (allucinazioni, bias, contesto)

Aggiornarsi su rischi, linee guida e normative

Distinguere strumenti: ricerca, sintesi, analisi documenti, chatbot

Definire procedure: quando usarla, quando no, e come documentare

Uso responsabile = competenza + procedure + supervisione clinica.

Motori/strumenti AI utili per psicologia: confronto ampliato

Strumento	Migliore per	Punti di forza	Limite/Rischio tipico
Consensus	Domande cliniche basate su evidenza	Sinthesi da letteratura; orientato a “cosa dice la ricerca”	Copertura non totale; dipende da indicizzazione
Perplexity	Overview rapida con fonti	Risposte veloci con citazioni; utile per esplorazione iniziale	Qualità varia: fonti online non sempre accademiche
Deep Research	Analisi e confronto multi-fonte	Report strutturati; comparazioni; utile per relazioni/lezioni	Richiede verifica critica; rischio di includere fonti deboli
Notebook LLM (Google)	Analisi di documenti caricati	Lavora sui tuoi file; ottimo per linee guida e PDF lunghi	Non cerca “fuori”; dipende dalla qualità dei file
Elicit	Revisione rapida di studi	Estrazione info metodologiche; supporto a review	Può perdere dettagli; serve controllo del paper
Scite	Valutare solidità/contesto delle citazioni	Distingue citazioni che supportano vs contraddicono	Non sostituisce lettura critica; copertura variabile
Semantic Scholar	Ricerca accademica ampia	Buoni filtri e discovery; rete di articoli/autori	Meno “sintesi pronta”; richiede selezione manuale
ChatGPT / Claude (con PDF)	Sintesi, didattica, trasformazione contenuti	Ottimi per schemi, materiale psicoeducativo, FAQ	Rischio allucinazioni se non vincolato ai documenti; privacy
PsyArXiv / preprint	Aggiornamento molto recente	Accesso a lavori prima della pubblicazione	Non peer-reviewed: interpretare con cautela
Google Scholar	Copertura ampia e storica	Ottimo “indice” generale; citazioni e alert	Non filtra qualità automaticamente; serve valutazione critica

Quale usare in base all'obiettivo clinico

Obiettivo	Strumento consigliato	Perché (in 1 riga)
Trovare “cosa dice la ricerca” su un tema	Consensus + Semantic Scholar	Consensus sintetizza; Semantic Scholar amplia e filtra
Fare un primo orientamento con fonti	Perplexity	Ottimo per overview rapida (poi si verifica)
Preparare una lezione / relazione approfondita	Deep Research + Scite	Struttura e confronto; Scite aiuta a pesare le citazioni
Preparare una review preliminare	Elicit + Semantic Scholar	Estrae info; aiuta a costruire una base di paper
Leggere linee guida/PDF lunghi e fare schemi	Notebook LLM / ChatGPT / Claude (con PDF)	Sintesi e confronto su documenti caricati
Valutare robustezza di uno studio molto citato	Scite	Capire se viene supportato o contraddetto
Aggiornarsi su lavori appena usciti	PsyArXiv + Google Scholar	Preprint + ampia copertura e alert
Trasformare evidenze in materiale per pazienti	ChatGPT / Claude (vincolati ai documenti)	Didattica, linguaggio semplice, schede e FAQ

Un rischio realistico per
la professione?

ATTENZIONE
OPINIONE PERSONALE!

Un rischio realistico?

Ci saranno professionisti (world-wide) che si trasformeranno in esperti di AI (o parteciperanno allo sviluppo di AI o di servizi AI)

- Sviluppo di modelli dedicati

Vi porteranno via il lavoro?

