

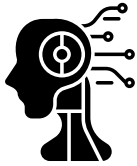
A.I. IS INCREDIBLE:

DAVID'S 7 MOST IMPORTANT TIPS
FOR HOW TO USE IT SAFELY
AND EFFECTIVELY

UPDATED: JULY 28, 2026



nonprofithelp



Written by a human (David Hartley of nonprofithelp.ca) with 25+ hours of research, checks and revisions using multiple leading AI tools — and verified against primary sources. This is the first in David's series of documents on this topic. It draws on the IPC and OHRC's 2026 Principles for the Responsible Use of Artificial Intelligence and the IPC's January 2026 guidance on AI scribes in health care.



TIP 1: Match your checking to the stakes — and verify important work

A) ERRORS/OUTDATED INFO: AI is unbelievably helpful, but every AI tool makes regular errors, big and small. They "hallucinate" — confidently invent facts, sources or details that sound completely real. Separately, AIs with memory features can drag stale content from a chat you had weeks ago into today's work. (AI scribes raise their own distinct issues — consent, privacy and accuracy checking — and will be covered in a separate document in this series.) Often the errors are subtle, and your own human brain will have no idea something is wrong or outdated.

B) INCOMPLETE INFORMATION: One AI gives you one valuable perspective. A second AI often catches what the first missed. But AIs are trained on similar data and are agreeable by design — they can, and do, agree on the same wrong answer. Agreement between AIs is a good screen, not proof.

C) KNOW WHAT YOU ARE ACTUALLY USING: "ChatGPT," "Claude" and "Copilot" are product names — not descriptions of the AI doing your work. Behind each product sits a subscription plan, a model, and an effort (reasoning) setting, and each one changes what you get. For important work, know all four: product --> plan --> model --> effort setting.

In my experience, Claude is still my first pick for writing that sounds like a human who cares, and for polished final Word documents.

ChatGPT is my challenger — I use it to attack the draft: hunt for factual errors, unsupported claims, omissions and contradictions.

For current facts, use any leading AI's web search or research mode — then open and read the sources it cites. AI helps you find the

evidence; it is not itself the evidence.

The leaders trade places every few months. The checking habits in these tips matter more than the brand names.

David's current ai setup (as of July 28, 2026 — this will change!): 1) Claude → buying Max plan → then in the chatbox choosing Fable 5 (or Opus 5) model with High effort. 2) ChatGPT → buying Plus plan → then in the chatbox choosing GPT-5.6 Sol model with High reasoning. You don't need that exact setup but it is really good. For everything in the Green Zone (page 2), the approximately \$20/month tier of either tool is likely plenty — and any paid tier beats the free version. These are my current choices, not permanent rankings.

D) DOWNTIME: Every AI has slow, overwhelmed periods — having a second AI open is very helpful.

David's RULE = scale your checking to the stakes. Routine, low-risk drafting (emails, agendas, rewrites): one strong AI plus your own careful read. Important non-confidential work: have a second strong AI check the draft against your source documents, hunting for errors, unsupported claims, omissions and contradictions (see Tip 6). High-stakes work: also verify the key facts against a primary source (the actual legislation, the Ministry document, the IPC website) and get knowledgeable human review. Agreement between AIs is a useful screen — not proof.



nonprofithelp

SCROLL TO NEXT PAGE ↓



TIP 2: Taking care when using any AI

This page is about general-purpose AI chat tools (ChatGPT, Claude, Copilot Chat, Gemini, Perplexity and similar). Purpose-built clinical tools your FHT has formally approved – such as an approved AI scribe – follow different rules and are covered in Tip 7.

● GREEN ZONE: LOWER RISK* & HELPFUL

- Emails, agendas, summaries (*never provide AI with any sensitive data)
- Plain-language rewrites
- Templates (policies, workflows, forms)
- Non-patient education materials
- Brainstorming & ideas

● YELLOW ZONE: SLOW DOWN AND VERIFY

- Ministry or Board documents
 - QIP (Quality Improvement Plan) narratives
 - HR or performance language
 - PHIPA / OHIP questions
- ➔ **If it is important:** verify against the original source document (the legislation, IPC guidance, your Ministry agreement) or the right professional – your Privacy Officer, lawyer or accountant. AI cross-checking is a supplement, never the verification.

● RED ZONE: NEVER IN A GENERAL-PURPOSE AI TOOL

- Names, DOBs, OHIP numbers, any identifiers
- EMR notes, labs, consults
- Recordings or transcripts of patient visits, clinical conversations or confidential meetings
- Identifiable or unique patient or staff stories

Keep out of personal or public AI accounts (unless your FHT has specifically authorized the tool and the use): budget drafts • HR investigations or performance concerns • union issues • vendor contracts • internal conflict documentation. De-identified summaries are safer than raw text – but a summary of a unique situation can still identify a person. Tip 7's test applies to summaries to.



