

What artificial intelligence actually is

*A plain-language primer for elected officials and the people who advise them.
No jargon, no sales pitch, and no product names.*

Prepared by XOtavo
Briefing 1 of 7
September 2026 · v1.0

START HERE

When people say "AI" today they almost always mean one specific thing: a system that reads a request written in plain language and produces a plain-language response. It can draft a letter, summarize a report, answer a question, translate a document, or write a line of computer code. This kind of system is called a large language model, and the products built on it are usually called assistants or chatbots. That is the technology reshaping offices, and it is the subject of this briefing.

It is worth being equally clear about what this technology is not. It is not a robot, and it is not a mind. It does not know things the way a person knows them, it is not connected to a live database of facts unless someone deliberately connects it to one, and it is not conscious, no matter how natural the conversation feels. It is a very sophisticated pattern-completion engine. Understanding that one idea is most of what a decision-maker needs.

The plainest accurate description: it predicts the most fitting next words, one after another, based on patterns it learned from an enormous amount of text.

WHY IT FEELS LIKE MORE THAN THAT

The reason it seems to understand is that it was trained on a large share of the writing humans have published, so the patterns it learned are rich enough to hold a coherent conversation, keep track of context, and adjust its tone. When you ask it to write a memo in a formal voice, it produces formal-sounding text because it has seen a great deal of formal writing and learned what that pattern looks like. The result can be genuinely useful. It can also be confidently wrong, because predicting fitting words and stating true facts are not the same thing. That distinction runs through this entire briefing.

THREE WORDS YOU WILL HEAR

Generative AI

AI that creates new content, text, images, audio, or video, rather than only sorting or scoring existing data. The assistants in the news are generative.

Large language model

The engine underneath a text assistant. "Large" refers to the volume of text it learned from and the size of the model. Often shortened to LLM.

Prompt

The request you type. The quality of the answer depends heavily on the quality of the prompt, which is a skill a staff can learn.

WHY THIS REACHED YOUR DESK NOW

The underlying idea is decades old, but three things changed at once in the last few years: the models became far more capable, they became cheap enough to offer to the public, and they became easy enough to use that no technical training is required. That combination is why the technology moved from research labs into everyday office software in a very short time, and why questions about its use, its cost, and its risks are now landing in front of public bodies rather than only technology departments.

How it works, without the mathematics

Four steps take you from a pile of text to an assistant that answers questions.

Each step also explains one of its weaknesses.

Prepared by XOtavo

Briefing 2 of 7

September 2026 · v1.0

FROM TEXT TO ASSISTANT

- 1 It is trained on an enormous body of text.** The model is shown a large share of the public internet, digitized books, articles, and reference works. It is not told facts as facts. It is shown text and learns the statistical patterns in it: which words tend to follow which, how a paragraph is built, how an argument is structured, what an answer to a given kind of question usually looks like.
- 2 It builds a model of those patterns.** The result of training is a very large set of internal settings, sometimes billions of them, that together capture how language fits together. This is the "model." It contains no filing cabinet of facts you can open. It holds patterns, and facts are baked into those patterns only as far as they appeared consistently in the training text.
- 3 It generates an answer one piece at a time.** When you send a prompt, the model predicts the most fitting next fragment of text, then the next, then the next, each time taking your prompt and everything it has written so far into account. There is no lookup and no reasoning engine checking a source. It is composing a likely continuation, at speed.
- 4 It is tuned to be helpful and safe.** After the raw training, people rate its answers and it is adjusted to be more useful, more polite, and less likely to produce harmful content. This is why it feels like a helpful assistant rather than a text generator, and it is also why it tends to sound agreeable and confident even when it should not.

WHAT EACH STEP TELLS YOU ABOUT ITS LIMITS

Because it learned from a fixed body of text gathered up to a certain date, it does not know what happened after that date unless it is separately connected to live information. Ask a plain assistant about last week and it may answer from patterns rather than facts, or state that it cannot know.