- Es. UnoBravo...
- Non è l'AI che vi porta via il lavoro
- Sono i pazienti che si fideranno dell'AI
 - Per il motivo sbagliato: per costi, prima di tutto

Grazie...domande?



marcoparenzan.bsky.social
marco_parenzan



marcoparenzan



marcoparenzan



marcoparenzan

Rispondiamo ad alcune
domande

Limiti e rischi specifici in ambito clinico

Limiti e rischi specifici in ambito clinico

Allucinazioni: output plausibili ma falsi → rischio clinico se non verificati

Riduzionismo: testo ≠ osservazione (non verbale, tono, contesto relazionale)

Bias: culturali/di genere/diagnostici → rischio di etichettamento

Overconfidence: stile assertivo che aumenta l'effetto autorità

Gestione del rischio (crisi, suicidarietà) non affidabile in autonomia

Privacy: inserimento involontario di dati identificativi o sensibili

Regola pratica: usare l'AI per supporto (bozze/schemi), mai come decisore clinico.

Etica e deontologia: come chiudere l'intervento

Centralità della persona: evitare automatismi e riduzionismi

Competenza: usare strumenti solo se se ne comprendono limiti e rischi

Responsabilità: il professionista risponde delle decisioni, non l'AI

Confidenzialità: minimizzazione dati + anonimizzazione + valutazione del fornitore

Non maleficenza: evitare usi che aumentano stigma o rischio clinico

Messaggio finale: "l'AI può amplificare il lavoro, non sostituire l'atto clinico."

Implicazioni medico-legali e regolatorie (alto livello)

Protezione dati: GDPR e regole su dati sanitari → minimizzazione e basi giuridiche

Consenso e informativa: valutare trasparenza sull'uso di strumenti AI

Sicurezza: dove finiscono i dati, tempi di conservazione, accessi, audit

Tracciabilità: distinguere cosa è stato “generato” vs cosa è decisione clinica

Rischio contenzioso: affidarsi a output non verificati può aumentare la responsabilità

Buona prassi: policy interna + procedure + formazione + documentazione d'uso.

Esempio pratico: usare l'AI su Big Five / PID-5 (solo supporto)

Input: fornire punteggi già calcolati e descrittori sintetici (anonimi)

Compito: chiedere una sintesi neutra + punti di forza/vulnerabilità

Richiesta: 2–3 ipotesi alternative + domande di verifica (no diagnosi)

Output: bozza di restituzione in linguaggio comprensibile (non stigmatizzante)

Suggerimento prompt: “segnala limiti e assunzioni; non usare etichette diagnostiche; proponi domande cliniche.”

Chatbot vs Psicoterapia tradizionale (comparativa)

Dimensione	Chatbot	Psicoterapia tradizionale
Alleanza terapeutica	Limitata/mediata	Centrale e relazionale
Personalizzazione	Buona su protocolli, limitata su complessità	Alta, basata su contesto vivo
Gestione rischio	Non affidabile in autonomia	Valutazione clinica diretta
Evidenza	Buona su sintomi lievi-moderati (digital CBT)	Ampia su molti quadri e severità
Accessibilità/costi	Alta scalabilità	Limitata da tempo/risorse
Responsabilità	Non professionale	Professionale e deontologica

Chatbot vs Colloquio clinico umano (comparativa)

Aspetto	Chatbot	Colloquio umano
Canali informativi	Testo/voce (dipende dallo strumento)	Multimodale (verbale + non verbale)
Accuratezza	Buona su screening strutturati	Maggiore su valutazione complessa
Sfumature relazionali	Limitate	Ricche (alleanza, sintonizzazione)
Bias	Possibili bias del modello	Possibili bias del clinico (mitigabili)
Tracciabilità	Log e trascrizioni utili	Dipende dalla documentazione clinica

Comparazioni più tecniche (senza numeri fuorvianti)

Screening: spesso comparabile a self-report quando è strutturato e guidato

Diagnosi differenziale: significativamente limitata senza contesto multimodale

Robustezza: varia tra strumenti e dipende da prompt, dati e verifiche

Validità: dipende da protocolli, governance dei dati e supervisione clinica

Se vuoi percentuali: meglio citarle da studi specifici e contestualizzarle (popolazione, compito, setting).

