

From: [REDACTED]
To: [Desmond, Jim](#); [Supervisor Joel Anderson District 2](#); [MontgomerySteppe, Monica](#); [BOS, District1Community](#); [Lawson-Remer, Terra](#)
Cc: [FGG, Public Comment](#)
Subject: [External] AI POLICY
Date: Friday, August 29, 2025 12:40:35 PM

Good Afternoon, Supervisors,

The use of AI is so alarming to me, that I am sending this in advance of the next Board Meeting or even agenda. It is too complex an issue to be buried in the Consent Calendar. And the unrestrained use of a partially developed, therefore flawed, and dangerous technology should be avoided if at all possible (there are better ways to conduct research, answer questions, etc.)

Short History of Artificial Intelligence - The concept goes back to 250 BC when, although not AI or near it, the Greek Ctesibius built a self-regulating water clock, and Ismail al-Jazari invented over 100 automated devices in the 1200's. Among them was a musical automaton, which had a programmable drum machine with pegs (cams) that bump into little levers that operated the percussion. The drummer could be made to play different rhythms and different drum patterns if the pegs were moved around. He is considered the father of robotics.

AI had its first public and practical use in a chess game by Leonardo Quevedo from Spain shown at a 1914 Paris expo, not 1950 Alan Turing as we'd like to think. AI uses old, very old, algorithms.

AI Safeguards – What a joke – AI STILL has no conscience, and the underlying algorithms are riddled with errors. It pretends to have a human face, but do you really want one error at the risk of it messing up the lives of 3 million County residents? Maybe your pet stakeholders and State residents too. Would you like the AG at your throats because

the AI you chose threw an innocent person in jail or someone died because of a mistaken denial of care?

From Psychology Today: Human complacency fuels AI deception. Apollo Research showed GPT-4 executing an illegal insider-trading plan and then lying to investigators about it. AI company Anthropic, along with Redwood Research, recently demonstrated that advanced AI models can produce apparently reliable answers while secretly planning to do the opposite once oversight weakens. And the busier the decision context (like a County office,) the more tempting it is to click accept and move on.

From ABC news: AI phone scam calls can now use clips of your voice as short as 3 seconds to create convincing fakes for robocalls and impostor scams. The example they give is that "Tech experts created an AI-generated voice of our own ABC7 Reporter Jason Knowles, asking, "Can you send me \$500? It's Jason, I've been in a car accident here on my vacation in Mexico. I'm okay for the most part, but I am at some hospital and they say I need to pay money for x-rays and treatment. My credit card isn't working and I don't have enough cash." But he never actually said those words."

Apple introduced its generative-AI tool in the US in June, boasting about the feature's ability to summarize specific content. It erroneously summarized a New York Times story, claiming that Israeli Prime Minister Benjamin Netanyahu had been arrested. In fact, this was when the International Criminal Court issued a warrant, but there was no arrest. It also falsely summarized a BBC report that Luigi Mangione, the suspect behind the killing of the UnitedHealthcare chief executive, had shot himself.

So AI lies too. Do you really want to promote this sort of thing? Should

the County aid these schemes as a side effect of AI when there are more reliable but slower non-AI algorithms?

Deception is a social art as much as a technical feat. AI systems master it by predicting which stories we are willing to believe.

Two people died after being denied coverage by UnitedHealth. The plaintiffs argue that the health insurance company should have known how inaccurate its AI was, and that the provider breached its contract by using it.

Their grievances are corroborated by an investigation from Stat News into UnitedHealth's internal practices at its subsidiary NaviHealth, which found that the company forced employees to unwaveringly adhere to the AI algorithm's questionable projections on how long patients could stay in extended care.

The latest is that it arrested schoolkids for crude jokes. Misidentification is routine. AI is well known for deepfakes. I have seen the same woman changed by AI from white to Asian to Black.

I myself asked the AI q/a tool my HOA's parking company uses a simple question, although not in the preset list they provide, and got a totally bogus nonanswer.

And I see ads for deepfake generators on Google.

Do we really need AI playing havoc with our lives at the County level.
Level of liability? Especially since you just reduced the actual reserves.

Do we really need this playing havoc with our lives at the County level.
Level of liability? Especially since you just reduced the actual reserves.

How much reliance do we want to put in a system that is not only flawed
but can do serious societal harm without so much as the conscience to
fix itself?

Who will control it when it starts lying, murdering, or cheating?

What you want is something to deter the bad actors as well as the
incorrect coding, logical errors, or similar.

It is good to have a policy, but this is an interstate issue. I know there
are guidelines and laws in place already which may, by the way, negate
anything the County does. You should examine these laws or guidelines
before instituting your own, given that you seem to have no idea what
financial burden this might impose. And I think that the best and safest
policy is not to use it.

You can get the same results, slower but a lot safer, by using the old
algorithms.

Regards,

Paul Henkin