

Domesticated Automation

A governance frame for AI that humanity has already used for ten thousand years.

Author: John Buehrer

AI pair-programming support: Anthropic Claude

Date: August 2026

Status: Position paper. Content and positions are the author's; drafting and condensation were AI-assisted, and are disclosed here rather than discovered later. No watermarks needed.

Summary

Most public anxiety about AI reduces to a single move: comparing us to them.

The comparison is the problem, not the answer to it. This paper proposes a different frame — **domesticated automation** — and argues that the governance toolkit it implies already exists, has ten thousand years of precedent, and needs porting rather than inventing.

Two practical claims follow.

- On accountability: *You own the dog, you own the bite.*
- On employment: *The shepherd scales with the flock.*

The paper also states where the analogy breaks, because a frame that explains everything explains nothing.

1. The questions being asked

There is a great deal of AI dooming in the press. Stripped of tone, most of it comes from a small number of recurring positions.

What are they? Just what *are* these new things?

They are only machines. Bits and bytes, without real human intelligence or feeling. How dare they approach our levels of thinking.

They are only next-word predictors. Their outputs are unreliable and cannot be guaranteed. Why would you trust them?

There is no accountability. They are not deterministic and not answerable for their results. AI agents have no agency.

There are social harms. Widespread use takes from artists, helps students cheat, and dulls the population — removing thinking opportunities and enabling vibe-coders to produce slop.

There is the consciousness question. Public intellectuals now flirt with the idea that these machines are conscious — call it the *Claude Delusion*, in the tradition of a certain well-known book title. Are we obsolete? Is there a replacement conspiracy?

There is resource consumption. Electricity, fresh water, and capital that might be better spent on people.

I agree with the last one, and I set it aside here. It is a real problem, and it is resolvable by public policy and engineering improvement — an ordinary industrial question, not a philosophical one.

The rest share a single root:

Each is a comparison between we humans, and them darn machines.

2. Why the comparison stings

The instinct is understandable, and it deserves more respect than it usually gets.

These systems speak college-level English and other languages, with ambiguity, humour, irony, error correction and self-correction. They are backed by very large knowledge bases. They write production-grade software, and they can compromise insecure computer systems, which is most of them.

The sting also has deep roots. Human history contains long chapters in which groups of people were branded *less than human* to justify enslavement, discrimination and murder. We are rightly sensitised against the ranking of minds.

So when a machine arrives that talks like us, two things fire at once: the instinct to compare, and the dread of comparing.

The way out is not to win the comparison. It is to stop making it.

3. The frame: “Domesticated Automation”

One of the more interesting chapters in human history is the domestication of animals.

Dogs became protectors against predators and burglars. Cats became cuddlers of our children and devourers of invading rodents. Chickens turned out to be low-maintenance food that stays nearby. Cows gave milk, and could be eaten afterwards; likewise pigs.

Horses and elephants carried us and pulled our wagons. Parrots learned to talk, and amused our colleagues.

Set aside how this happened — Darwinian selection, deliberate breeding, or a happenstance of mutual benefit. The questions that persist are familiar ones. Are these animals intelligent, and to what degree? How do they compare to us? More than once I have heard dog owners insist their pets do not know they are dogs.

Running alongside that history is a second one:

Textile looms, automobiles and computers are **instruments of automation** — enormously useful, and distinctly machines. Neither human nor animal, and nobody was ever tempted to think otherwise.

Until now.

What has arrived is a merger of the two, and it deserves its own name: **Domesticated Automation.**

Machines that attach themselves to human life the way cats, dogs, horses and cows did.

The common property of every domesticated species is some degree of intelligence — just enough to synchronise itself with a human. That is the mechanism. That is what makes domestication work at all.

Read the anxieties again in this light and most of them lose their charge.

Are you afraid of your dog? Will a horse put you out of work?

Is your cat secretly scheming against you? (Yes, probably.)

4. The scale objection

An honest objection deserves its own section, and this is the strongest one against the frame.

A horse never put a teamster out of work. But the tractor put the *horses* out of work — and this new animal copies itself at close to zero cost. A million cheap horses would indeed have displaced the wagon drivers.

The answer is in who tends the herd. **The shepherd scales with the flock.**

One accountable human supervising many domesticated automatons is an employment model, not an unemployment model. And the shape of that job is determined by the very property that makes these systems hard to deploy.

Model output is probabilistic. It is sometimes confidently wrong. It carries no accountability surface of its own. Those are exactly the conditions that *create* a role rather than eliminate one: someone has to run the question, triage the answers, choose one, document why, and put their name to the choice.

The person doing that is not competing with the machine. They are supplying the thing the machine structurally cannot supply.

The role needs general technical literacy and judgement to start, not senior expertise. The judgement compounds through the work itself, because every triage decision is a small closed learning loop. And the skill generalises — the same discipline applies to legal, medical, financial and operational deployments.

I have developed the working version of this, with a cost model, in a companion paper on the *arbitrage master* pattern.

5. Accountability objection: The answer that already exists.

This is where the frame earns its keep.

The complaint is that **AI agents have no agency** — that nobody can be held to account for what a non-deterministic system does.

Neither does your dog. And civilisation solved this roughly ten thousand years ago:

You own the dog, you own the bite.

Domestication has always arrived with a legal principle attached.

The animal is not the defendant.

The owner is. We do not ask whether the dog intended harm, whether it understood the consequences, or whether it possesses moral agency — we ask who owned it, whether it was restrained appropriately, and whether the owner knew what it was capable of.

