



INSIGHT

Gergana Ivanova

Partner | Transactions & Corporate Finance
EY-Parthenon

Musings on the responsible use of AI

As corporate finance professionals, we are all well-acquainted with deal fatigue, but I think the concept is applicable more broadly. Many are currently experiencing AI fatigue. Therefore, this is (yet another) high-risk, high-stakes AI article sharing further thoughts on AI and its place in the world of deals and corporate finance.

So far, the conversation has largely centred around faster insights, efficiencies, more time to do 'other' things... Now, I'm coming at this mainly from the lens of a transaction diligence professional and, in my view, as much as we will experience some of these benefits, ultimately, the responsibility for the integrity of the outputs will sit with us as humans; and with the volume and speed of information coming at us, it will require true skill and mastery in being able to sift out the truth.

As much as it doesn't sound like a naturally enjoyable topic at first glance, my team and I had some fun exploring the responsible use of AI because I was keen to explore what the 'rules' for responsible use of AI may look like for a transaction diligence team. As a long-time sci-fi fan who has enjoyed reading the works of Isaac Asimov, I was keen to shape this in the image of Asimov's Three Laws of Robotics, which he explores extensively in his writing.

Here's the first cut of what we came up with as part of this thought experiment: The Three Laws of The Responsible Use of AI in Transaction Diligence (catchy!):

1. AI must never replace professional judgement.
2. AI must protect confidential information.
3. AI must be transparent and traceable.

After some discussion, we agreed that these sound sensible enough as a starting point.

Now for those who are familiar with Asimov's work, you would know that in his short stories, he explored extensively how the Three Laws of Robotics can be avoided, superseded and evaded. The question for us is whether we understand how to apply our responsible use of AI laws in practice in our day-to-day work, and whether we can identify the scenarios in which the laws are not bullet-proof and where we may need additional safeguards.

Therefore, it will be the consistency and rigor with which the laws are applied in practice that ultimately determines their success. The scenarios reach far and wide, and some are more straightforward than others. For example, corroborating facts to sources, which can mean a number of things, including:

- Does the source exist?
- Does the source say what the AI interpreted it said?
- Is it a credible source?
- Would I be comfortable referencing this source?

That's a lot of work. But, sometimes, things can be a bit more grey because AI outputs often appear quite polished. This can seem convincing and can easily take you off your guard, but just because something looks and sounds convincing does not mean it's true. Therefore, enhanced professional scepticism and verification back to source is required.

In my mind, thinking critically and questioning what we are presented with has become even more important in the world of AI.

Coming back to Isaac Asimov, after much exploration of the laws of robotics and how these can fail and be conflicted, he came up with another law – known as the Zeroth law – which was intended to protect humanity as a whole.

What would a Zeroth Law for responsible use of AI in transaction diligence look like?

This could be centred on transaction quality ('transaction quality comes above all else'). It could be centred around client trust ('client trust should always be maintained'). Or it could be focused on the capital markets ('to preserve trust in the capital markets'), with many more possibilities and iterations. Ultimately, we are probably still in the sandpit phase and have not accumulated enough practical experience to come up with a unifying Zeroth Law. Perhaps it will still come down to doing what's best for humanity as a whole.

Using this thought experiment to frame how we think about safeguards for the use of AI in our day-to-day work has been interesting (and fun). And as much as this was a thought experiment, it doesn't take away from the fact that we should be creating and implementing guardrails with respect to the use of AI, in parallel with ramping up usage and embedding AI in our day-to-day work. Perhaps these are guidelines or guiding principles, rather than laws. But in whichever way one may want to frame it, when it comes to transaction diligence, I quite like human professional judgement, confidentiality and traceability as guiding principles.

I firmly believe that people should remain in the driving seat. AI can certainly help us to work faster and/or be better at what we do, but at the end of the day, it remains just another tool in our toolbox.

My AI assistant helped me summarise some of this thinking around the Zeroth Law, and it came up with the following guideline when it comes to tension between efficiency and quality:

"Quality beats speed.

Evidence beats eloquence.

Judgement beats automation."

Perhaps this is true, at least for AI. And at least for now. 