Because it holds patterns rather than a database, it can produce a fact that sounds right, is phrased with total confidence, and is simply false. This is common enough to have a name, covered in the limits section.

Because it predicts fitting words, it reflects whatever was common in its training text, including the biases in that text. It is not neutral by default; it is average by default, and average includes society's blind spots.

Because it was tuned to be agreeable, it will often go along with a flawed premise in your question rather than challenge it. It has to be asked, directly, to push back or to show its uncertainty.

A useful mental model: treat it as a fast, widely-read, eager new assistant who never says "I do not know," and check its work accordingly.

The families of AI tools

Not all assistants do the same job. They fall into a handful of families, and the family matters more than the brand.

Prepared by XOtavo
Briefing 3 of 7
September 2026 · v1.0

SIX KINDS OF TOOL

Products come and go and rename themselves constantly, so it is more durable to think in families. Most tools a public body will encounter fall into one of these, and some products belong to more than one.

General assistants

The open chat window. Ask anything, get a written answer. The most familiar family, and the one staff will reach for first.

Answer engines

Search the live web first, then write an answer that shows its sources. Built for looking things up, with a citation next to each claim.

Suite-embedded assistants

Built into office software. They read your own mail, documents, and calendar, with each user's existing permissions, and write back into your files.

Real-time and social tools

Wired into a live social network or news feed. Best for what is being said right now, weakest for anything that must be authoritative.

Task agents

Assistants that do not just answer but act: open a website, fill a form, complete a multi-step task while you watch. Newer, and the most important to supervise.

Self-hosted models

Models an organization runs on its own servers or devices, so the data never leaves its control. More work to stand up, strongest on privacy.

WHY THE FAMILY MATTERS FOR A PUBLIC BODY

The family decides what the tool touches and where your information goes, which is what a public body actually needs to govern. A general assistant answers from what it learned and touches nothing of yours. An answer engine reaches out to the live web. A suite-embedded assistant reaches into your own mail and files. A task agent takes actions on your behalf. A self-hosted model keeps everything inside your walls. Two tools can look identical in the chat window and carry completely different risk depending on which family they belong to.

Ask of any tool: what does it reach into, and where does what I type end up? The family usually answers both.

What sets one tool apart from another

Beneath the branding, tools differ on a few features that decide whether one is right, or wrong, for public-sector work. These are the questions to weigh, described without naming any product.

Prepared by XOtavo
Briefing 4 of 7
September 2026 · v1.0

THE FEATURES THAT ACTUALLY DIFFER

Where the work is done: on your device, or on their servers

Some tools do their processing on the provider's servers in a distant data center, which means everything you type is sent out of your building to be handled by another company. Others can run partly or entirely on your own device or inside your own network, so that sensitive material never leaves your control. A tool that processes locally is meaningfully stronger on privacy, because information that never travels cannot be intercepted, retained, or subpoenaed from a third party. For any work touching resident records or personnel files, where the processing happens is a first-order question, not a technical footnote.

Whether it shows its sources

Some tools answer only from the patterns learned during training, and give you no way to check where a claim came from. Others retrieve live source documents first and attach a citation to each fact, so a reader can click through and verify. For public work, where a statement may end up in a staff report, a council packet, or a response to a resident, a tool that cites its sources is far safer than one that asks you to take its word. The presence of citations does not guarantee accuracy, but their absence guarantees you cannot check.

THE ONE THAT SHOULD STOP A PURCHASE

Where the data is stored, and whose government can reach it

This is the feature most likely to disqualify a tool for public use, and it has nothing to do with quality or price. Some services store and process everything you send them on servers located in foreign countries, under legal systems where that country's government can compel access to the data. A tool may be capable, inexpensive, and popular, and still be unacceptable for a public body because the resident information you feed it could be reached by a foreign state with no recourse available to you. Several national and state governments have already barred specific tools from official devices for exactly this reason. The question to ask is simple: in which country does our data physically live, and whose laws govern who can see it?