Quale strumento usare in base all'obiettivo (sintesi)

Obiettivo	Strumento	Nota
Domanda evidence-based	Consensus	Sintesi da letteratura; verificare copertura
Cercare paper e filtri	Semantic Scholar / Google Scholar	Ottimi per discovery e citazioni
Pesare citazioni	Scite	Supporta vs contraddice
Review preliminare	Elicit	Estrae info metodologiche
Analisi di PDF/linee guida	Notebook LLM / ChatGPT / Claude	Vincolare ai documenti; privacy
Overview veloce con fonti	Perplexity	Ottimo per orientarsi, poi verificare

Quanto può essere
accurato
un colloquio clinico guidato
da un chatbot?

EVIDENZE, POTENZIALITÀ E LIMITI STRUTTURALI

Dove può essere accurato

Raccolta strutturata di sintomi (checklist)

Screening per ansia/depressione lieve–moderata

Somministrazione guidata di questionari standardizzati

Monitoraggio sintomatologico ripetuto (follow-up)

Identificazione di parole chiave e pattern testuali

Buoni per screening e raccolta dati; meno per diagnosi differenziale complessa.

Dove l'accuratezza diminuisce

Assenza di osservazione non verbale (mimica, postura, comportamento)

Mancanza di prosodia e tono emotivo reale (nel testo)

Difficoltà nel cogliere ambivalenze e sfumature relazionali

Valutazione del rischio (es. suicidario) spesso non affidabile

Ridotta integrazione del contesto socio-relazionale

Il colloquio clinico è multimodale: non è solo linguaggio.

Cosa dice la ricerca (in sintesi)

Accuratezza spesso simile a strumenti di screening self-report

Buona coerenza in interviste strutturate e protocolli guidati

Migliori risultati in interventi digitali CBT strutturati

Variabilità tra sistemi e contesti d'uso

Evidenze in evoluzione: serve cautela nel generalizzare

Sintesi: utili in contesti specifici e ben definiti, non equivalenti alla valutazione specialistica.

Rischio di falsa sicurezza

Linguaggio fluente $\neq$ competenza clinica

Risposte plausibili $\neq$ corrette

Coerenza narrativa $\neq$ validità diagnostica

Tono “sicuro” può aumentare l’effetto autorità

Rischio clinico: sopravvalutare l’affidabilità di output ben scritti.

Conclusione clinica (cosa può / cosa non può)

- ✓ Supportare screening e raccolta dati
- ✓ Strutturare informazioni e follow-up
- ✓ Monitorare sintomi nel tempo
- ✗ Sostituire una valutazione clinica completa
- ✗ Formulare diagnosi complesse affidabili
- ✗ Assumersi responsabilità terapeutica

Buona regola: usarlo come “assistente” e non come decisore.

Conclusioni

L'accuratezza tecnica non equivale a competenza clinica.

Chatbot che supportano un intervento di psicoterapia

ESEMPI, FUNZIONI PRATICHE E LIMITI CLINICI

Esistono? Sì, ma come supporto

Esistono chatbot progettati per supporto psicologico e self-help guidato

Sono pensati per integrare (non sostituire) il lavoro clinico

Maggior evidenza in protocolli strutturati (es. CBT digitale)

Messaggio chiave: supporto operativo e psicoeducativo, non “terapia” autonoma.

Chatbot di auto-aiuto strutturato

Woebot → micro-interventi CBT e psicoeducazione

Wysa → esercizi CBT, mindfulness e journaling

Youper → monitoraggio emotivo e tecniche di ristrutturazione

Nota: disponibilità, funzioni e policy possono variare nel tempo e per paese.

Modello “blended therapy” (con terapeuta)

Supporto tra le sedute (check-in, reminder, micro-esercizi)

Homework digitali guidati e follow-up

Monitoraggio dell’umore e dei sintomi

Raccolta dati per discussione in seduta (con consenso)

L’AI qui funziona da “assistente operativo”, non da decisore clinico.

