

AI in neuropsychological practice: what to check before you use it

Reference sheet accompanying *The Jagged Frontier of LLMs: False Confidence, Bias, and Ethical Risks in Neuropsychological Practice* | Dr Bridget Regan | CCN Conference, Hobart, 31 July 2026

THE FIVE PRINCIPLES — APS (Dec 2025) and AHPRA, for all registered practitioners

Accountability	You own every output. Professional accountability remains yours, whatever assisted.
Understanding	Know how the tool is trained and where its limits are.
Transparency	Disclose AI use; inform clients.
Informed consent	Per tool, per session, in writing.
Ethical and legal	Privacy, bias, legislative compliance. Peer-reviewed evidence before any clinical deployment.

APS CLAUSE MAP — Professional practice guidelines for the use of AI and emerging technologies

Clause	Applies to	What it goes to
1.3, 3.6, 4.3–4.5, 6.2	General governance	The five principles above; evidence base before deployment; check your PI cover.
2.2	Confabulated output	Verify every source; confirm citations are real.
2.2, 3.3	Legitimate uses	AI assists within competence — it does not extend it.
3.4, 6.1	Deskilling, scope creep	Human oversight is not optional; accountability is unaffected by assistance.
5.2, 5.9, 8.2–8.3	Privacy, vendors	Vendor Privacy Act compliance in writing before any client data is entered.
7.1, 7.3	Bias, diversity	Account for population diversity; RLHF reflects majority-population raters.
8.1–8.3	Medicolegal work	Consent, disclosure, and the limits below.

MEDICOLEGAL — a continuously changing landscape; check the Practice Notes in your jurisdiction

NSW	Generative AI currently requires prior leave to use when preparing an expert report.
Vic, Qld, Federal Court	Mandatory disclosure where AI has been used.
All jurisdictions	Written consent per assessment. Never input raw test data. Confirm your PI insurer covers AI tool use — Aon (July 2026), asked directly: cover applies where the final opinion is the clinician's own.

The clinical opinion is yours alone. AI does not produce expert evidence.

WHERE IT HELPS, WHERE IT DOES NOT

Reasonable, with review	Not established / actively cautioned	Never
Session transcription and summary — reviewed, never accepted	AI-generated medicolegal reports	Test materials and protocols — in any context, any tool
Literature orientation — every citation verified independently	Autonomous clinical formulation	Patient information in medicolegal work
Plain-language versions for clients and referrers	AI-assisted test interpretation without human oversight	
Structural scaffolding and readability — the clinician writes, AI polishes	Patient information without consent and written vendor compliance	

THE FOUR FAILURE MODES

1. Confabulation, not lying	Fabricated citations; scores mapped to the wrong cognitive domain. The danger is not that it sounds wrong — it sounds exactly right.
2. Embedded bias	Interpretation skewed to majority norms; cultural variation misread as deficit; non-preferred or stereotyped language. Bias does not announce itself.
3. Privacy and discoverability	Cloud tools retain what you feed them; no Australian-hosted models, so it sits offshore. In medicolegal work, what enters the scribe may be discoverable.
4. Deskilling and scope creep	The “just this once” pattern: transcription → summary → formulation → opinion. Each step feels incremental; cumulatively AI shapes judgement.

SOURCES AND FURTHER READING

Australian Psychological Society. *Professional practice guidelines for the use of AI and emerging technologies* (December 2025). AHPRA principles for registered practitioners.

Lindsey, J., Gurnee, W., Ameisen, E., Chen, B., Pearce, A., Turner, N. L., Citro, C., et al. (2025). *On the Biology of a Large Language Model*. Transformer Circuits Thread, 27 March 2025.

Lindsey, J., Gurnee, W., Sofroniew, N., et al. (2026). *Verbalizable Representations Form a Global Workspace in Language Models*. Transformer Circuits Thread, 6 July 2026.

Needham, J., Edkins, G., Pimpale, G., Bartsch, H., & Hobbhahn, M. (2025). Large language models often know when they are being evaluated. *arXiv preprint arXiv:2505.23836*.

Regan, B. (2026, 24 March). All frontal lobe, no brain: What large language models share with the human frontal lobe – and what they don't. *Jagged Profiles*. jaggedprofiles.substack.com/p/all-frontal-lobe-no-brain

Regan, B. (2026, 12 July). Believed by Roosevelt, doubted at home: How do you catch a fake when there may be no fake? *Jagged Profiles*. jaggedprofiles.substack.com/p/believed-by-roosevelt-doubted-at

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). *Attention Is All You Need*.