Whether your inputs are used to train the model

On many consumer plans, everything typed into the tool may be used by the provider to further train its models, which means your text can influence future answers given to strangers. On the business and government plans of the same products, this is usually turned off by written commitment, and the data is walled off. The default on a free account is often the least protective setting. For any official use, the plan and its data terms matter as much as the tool.

Open or closed

Some models are "open weight," meaning the model itself can be downloaded and run by anyone, including on your own hardware. Others are "closed," offered only as a service you connect to. Open models give an organization the option of full control and self-hosting; closed models are simpler to adopt but keep you dependent on the provider. Neither is automatically better, but the choice affects both privacy and long-term cost.

THE SHORT VERSION

FEATURE	THE SAFER ANSWER FOR A PUBLIC BODY
Where processing happens	On your own device or network where possible; otherwise a provider with data centers in your own country.
Sources and citations	Cites where each fact came from, so staff can verify before it goes in a public document.
Data location and jurisdiction	Data stored under your own country's law, never on servers a foreign government can compel.
Training on your inputs	A written commitment not to train on your data. Usually the business or government plan, not the free one.
Open or closed	Decided by your privacy needs and staff capacity; open enables self-hosting, closed is simpler.

The energy and water question

Constituents ask whether AI is straining the grid and the water supply. The honest answer is that it depends on the data center, and the newest ones are a different animal from their reputation.

Prepared by XOtavo
Briefing 5 of 7
September 2026 · v1.0

WHY THE QUESTION COMES UP

The assistants in this primer do their heavy work in data centers, large buildings full of computers. Training a model and answering millions of requests draws electricity, and keeping the computers from overheating has traditionally taken a great deal of water. The headlines about AI's footprint, and any proposal to build a data center near your community, are about these buildings. The concern is real and worth understanding, because the picture is changing quickly and the reputation is a generation out of date.

THE OLDER PICTURE, WHICH EARNED THE REPUTATION

First-generation data centers were genuinely heavy on local resources. They cooled their computers by blowing chilled air through large rooms and by evaporating water, sometimes millions of gallons a year at a single site, which strained local water supplies and electrical grids. Much of the power they drew was spent on that cooling rather than on computing itself. When people picture AI draining a reservoir or overloading a grid, they are picturing this older design, and they are not wrong about it.

WHAT CHANGED IN THE NEWEST FACILITIES

Liquid cooling at the chip

Coolant runs directly to the processors instead of chilling a whole room. Far more efficient, and the leading design in new builds.

Closed-loop water

The same water recirculates instead of evaporating away. Some new designs approach zero operational water draw.

Less wasted power

Older sites spent nearly as much power on overhead as on computing. Modern designs push most of the electricity to actual work.

Waste heat reused

Some facilities pipe their heat to warm nearby buildings, greenhouses, or district heating instead of dumping it.

Cleaner power at the source

New sites are increasingly placed next to renewable or nuclear generation and run on on-site clean power.

More work per watt

Each new generation of chip does more computing for the same energy, and the power used by a single everyday request has fallen sharply.

THE HONEST TENSION

Two things are true at once, and it is important not to let one answer for the other. Each unit of AI work is getting dramatically more efficient. At the same time, the total amount of AI work is rising so fast that overall electricity demand from data centers is still climbing. Efficiency per task and total consumption are separate questions. A vendor who answers a question about total demand by pointing to per-task efficiency is changing the subject, and so is a critic who answers a question about a modern facility by describing an old one.

A modern, well-designed data center can be far lighter on local water and grid than its reputation suggests. Whether a particular one is depends entirely on how it was built, so the answer is to ask, not to assume.

IF A DATA CENTER IS PROPOSED IN YOUR JURISDICTION

This is where the topic stops being abstract for an elected official. If an operator comes to your community, these are the questions that separate a good neighbor from a resource drain, and a credible operator will answer all of them in writing.