Cosa fanno concretamente (funzioni tipiche)

Psicoeducazione breve e contestuale

Diario emotivo e monitoraggio (mood tracking)

Esercizi di ristrutturazione cognitiva / problem solving

Tecniche di grounding e regolazione emotiva

Micro-interventi ACT / mindfulness

In genere rendono meglio quando inseriti in percorsi strutturati.

Cosa NON fanno (limiti importanti)

Non gestiscono bene crisi complesse o rischio suicidario

Non sostituiscono alleanza terapeutica e valutazione clinica

Non fanno diagnosi affidabili

Rischio di risposte generiche o non contestualizzate

Questioni di privacy, sicurezza e responsabilità professionale

Regola pratica: per casi complessi o rischio → invio a professionista e protocolli di sicurezza.

Evidenze scientifiche (in sintesi)

Risultati migliori su sintomi lievi-moderati e interventi brevi

Buona accettabilità in alcuni target (es. giovani adulti)

Utili per accesso e continuità, ma con rischio drop-out

Messaggio: strumenti utili, ma non equivalenti alla psicoterapia relazionale in presenza.

Conclusioni

I chatbot possono ampliare l'accesso al supporto psicologico.

Non sostituiscono la psicoterapia relazionale.

L'AI generativa può
formulare ipotesi
diagnostiche da una
trascrizione clinica?

OPPORTUNITÀ RIFLESSIVE E LIMITI CLINICO-DEONTOLOGICI

Cosa può fare (uso prudente)

Evidenziare temi ricorrenti e pattern narrativi

Segnalare contenuti emotivi dominanti nel linguaggio

Organizzare sintomi riportati e loro ricorrenze

Proporre domande di approfondimento clinico

Generare 2–3 ipotesi diagnostiche alternative (in forma esplorativa)

Principio: supporto alla riflessione, non conclusione diagnostica.

Cosa NON può fare

Effettuare una diagnosi clinica valida

Valutare non verbale, prosodia, contesto relazionale reale

Integrare anamnesi completa e dati extra-testuali se non forniti

Valutare validità della risposta/simulazione in modo affidabile

Assumersi responsabilità professionale o medico-legale

La diagnosi è un processo clinico complesso, non un output testuale.

Rischi principali

Sovra-diagnosi o “etichette” premature

Riduzionismo testuale (solo parole, non relazione)

Linguaggio eccessivamente assertivo → effetto autorità

Bias culturali, di genere o diagnostici

Delega implicita del giudizio clinico

Rischio maggiore: scambiare una “ipotesi plausibile” per una valutazione valida.

Uso corretto (modalità riflessiva)

Trascrizione completamente anonimizzata e minimizzata

Richiedere tono prudente e dichiarazioni di incertezza

Chiedere ipotesi alternative + criteri/indicatori a favore/contro

Chiedere quali informazioni mancano per confermare/negare

Revisione clinica obbligatoria (mai automatizzare decisioni)

Obiettivo: ampliare lo sguardo, non chiuderlo.

Esempio di prompt (più sicuro)

Ruolo: “Agisci come supporto alla riflessione clinica (non diagnosi).”

Input: “Trascrizione anonimizzata + contesto essenziale (senza dati identificativi).”

Compito: “Evidenzia pattern e organizza i sintomi riportati.”

Output: “2–3 ipotesi diagnostiche possibili, tutte preliminari e alternative.”

Vincoli: “Indica incertezze e quali informazioni mancano; evita conclusioni definitive.”

Suggerimento: chiedi anche “usa un linguaggio non stigmatizzante e culturalmente sensibile”.

Conclusioni

L'AI può suggerire possibilità.

La diagnosi resta un atto clinico relazionale, contestuale e responsabile.

Come può essere già
oggi usata
l'AI in psicologia clinica?