ONE MORE WARNING= A paid personal subscription is not an organization-managed workspace. Words like “Plus,” “Pro” and “Max” describe price and power – not privacy. Consumer plans have different data-training and retention arrangements than enterprise workspaces, they vary by vendor, and they change often. Check the actual settings (Tip 3).

GENERALLY LOWER RISK FOR AI	NEVER OK IN AN UNAPPROVED GENERAL-PURPOSE AI TOOL
General research	Anything at all that might possibly identify a patient, staff member or provider
Generic drafting	OHIP numbers, SINS, banking info – any info you wouldn't want found by accident on the floor of your clinic waiting room
Brainstorming	Screenshots containing any patient, staff or internal information – and check the edges: screenshots capture open tabs, notifications and email previews you didn't notice
De-identified examples	Internal HR or conflict issues that are not general in nature



TIP 3: Change the privacy/data settings — on every account



Personal AI accounts usually offer data controls; organization-managed accounts also have administrator, retention and compliance settings. Two rules: (1) On personal accounts, turn off model-training and data-sharing options wherever they exist. (2) On work accounts, have your IT or privacy lead confirm which protections actually apply. And do not ask a chatbot to describe its own privacy settings — AI models are unreliable narrators about themselves, because settings change after the model was trained. Go to the vendor's official settings and privacy pages (search "[tool name] data controls"), and re-check every few months — vendors change these often.



TIP 4: AI does not know your context unless you tell it

Without direction, AI defaults, for example, to a generic U.S. healthcare model.

CONTEXT TO PROVIDE:

- Ontario Family Health Team / interprofessional primary care model
- Rural, northern, or underserved population realities
- Publicly funded, and accountable to the Ministry under your funding agreement (Schedule A context)
- Ontario privacy and PHIPA expectations (without providing PHI)
- Your audience, your purpose, and anything the AI must not assume

ASSIGN A ROLE AND CONSTRAINTS:

Generic prompts produce generic answers. Example: "Act as an experienced Ontario Family Health Team Executive Director. Draft a response to a patient complaint about wait times. Use plain language, acknowledge the frustration, explain triage without making unsupported claims, avoid defensive wording, and flag anything I should verify before sending."

David's Rule = AI forgets, sometimes painfully in the middle of a long chat. Most AI tools now have a settings area for standing context (often called custom instructions, preferences or memory) — put your constant context there so you don't retype it for every chat. Don't put anything in there that you wouldn't post on a patient waiting room wall.



TIP 5: Match the mode to the task — and know which model you're on

The names differ by product, but most leading AI tools now offer faster settings and deeper reasoning settings.

Fast / routine setting: emails, agendas, rewrites, templates, job postings.

Deeper reasoning / higher effort setting: policy analysis, Schedule A funding reviews, QIP narratives, risk identification, governance materials.

For important work, check both the exact model and the effort level — and never assume the default is the strongest option your subscription includes.

Leadership note: if the task involves regulatory interpretation, avoid "fast" mode.





TIP 6: The biggest risk is treating AI output as final

For anything important:

- **Keep an accountable human in charge of the final document or decision.** A subject-matter expert — your Privacy Officer, your lead physician, your ED, your accountant — is the single strongest check on this list.
- Verify key claims against a primary source — and open and read the sources rather than trusting the AI's description of them.

- Ask a second strong AI to hunt specifically for factual errors, unsupported statements, omissions, contradictions and implementation problems — and arm it properly (see David's Rule below).
- Protect confidential and sensitive information throughout.

And keep it proportionate: routine drafting doesn't need all of this. Save the full treatment for documents where being wrong actually costs you.

DAVID'S RULE = GIVE THE CHECKER THE SOURCES. Passing a draft between AIs without the source material is a comfort ritual: the second AI can only judge whether the draft sounds plausible — and plausible is exactly what a good error sounds like. Instead, give the checking AI the draft AND the source documents, with a hunting brief: "Check every claim, quote and attribution in this draft against the attached sources. Flag anything stated more strongly than the original said it, anything attributed to the wrong person or document, and any claim with no supporting passage." Two cautions. (1) AIs tend to err in the same direction — for example, hardening "I sometimes wonder about coverage" into "staff expressed concern about coverage" — so even here, agreement is a screen, not proof; your own read of the sources is the final check. (2) Every AI you hand source material to is another place that material lives. De-identify first, and confirm the data settings on every account you use (Tip 3).



