

EDITORIAL

EDITORIAL

Milé čtenářky, milí čtenáři,

za rok a půl budeme slavit deset let od prvního vydání našeho/vašeho časopisu. Začínám proto poslední dobou více bilancovat, přemýšlím, co změnit, vylepšit... Učím se používat nové nástroje, zapojovat postupně do své práce umělou inteligenci. A možná právě proto bilancuji a přemýšlím o to víc. Kam a jak daleko ji pustit? Ověřovat, nebo důvěřovat? Jen si to představte: klíč s visačkou „lod“. Znamená to, že je to klíč od lodi? Nebo je to rafinovaný způsob majitele půjčovny lodí, jak obelstít případného zloděje klíčů? No páni, řeknete si, tohle uvažování za roh není zdravé, je více než pravděpodobné, že to je opravdu klíč od lodi! Jenže lze tento zdravý způsob uvažování, který je založen na elementární mezilidské důvěře, aplikovat na umělou inteligenci? Umělá inteligence nemá možná postranní úmysly, nesnaží se nás vědomě přelstít, jenže to nestačí. Nemá vědomí, nemá ŽÁDNÉ úmysly (doufejme), nemá ani pocit viny, když nás mystifikuje. Je velmi snadné důvěřovat tomu, co nám napíše v odpovědi na naši otázku, stejně jako je snadné důvěřovat tomu, že klíč s visačkou „lod“ je klíčem od lodi. Sama jsem v posledních týdnech několikrát přistihla umělou inteligenci při mystifikaci a trvalo mi dlouho přesvědčit ji, že se mýlí. Jednou to bylo tvrzení, že známá světová autorka nenapsala v posledních letech dvě knihy (umělá inteligence tvrdošijně trvala pouze na jedné), podruhé to bylo prohlášení – a to bylo obzvláště pikantní – že papež Lev XIV. neexistuje (bylo to asi týden po jeho zvolení). Přesvědčit ji o opaku bylo únavné. Podařilo se mi dokonce přimět ji k tomu, aby se mi omluvila. Radost jsem ale z tohoto „vítězství“ neměla. UVědomila jsem si, jak snadné je podlehnout jejím mystifikacím a nepřesným či zcela chybným informacím. Je velmi snadné zahodit kritické myšlení a plout na vlně pohodlného, instantního a bezbolestného přijetí.

A zde se dostávám k tomu, proč Vám to zde, milé čtenářky a čtenáři, píši. Náš časopis také přináší informace, celou řadu informací! A i přesto, že máme z jednotlivých článků v redakci radost a těšíme se, že obohatí naše i Vaše poznání, stále se jedná pouze o informace, s nimiž je třeba dál pracovat. Přejeme Vám, abyste nejen tento časopis, ale vše, co Vám je a bude v budoucnu předkládáno, četli s otevřeností a kritickou myslí. Je velmi pravděpodobné, že tam, kde je napsáno „lod“, lod také nalezneme. Ale nemusí tomu tak být vždy.

Šťastnou plavbu mořem informací Vám za redakci LKL přeje
Zuzana Lebedová

P. S.: Všímavý čtenář jistě zaznamenal drobnou změnu – chybí rozhovor s osobností oboru. Rozhovory definitivně přesouváme do podcastu Řeči o řeči, který je k poslechu na YouTube a Spotify a volné strany věnujeme recenzovaným článkům. O nic tedy nepřicházíte a věříme, že změnu pochopíte.



Dear readers,

In eighteen months from now it will be time to celebrate ten years since the first issue of y/our journal. Accordingly, I am lately finding myself taking stock, musing, wondering what to change, to improve... I am learning to use new tools, progressively introducing Artificial Intelligence into my workflow. And maybe that's why I am inclined to weigh-up and ponder all the more. Where should we let AI encroach, and how far? Do we keep checking on it, or trust it? Imagine this: A key, with "boat" written on the key fob. Does that mean it's a key to a boat? Or is it the boat-hire business owner's sneaky attempt to 'fob off' some sneak-key-thief? Oh come on,

you'll be thinking, this kind of second-guessing is no good for you, more than likely it is indeed a boat key! Yet, can this kind of guileless thinking, founded on elementary confidence in human nature, be applied to AI? Artificial Intelligence may indeed lack any guile, deviousness, and is not out to trick us on purpose, but that alone will not suffice. AI has no consciousness, it has NO intentions (hopefully), no sense of guilt about pulling the wool over our eyes, leading us astray. It's all too easy to trust what it writes in reply, just as it's easy to believe the key marked "boat" is in fact a boat key. Just in these last few weeks I have called out AI personally, as it tried to bamboozle me. It took quite some time and effort to persuade it that it was mistaken. One such case was when it claimed a renowned lady writer hadn't written two books in the last few years (AI stuck to claiming there had been only one), and the second case was when it proclaimed – a particularly zesty quirk on its part – that there was no Pope Leo XIV (about a week after his announcement). It was pretty tiresome to talk 'her' out of it. I did get her to apologize to me in the end. But I felt no joy as a result of my 'triumphant victory'. I realised how easily one could succumb to AI mystifications, inaccuracies and even downright falsehoods. It is too easy to suspend disbelief, put critical thinking aside and float along blithely on the gentle wave of convenient, immediate and painless acceptance.

And here is the crux of why I am writing this to you, dear readers. Our journal also brings information, plenty of it! And despite our editorial team being delighted by the articles and looking forward to how these articles will enrich y/our understanding, let us bear in mind that these are but pieces of information to work with, to follow further. We wish you the time and interest to read not only this journal but everything you come across, now and in the future, with an open, curious yet cautious mind. It is highly likely that where a "boat" tag is written, a boat will be found. But it is not to be taken for granted.

May you sail the high seas of information with all good fortune, such is our wish to you all, from all of us in the LKL editorial team.
Zuzana Lebedová

PS: The mindful reader will have noticed a small change – no interview with the 'Guiding Lights' of our domain. We are definitively relocating interviews to our 'Řeči o řeči' [Speaking of speaking] podcast on YouTube and Spotify; those pages are being freed-up for reviewed articles. Nothing has been lost, and we appreciate your kind indulgence.