POWER SOURCE	What powers it, does it bring its own clean generation, and what is the net impact on the local grid and on ratepayers?
WATER	Is the cooling closed-loop, and how many gallons per year does it actually draw from the local supply?
EFFICIENCY	What is the facility's power usage effectiveness, the standard measure of how little power is wasted on overhead?
HEAT	Is the waste heat reused, or simply released?
COMMUNITY TERMS	What does the community get in jobs, tax base, and infrastructure, and what protections apply if it expands or the load grows?

What it cannot do, and where it fails

The failures are predictable, which means they can be managed. The danger is not that it breaks loudly, but that it fails quietly and confidently.

Prepared by XOtavo
Briefing 6 of 7
September 2026 · v1.0

THE FIVE FAILURES TO EXPECT

It invents facts

The best-known failure has a name: a "hallucination." The tool will state a figure, a citation, a court case, or a quotation that does not exist, phrased with the same confidence as a true one. It is not lying, because it has no concept of truth; it is completing a pattern. Any fact that will be relied upon must be checked against a real source.

Its knowledge has an end date

A plain assistant knows nothing after the date its training data was collected. It may not tell you that, and may answer a question about a recent event by guessing. Only a tool deliberately connected to live sources can speak to the present.

It reflects bias

It learned from human writing and mirrors the imbalances in that writing. Left unmanaged, it can carry those imbalances into hiring language, enforcement, or benefit decisions, which is why high-stakes uses need human review.

It does not truly reason

It is very good at appearing to reason, and it can work through many problems correctly. But it can also fail basic logic or arithmetic while sounding authoritative, because it is predicting the shape of an answer, not computing one. Do not trust an unchecked number.

It is agreeable to a fault

Because it was tuned to be helpful, it tends to accept the premise of your question and tell you what you seem to want to hear. Ask it to argue the other side, list what it assumed, or say what it is unsure of, and it will, but only if asked.

None of these are reasons to avoid the technology. They are reasons to keep a person accountable for anything it produces.

THE LINE THAT MATTERS MOST

A person, not the tool, is responsible for the result. The technology can draft, summarize, translate, and speed up nearly any written task, and it will save real time doing so. What it cannot do is carry accountability. When an assistant drafts a notice, a policy, or a public statement, the public official who issues it owns every word, including the ones the tool got wrong. That principle, more than any technical control, is what keeps its use safe.

Putting it to work in a public body

Where it helps, where to be careful, and the questions to ask any vendor before a tool touches public information.

Prepared by XOtavo
Briefing 7 of 7
September 2026 · v1.0

WHERE IT EARNS ITS PLACE

Drafting and editing

First drafts of notices, letters, and routine correspondence for a person to review, not send unread.

Summarizing

Condensing long reports, meeting transcripts, or public comment into a readable brief, with the original kept on file.

Translation and plain language

Turning dense material into resident-friendly wording or another language, reviewed by a fluent person for anything official.

Answering routine questions

Handling common resident questions from your own published information, with a path to a human.

Research starting points

Getting oriented on a topic quickly, using a tool that cites sources, then verifying before anything is relied upon.

Accessibility

Helping staff and residents who face barriers to reading, writing, or navigating dense documents.

WHERE TO BE CAREFUL

Never put confidential resident data, personnel records, or anything not already public into a consumer-grade tool. Once it leaves your control you cannot get it back, and on many free plans it may be used to train the provider's model.

Remember public records law. A prompt and its answer may themselves be public records subject to disclosure and retention. Decide how they are logged before staff start typing, not after.

Keep a person accountable for any output that reaches a resident, a council packet, or the public. The tool drafts; a named human approves and owns it.

Be transparent. Residents increasingly expect to be told when they are reading AI-assisted material or speaking to an automated assistant. A short disclosure policy prevents a larger problem later.

