



Delft University of Technology

Large Language Models: de toekomst van mobiliteit?

Gao, Ting; Movaghar, Mahsa; Hari, Theivaprakasham; Roocroft, Alex; Rinaldi, Marco; Xin, Yanan; Hoogendoorn, Serge

Publication date

2025

Document Version

Final published version

Published in

NM Magazine

Citation (APA)

Gao, T., Movaghar, M., Hari, T., Roocroft, A., Rinaldi, M., Xin, Y., & Hoogendoorn, S. (2025). Large Language Models: de toekomst van mobiliteit? *NM Magazine*, 2025(2), 31-33. <https://www.nm-magazine.nl/artikelen/large-language-models-de-toekomst-van-mobiliteit>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Large Language Models: de toekomst van mobiliteit?

Large Language Models zijn AI-systemen die menselijke taal begrijpen en zich er ook in kunnen uiten. Ze zijn de basis onder populaire applicaties als ChatGPT, Gemini en Copilot. Maar inmiddels is de technologie zó breed inzetbaar dat ze ook doordringt in de mobiliteitssector. Hoe werken de Large Language Models? Hoe kunnen ze van nut zijn in ons vakgebied? En wat zijn de mitsen en maren?

Wat is een LLM?

Het *Large Language Model*, LLM, is een vorm van kunstmatige intelligentie die is getraind met enorme hoeveelheden tekstdata. Tijdens die training herkent het model patronen die het opslaat. Als het maar met genoeg tekstdata wordt gevoed, leert het taal 'begrijpen' en kan het zelfs teksten genereren.

Dankzij LLM-modellen als het *Generative Pretrained Transformer*-model, ontwikkeld door OpenAI, kunnen we ook wat met die intelligentie. Een tekst samenvatten of vertalen, een mailtje of brief componeren of zelfs hele werkstukken schrijven – een beetje model draait er zijn hand niet voor om.

De techniek

Maar hoe gaat dat leren, begrijpen en genereren in zijn werk? We belichten een paar aspecten.

Het is allereerst belangrijk dat teksten 'verwerkbaar' zijn voor een computer. Dat betekent dat die woordenbrij van websites, boeken, social media en videotranscripts omgezet moet worden naar numerieke data. Dat gebeurt door aan elk woord een unieke, cijfermatige ID toe te kennen, een *token*. Een zin, alinea, hoofdstuk of boek leest voor een LLM daarmee als een lange rij cijfers.

Het model kan die numerieke data nu analyseren, naar patronen zoeken en relaties leggen. Het model checkt en update daarbij voortdurend zijn eigen kennisbasis en wordt zo ook hoe langer hoe intelligenter.

In dit proces speelt *self-attention* een belangrijke rol. Dit mechanisme helpt een LLM de relaties te bepalen tussen woorden (tokens) in een of meer zinnen. Oudere modellen verwerkten woorden nog één voor één – een beetje zoals een beginnend lezer met z'n vinger die woorden langsgaat. Maar met *self-attention* kijkt het model naar het geheel: van elk woord in een zin wordt eerst vastgesteld hoe belangrijk de andere woorden in de zin zijn voor z'n betekenis. Het resultaat van die verwerkingen wordt opgeslagen in een aantal vectoren. Elk token in een zin beschikt daarmee over informatie van zichzelf én de context. Met dat complete beeld verwerkt en interpreteert het LLM de zin. Ongeveer zoals een ervaren lezer dus: die leest ook niet 'lineair', maar scant voortdurend de context en focust daarbij op sleutelwoorden.

Een laatste aspect dat we willen aanstippen, betreft het genereren van teksten. Een LLM gaat daarbij in essentie te werk als de *autocomplete*-functie op een smartphone: op basis van een database kan die woorden suggereren die waarschijnlijk volgen op het voorgaande woord, zoals 'dag' na 'fijne'. Het verschil met de smartphonefunctie is dat een LLM is getraind op heel veel meer data, dankzij z'n analyses een goed beeld heeft van de context en slim variatie aanbrengt. Dat laatste gebeurt met *sampling*. Het LLM kiest dan niet steeds de woorden met de hoogste waarschijnlijkheid, maar overweegt meerdere kansrijke tokens. Zo ontstaan gevarieerdere, rijkere teksten.

Toepassingen in de mobiliteitssector

Tot zover de (vereenvoudigde) blik op de techniek. Nu de vraag wat LLM's voor ons werkveld mobiliteit kunnen betekenen. In het onderstaande noemen we een aantal mogelijke toepassingen. Die zijn niet uit de lucht gegrepen, maar zijn gebaseerd op lopende onderzoeken bij universiteiten en onderzoeksinstituten. Enkele kansen worden zelfs al in de praktijk gebruikt. Dat vermelden we dan specifiek.