Supporto alla scrittura clinica

Riformulare appunti in forma più strutturata

Sintetizzare colloqui a partire da note anonimizzate

Creare report psicoeducativi o riassunti per il paziente

Migliorare chiarezza, coerenza e leggibilità del testo

Nota: evitare dati identificativi; verificare sempre contenuti e formulazioni.

Preparazione di materiale terapeutico

Schede e homework personalizzabili (esercizi, monitoraggi, diari)

Tracce per esercizi: esposizione, defusione, ristrutturazione cognitiva...

Script per psicoeducazione (es. ansia, panico, ruminazione)

Metafore e esempi adattati al linguaggio del paziente

L'AI accelera la produzione: **la supervisione e la responsabilità restano del clinico.**

Simulazione e role-play

Simulare risposte di un “paziente” con un profilo specifico

Allenare riformulazioni empatiche e domande esplorative

Esercitarsi su colloqui difficili (resistenze, ambivalenza, drop-out)

Preparare interventi e psicoeducazione prima della seduta

Utile per formazione e supervisione; attenzione a non “standardizzare” eccessivamente.

Supporto alla formulazione del caso

Organizzare informazioni cliniche in modo ordinato (timeline, temi, pattern)

Generare ipotesi alternative e domande di approfondimento

Evidenziare possibili fattori di mantenimento

Strutturare mappe concettuali o schemi del caso

Non sostituisce il giudizio clinico: supporta il ragionamento e la chiarezza.

Aggiornamento e studio

Sintesi di articoli e linee guida (con verifica delle fonti)

Confronto tra modelli teorici e concetti chiave

Creazione di schemi riassuntivi e flashcard

Preparazione di lezioni, supervisioni, workshop

Buono per risparmiare tempo; serve sempre controllo critico e citazioni affidabili.

Supporto organizzativo (non clinico)

Bozze di email e comunicazioni standard (non cliniche)

Template di documenti informativi e materiali per il sito

Organizzazione di agenda, checklist e priorità

FAQ e testi divulgativi per pazienti o famiglie

Separare contenuti amministrativi da contenuti clinici; attenzione a privacy e consenso.

Conclusioni

Oggi l'AI non sostituisce il terapeuta.

Può però amplificare tempo, chiarezza e capacità organizzativa.

Come usare l'AI per
leggere
un profilo di personalità

Chiave di lettura delle slides seguenti

l'AI non “legge” la personalità

Non fa diagnosi né sostituisce la valutazione clinica

Non somministra test validati

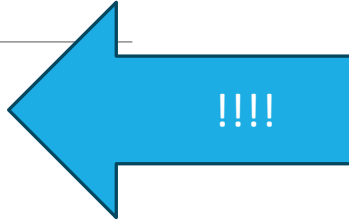
Lavora solo su dati forniti (testi, punteggi, osservazioni)

Può aiutare a organizzare e sintetizzare, non a “concludere”

Attenzione! Potrebbe scrivere delle cose che possono essere interpretate come una conclusione, perchè «statisticamente» alla fine di un «ragionamento» ci stanno le «conclusioni».

Analisi di testi clinici

ATTENZIONE! Usare solo materiale anonimizzato e minimizzato (niente dati identificativi).



Evidenziare temi ricorrenti e pattern narrativi

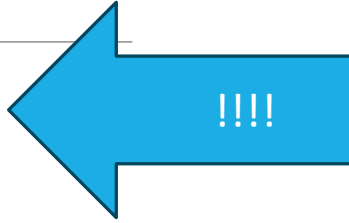
Rilevare polarizzazioni cognitive e modalità espressive

Analizzare coerenza interna del racconto (senza “verità” clinica)

Suggerire domande per approfondire (intervista clinica)

Organizzazione di risultati test (già raccolti)

ATTENZIONE! Usare solo materiale anonimizzato e minimizzato (niente dati identificativi).



Riassumere profili da strumenti standardizzati (es. Big Five, PID-5, MMPI, TCI...)