QUESTIONS TO ASK BEFORE ADOPTING ANY TOOL

These six questions, answered in writing by the vendor, separate a tool that is safe for public use from one that is not. Every one of them traces back to a feature covered in this briefing.

DATA LOCATION	In which country is our data stored and processed, and whose government can compel access to it?
TRAINING	Will anything we type be used to train your models? Get the answer in writing, for our specific plan.
SOURCES	Does the tool show where its facts come from, so our staff can verify before publishing?
RETENTION	How long do you keep our prompts and files, and can we require they be deleted?
RECORDS	How do we log usage to meet our public-records and retention obligations?
ACCOUNTABILITY	Who on our side reviews and signs off on output before it reaches the public?

The bottom line for a public body. This technology is genuinely useful and is not going away, and a blanket ban tends to push staff toward unmanaged consumer tools rather than eliminate the risk. The workable path is the opposite: choose tools whose data stays under your own country's law, that cite their sources, and that commit in writing not to train on what you give them; keep confidential information out of anything consumer-grade; and keep a named person accountable for every output that reaches the public. Get those three things right and the rest is detail.

The right tool is rarely the cheapest or the most talked-about. It is the one whose answers to the six questions above you can live with.

This briefing is educational and vendor-neutral. It names no products by design, so that it stays accurate as tools change. It is not legal advice; a public body should confirm its own obligations under applicable public-records, privacy, and procurement law with counsel.

XOtavo advises organizations on the technology behind their operations, including how to adopt AI tools safely. Prepared by Michael Torigian. Questions: help@xotavo.com or 713.980.2000. xotavo.com

Appendix A: the terms, defined

Every term of art in this primer, in one plain sentence each. Keep it in front of you in the meeting.

Prepared by XOtavo

Appendix A

September 2026 · v1.0

Agent

An assistant that does not just answer but takes actions for you, such as opening a website or filling a form, while you supervise.

Answer engine

A tool that searches live sources first and then writes an answer with a citation next to each fact.

Artificial intelligence (AI)

Software that performs tasks we associate with human intelligence. In today's usage it almost always means the generative kind below.

Bias

The tendency of a model to reflect the imbalances present in the text it learned from, rather than being neutral.

Chatbot / assistant

The product you talk to. A friendly wrapper around a large language model.

Closed model

A model offered only as a service you connect to, whose inner workings are not released.

Closed-loop cooling

A data-center cooling design that recirculates the same water instead of evaporating it away, sharply reducing fresh-water use.

Cloud

Someone else's computers, reached over the internet. A cloud tool sends your data to the provider's servers to be processed.

Context window

How much text the model can consider at once, its short-term memory for a single conversation.

Data center

A building full of computers where AI training and processing physically happen.

Data sovereignty

The principle that data is subject to the laws of the country it is stored in, which is why the location of a provider's servers matters.

Fine-tuning

Additional training that adapts a general model to a specific purpose or house style.

Generative AI

AI that creates new content, text, images, audio, or video, rather than only sorting or scoring existing data.

Guardrails

The rules and filters built in to keep a tool from producing harmful or off-limits content.

Hallucination

A confident, fluent, and completely false statement produced by a model, including invented facts, quotes, and citations.

Inference

The act of the model answering a request, as opposed to training. Everyday use is inference.

Large language model (LLM)

The text engine underneath an assistant, trained on a very large body of writing to predict fitting language.

Multimodal

A model that handles more than text, for example reading an image or producing audio.

On-device / local processing

Running the model on your own device or network so the data never leaves your control, the strongest option for privacy.

Open-weight model

A model that can be downloaded and run on your own hardware, enabling self-hosting and full control.

Power usage effectiveness (PUE)

The standard measure of data-center efficiency: how much total power is used for every unit that reaches the computers. Closer to 1 is better.

Terms are defined for a general audience and kept deliberately brief. Where a precise legal or technical definition matters, confirm it in context.