TIP 7: PHIPA is a non-negotiable red line



The stakes went up on January 1, 2024: the Information and Privacy Commissioner of Ontario can now directly impose administrative monetary penalties of up to \$50,000 for individuals and \$500,000 for organizations that contravene PHIPA — and the most serious cases can still be prosecuted, with fines of up to \$200,000 for individuals and \$1,000,000 for organizations. The IPC issued its first penalties in 2025 (against a physician and his private clinic) and its second in April 2026 — a \$2,000 penalty against a hospital clerk, personally, for snooping in patient records. Front-line staff, not just organizations, are on the hook. (Source: Information and Privacy Commissioner of Ontario.)

De-identification is NOT simply removing a patient's name. Under PHIPA, information is de-identified only when there is **no reasonable possibility of identifying an individual**, either alone or **in combination with other information**. Context can identify someone even when specific fields are removed.

TWO KINDS OF AI, TWO SETS OF RULES:

GENERAL-PURPOSE AI CHAT (ChatGPT, Claude, Gemini, Perplexity, Copilot Chat): unless your FHT has formally approved a specific tool and use case, limit these to templates, frameworks, and public, generic or properly de-identified information.

PURPOSE-BUILT CLINICAL AI (for example, an approved AI scribe): personal health information may be processed

only where the FHT, as health information custodian, has approved the specific tool and workflow after proper privacy and security assessment, vendor due diligence, contractual safeguards, consent procedures, access controls, staff training and ongoing monitoring. Ontario's provincial AI Scribe Program can help with procurement — but it does not replace your own due diligence. The custodian remains responsible.

IMPORTANT: Even in the protected version of Copilot (green shield showing on your screen), **PHI stays out**. The shield is an indicator of data protection — it is not approval to enter PHI, and it is not your FHT's approval of the tool.



FINAL NOTE: THE “GREEN SHIELD” DISTINCTION

If your FHT uses Microsoft 365, staff need to understand the two versions of Copilot:

1. Consumer / public Copilot (not signed in with your work account)

- Not governed by your FHT’s Microsoft 365 administrator, retention and compliance controls.
- Data may be used for training by the AI company.



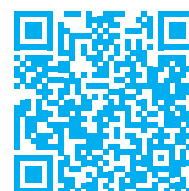
2. Protected Copilot Chat (green shield showing, signed in with your work account)

- The green shield is Microsoft’s indicator of Enterprise Data Protection: prompts and responses are handled under your organization’s Microsoft 365 protections and are not used to train public AI models.
- The green shield does not mean the AI is smarter; it means your data is more protected. And it does not mean “approved for everything” — it does not by itself authorize PHI or other confidential information.

THE RULE: Your FHT should decide centrally which AI tools and uses are approved, and set up staff accounts so the protected version is the only version staff land in. Make the safe path the default path — individual vigilance is the backup, not the plan. Either way: no PHI in Copilot, shield or no shield.

ONE LAST THING — WHO DOES WHAT. Tips 1–7 protect individuals. Protecting your organization takes a governance layer, with the right jobs in the right hands.

Your board: ensures AI is on the risk register, sets governance expectations, and receives reporting on significant uses, risks and incidents. **Your management:** approves tools and use cases, completes privacy, security and vendor assessments, sets operating procedures, trains staff, and responds to incidents. If your FHT doesn’t have those pieces yet, that’s the first fix.



SCAN THE QR CODE or click here for **David’s up to date current AI tools.**

BRING THIS TO YOUR BOARD. David delivers a 90-minute AI governance session for FHT boards — in person preferred, virtual available — paired with a ready-to-adapt AI Use Policy package for your team. Your board leaves with a clear policy, a risk picture, and confidence instead of anxiety. Email david@nonprofithelp.ca to hold a date.



MEET DAVID David has assisted **over 55 Family Health teams** to date, providing primarily: **(1) Strategic Planning/dashboards** including patient engagement; **(2) Board governance evaluation and training;** **(3) risk management assessment/dashboards;** **(4) all staff or small team retreats.**

Please see our website for **pricing and 200+ testimonials:** nonprofithelp.ca, along with **David’s full bio and references.**

IMPORTANT LEGAL DISCLAIMER: David Hartley (nonprofithelp.ca), who has assisted over 55 FHTs to date, developed this series using the IPC and OHRC’s AI guidance for Ontario, primary-source research, and multiple leading AI models over 25+ hours; most recently updated July 28, 2026, based on publicly available information as of that date. It provides governance guidance — not legal advice, a compliance certification, or an endorsement of any vendor. Before approving an AI use, an FHT should complete privacy, security, contractual and human-rights assessments proportionate to the information, tool and risks involved (a decision record showing you thought it through — not just paperwork). Where personal health information is involved, consult your Privacy Officer, complete an appropriate privacy impact assessment, and obtain legal or specialist advice where required.