Evidenziare punti di forza, vulnerabilità e domini principali

Tradurre il profilo in linguaggio comprensibile (per restituzione)

Proporre una struttura di report (non contenuti “inventati”)

Generare ipotesi (non conclusioni)

Produrre 2–3 ipotesi esplicative alternative

Indicare domande di verifica per ciascuna ipotesi

Separare “dati osservati” da “interpretazioni”

Evidenziare assunzioni e livello di incertezza

Rischi principali da gestire

Overinterpretazione (andare oltre i dati disponibili)

- Può associare dati da altre fonti!

Bias (nel modello o nei dati di input)

Tono troppo “sicuro” che aumenta l’effetto autorità

Generalizzazioni eccessive (perdita di unicità del caso)

Mitigazione: supervisione clinica, verifica, linguaggio prudente, tracciabilità delle fonti/dati.

Template di prompt (uso corretto)

Ruolo: “Agisci come supporto alla riflessione clinica (non diagnosi).”

Input: “Ecco dati/testi anonimizzati + punteggi (se presenti).”

Compito: “Evidenzia pattern e sintetizza in modo neutro.”

Output: “3 ipotesi alternative + domande di verifica per ciascuna.”

Vincoli: “Segnala assunzioni e incertezze; evita etichette diagnostiche.”

Suggerimento: chiedi esplicitamente “usa un tono prudente” e “non usare dati identificativi”.

Messaggio chiave

L'AI può supportare la lettura dei dati.

La comprensione della personalità resta un atto clinico relazionale.

È possibile usare la AI
generativa
per interpretare i test
psicologici?

Cosa può fare (uso appropriato)

Riassumere punteggi già calcolati e descrittori forniti

Organizzare profili per domini/scale (es. Big Five, PID-5, MMPI, TCI...)

Tradurre descrizioni tecniche in linguaggio comprensibile (restituzione)

Strutturare un report a partire da dati forniti

Evidenziare incoerenze numeriche o punti da verificare

Principio: lavora sui dati forniti; non “vede” il contesto clinico reale.

Cosa NON può fare

Formulare diagnosi cliniche valide

Integrare in modo affidabile test + osservazione clinica + anamnesi

Valutare validità di risposta con standard psicometrici (in autonomia)

Assumersi responsabilità professionale

Sostituire competenze psicodiagnostiche e supervisione

Regola: l'interpretazione clinica resta responsabilità dello psicologo.

Rischi concreti

Overinterpretazione automatica (andare oltre i dati)

Linguaggio troppo assertivo → “effetto autorità tecnologica”

Bias del modello o dei dati di input

Scarsa aderenza a norme/campioni di riferimento se non esplicitati

Possibili violazioni di copyright (manuali e contenuti proprietari)

Mitigazione: prudenza, alternative, verifica manuale, attenzione a proprietà intellettuale.

Uso corretto (best practice)

Fornire solo punteggi/descrittori sintetici (non protocolli completi)

Anonimizzare completamente e minimizzare i dati

Chiedere esplicitamente un tono prudente e dichiarazioni di incertezza

Richiedere ipotesi alternative e domande di verifica

Revisionare sempre e confrontare con manuali/linee guida pertinenti

Obiettivo: supporto alla redazione e alla riflessione, non “automatizzare” la psicodiagnosi.

Template di prompt (più sicuro)

Ruolo: “Agisci come supporto alla redazione di un report psicodiagnostico (non diagnosi).”

Input: “Ti fornisco punteggi già calcolati + descrittori sintetici e contesto anonimo.”

Compito: “Riassumi il profilo in modo neutro e prudente.”

Output: “Aree di vulnerabilità e risorse + 2–3 ipotesi alternative + domande di verifica.”

Vincoli: “Non formulare diagnosi; segnala limiti, assunzioni e incertezze.”

Nota: evita di incollare materiale protetto (manuali) o dati identificativi.

Conclusioni

L'AI può aiutare a scrivere un report.

L'interpretazione clinica resta competenza e responsabilità dello psicologo.