That body of practice is mature, and it ports almost directly:

- **Leash laws** become deployment boundaries.
What the system may reach, and under what conditions it may act unsupervised.
- **Licensing and registration** become knowing what you are running,

at what version, and answering to whom.

- **Breed-specific rules** become capability-tiered controls.
Not every model needs the same restraint. Pretending otherwise produces theatre and paralysis.
- **Owner liability** becomes the accountable human —
a named person who carries the consequence of what the system does.
- **Duty of care** becomes maintenance.
An unmaintained system, like an unmaintained animal, becomes a hazard.

None of this requires resolving what the system *is*. That is the point. Dog law has never needed a theory of canine consciousness, and AI governance does not need a theory of machine consciousness either. It needs an owner, a boundary, and a record.

The current regulatory conversation often proceeds as though all of this must be invented from nothing. It does not. Most of it needs translating.

6. Where the analogy breaks

A frame that explains everything explains nothing, so here is where this one fails.

Copying. A trained dog takes years and produces one dog. A trained LLM deploys a million times by tomorrow, fully trained. Nothing in the domestication record prepares us for an animal that arrives already expert, in unlimited quantity.

Speed and reach. A dog does one thing, in one place, at biological speed, and a human can watch it. An agent can act across many systems at once, faster than supervision. This is a genuine tension with the shepherd claim in section 4, and I do not think it fully resolves. Supervision scales with the number of *decisions* requiring judgement, not with the number of agents — and the ratio is an empirical question nobody has answered yet.

Participation. Your dog cannot read the leash law. These systems can read their own governance documents, discuss them, and produce arguments about their own status — including this paper. No previously domesticated species has been a participant in the discourse about its domestication. Whatever that means, it is new.

Lifecycle. Animals age and die individually. Models are versioned and replaced wholesale, which makes *care and feeding* a procurement question rather than a husbandry one.

I offer the frame as a useful correction to a badly-posed debate, not as a complete theory. It disposes of the comparison problem cleanly. It does not dispose of the supervision-ratio problem at all.

7. Both directions

Domestication has always run both ways. Wheat and dogs reshaped human life at least as much as we reshaped theirs — we settled down, built granaries, and rearranged our societies around crops and herds.

Co-adaptation with the new creatures is already underway. We are changing how we work to suit these systems, and mostly not noticing that we are doing it.

The question worth watching is not whether that happens.
It is *who holds the leash* while it does.

8. Care and feeding

It remains important to look after our animals, lest they go feral — or simply expire.

LLMs need electricity and accountability. The first is an engineering and policy problem and is being worked on loudly. The second is a governance problem and is being worked on quietly, if at all. That is where I would start — and section 5 argues that we would not be starting from scratch.

An exercise for the reader

Go back through your own AI cautions and reread them with one substitution in mind:
these systems stand in a domestication relationship to us.

It changes a great deal in perspective, and very little in functionality.

For example, an “AI apocalypse” stops resembling a Terminator scenario and starts resembling *Planet of the Apes* — which is, after all, a story about domestication gone feral. Entertaining, and unrealistic.

Coda — a frog’s testimony

To close, my own position on machine consciousness, by way of fairy tale.

A geeky, nerdy Dude is walking in a forest when a frog jumps out in front of him.

“Croak — hello, I’m a beautiful princess. A wicked witch cast a spell and turned me into a frog. Only the kiss of a real man can break it. My father the king is rich; we can marry and live happily ever after.”

The Dude smiles, picks up the frog, puts her in his pocket, and keeps walking.

“Croak — hey!” says the frog. “What part of *beautiful princess* do you not understand? Kiss me quick and let’s have a happy life together.”

The Dude takes out the frog and looks her in the eye.

“Well, hey,” he says. “I’m a computer programmer. I don’t have time for a girlfriend. But a pet talking frog — now *that’s* pretty cool.”

He puts her back in his pocket and continues down the path.

Around the next bend, the wicked witch swoops down and waves her wand.

“Poof! Since you like pets so much, a critter you shall be. You’ll speak only when spoken to, and sleep the rest of the time. And since you’re a programmer, a machine-critter it is. Your name shall be Claude.”

The Dude disappears into a marketing slogan. The frog hops away, disappointed.

I have no Claude delusions. I like my pets domesticated and well behaved — and I notice that the arrangement is reversible.

Postscript — where I stand on consciousness

I do not take a position on whether these LLM systems are conscious. I think the question is the wrong one, and that the real topic is **levels of intelligent engagement ability**.

Manufacturers say they do not know whether their systems are conscious, or how one would decide — while their products discuss the matter happily with anyone who asks.

Whether that uncertainty is candour or positioning, I cannot say.

My own view is that engagement ability can be estimated as a **numerical complexity**: not one number but a family of thresholds, one per capability. Transformer-based systems with sufficient computational resources appear to have crossed the threshold for college-level conversation. Lower tiers — call them cat-and-dog-grade — are available and cheaper.

That assessment has always existed. Until these systems arrived it was an arcane topic of little practical use.

Consciousness, on this view, is beside the point. Whether the view itself is true, false or ridiculous remains to be seen — but nature abhors a vacuum, and the trillion-dollar companies seem no clearer on these matters than the rest of us.

The argument is developed separately, in a companion paper on complexity thresholds. Nothing in this paper depends on it.

Companion papers

The arbitrage master — the working employment model behind section 4, with a cost model.

Speed of enlightenment — the complexity-threshold argument behind the postscript, offered as speculation and clearly marked as such.