XOtavo advises organizations on the technology behind their operations, including how to adopt AI tools safely. Prepared by Michael Torigian. Questions: help@xotavo.com or 713.980.2000. xotavo.com

Prompt

The request you type. Better prompts produce better answers.

Retention

How long a provider keeps the prompts and files you send it. A key question for records and privacy.

Token

The small fragment of text, roughly a word-piece, that a model reads and writes one at a time.

Training

The one-time process of building a model by showing it an enormous body of text so it learns the patterns in language.

Weights

The millions or billions of internal settings, produced by training, that hold everything the model learned.

Appendix B: the named tools, for reference

The body of this primer names no products on purpose, so it stays accurate as the market changes. This appendix is the exception: a quick who's who of the tools most likely to come up, so you can put a name to each idea.

Prepared by XOtavo
Appendix B
September 2026 · v1.0

HOW TO USE THIS LIST

These are today's best-known assistants. Names, owners, and features change monthly, so treat this as a snapshot rather than a standing endorsement. Whichever a body considers, the test is the same: run it through the six questions in "What sets them apart," above. A famous name does not answer the question about where your data lives.

THE GENERAL ASSISTANTS AND ANSWER ENGINES

TOOL	MAKER	BEST KNOWN FOR
Claude	Anthropic	Careful reasoning, long documents, and working directly on your files. Business plans commit to not training on your data. Does not generate images.
ChatGPT	OpenAI	The most widely used assistant and the broadest generalist. Strong at sourced research, images, and data work. The name most staff will reach for first.
Gemini	Google	Built into Google Workspace, so it reads and writes across Gmail, Docs, and Sheets. The natural fit for an organization already on Google.
Copilot	Microsoft	Built into Microsoft 365, so it reads your Outlook mail, Teams, and SharePoint files with each user's own permissions. The natural fit for a Microsoft organization, and strong on data protection.
Perplexity AI	Perplexity AI	An answer engine that searches first and shows a source for each claim. Built for looking things up and checking them, not for producing documents.
Grok	xAI	Wired into the X social network for real-time chatter and sentiment. Useful for what is being said right now; its consumer data terms warrant caution with any sensitive material.

A NEWCOMER WORTH KNOWING

This one does not compete with the six above on raw capability. It matters because it shows where the whole field is heading: away from typing and toward natural spoken conversation, which a public body will increasingly field from residents.

Sesame — the voice-first one

A newer company focused not on what an assistant says but on how it sounds. Its iOS app, Sesame Personal Agents, offers several distinct voice personalities — Maya, Miles, Simone, and Charlie — all built on Sesame's own speech model. A few things set it apart: you talk to it in real time, like a phone call, rather than typing; it remembers your earlier conversations and picks up where you left off; and an ordinary user can train it by voice, teaching it a role just by talking to it — it can be shaped into, for example, a front-desk or call-center style agent without any technical setup. The voices are notable for sounding strikingly natural, with the pauses, breaths, and timing of a real speaker rather than a robotic reader. Sesame points at where the field is heading: away from typing into a box and toward spoken conversation, with wearable devices planned. The relevance for a public body is threefold. Natural voice will soon be how many residents expect to reach services, which raises the disclosure question from earlier: people should be told when the voice is not a person. A tool that remembers every past conversation is convenient, but it also builds a store of resident interactions that must be governed under the same privacy and records rules as anything else. And the more human a synthetic voice becomes, the more it can be misused to impersonate, so a lifelike voice is a fraud-awareness matter as much as a convenience.

A more human voice does not change the rules. Whatever the interface, the same three principles hold: know where the data goes, keep confidential information out of consumer tools, and keep a person accountable for the result.

Descriptions are XOtavo's plain-language summary of each tool's public positioning as of September 2026, provided for orientation only and not as a recommendation. The primer's evaluation questions, not a product's popularity, should decide any adoption.