Autonome voertuigen

LLM's kunnen op termijn in (autonome) voertuigen terechtkomen. Een taalmodel lijkt misschien ver af te staan van de rijtaak. Maar net als een tekst kan worden vertaald in tokens, kunnen ook de waarnemingen van voertuigsensoren in een reeks tokens worden omgezet. Het is modellen als GPT-Driver en DriveGPT-4 op deze wijze gelukt autorijden als een taalmodelleerprobleem te benaderen.

Wat de *perceptie* betreft zit de meerwaarde van LLM's in het combineren van data van LiDAR, camera's, sensoren enzovoort met contextueel begrip. Dat betekent bijvoorbeeld dat ze voetgangers en fietsers niet alleen detecteren, maar ook begrijpen: waarom mensen zich bij een oversteekplaats verzamelen of wat een fietser met een handgebaar bedoelt. Dankzij dit redeneervermogen kunnen ze beter uit de voeten met zeldzame of onverwachte situaties – iets waar conventionele deep *learning*-modellen nog veel moeite mee hebben.

In de *besluitvorming* verbeteren LLM's de interactie met de passagiers. Stel dat die de opdracht geven: "Voeg in naar links". Op basis van waarnemingsdata kan een LLM zo'n commando beoordelen op haalbaarheid én met een goede uitleg reageren als het commando niet



mogelijk of niet veilig is. Op dit moment heeft het Britse bedrijf Wayve al een LLM in zijn rijstelsysteem geïmplementeerd waaraan passagiers vragen kunnen stellen als: “Waarom rijden we niet?” Het systeem geeft dan in real-time antwoord.

Ook *motion planning* kan dankzij LLM een sprong vooruit maken. Traditionele ‘bewegingsplanning’ werkt vaak met rigide, op regels gebaseerde systemen. Die zijn effectief in gestructureerde omgevingen, maar minder geschikt voor complexe of onbekende situaties. LLM-modellen zoals de genoemde GPT-Driver benutten het algemene redeneervermogen van LLM’s om vloeiende, contextgevoelige paden te volgen in uiteenlopende scenario’s.

Logistiek

In de logistieke sector zijn LLM’s vooral interessant voor *route-optimalisatie*. Doordat ze excelleren in patroonherkenning, zijn ze prima in staat historische dienstregelingen, realtime GPS-data en de vervoersvraag te analyseren en voorstellen te doen voor betere routes en dienstroosters.

Ook het vinden van de efficiëntste *bezorgtrajecten* lijkt een kolfje naar LLM’s hand. Naar verluidt heeft Uber Freight al een LLM ingezet om optimale volgordes van vrachtwagenleveringen voor te stellen. Het model leerde van verkeersinformatie en eerdere leveringen, en wist zo tijd en brandstof te besparen.

Openbaar vervoer

LLM’s patroonherkenning komt ook van pas in het openbaar vervoer. Verschillende studies lieten al zien dat de modellen op basis van historische reizigersaantallen en vertragingen *buslijnen* en *treinschema’s* kunnen optimaliseren.

Maar ze zijn ook inzetbaar, en worden ook al zo gebruikt, als *digitale assistent voor reizigers*. Chatbots op basis van GPT-modellen helpen een toerist de beste metroverbinding te vinden en waarschuwen forenzen voor vertragingen op hun vaste route – allemaal via een natuurlijk gesprek. Nog een pre: deze AI-assistenten communiceren meertalig. Zo krijgen niet-moedertaalsprekers gelijke toegang tot mobiliteitsinformatie.

Beleid

Binnen het domein van verkeerswetgeving kunnen LLM’s functioneren als *beleidsassistenten* of zelfs *juridisch adviseurs*. Planners en ingenieurs hebben te maken met steeds complexere wetgeving – denk alleen maar aan de regels waaraan een autonoom voertuig moet voldoen. LLM’s kunnen hen hierbij ondersteunen, bijvoorbeeld door vragen

over mobiliteitsbeleid te beantwoorden, nieuwe verkeerswetten te interpreteren of lange ontwerpdocumenten samen te vatten. Het is onderzoekers aan de Universiteit van Californië-Los Angeles ook gelukt om met LLM’s *autonome voertuigen* wijs te maken in de wetgeving: door een LLM te voeden en trainen met een bibliotheek aan verkeersregels, kon het voertuigstelsel steeds zelf beredeneren welke wetten van toepassing waren en wat het voertuig dus wel en niet kon doen.

Veiligheid en verkeersmanagement

Ook op het gebied van veiligheid zouden LLM’s concrete aanbevelingen kunnen doen in begrijpelijke taal. Onderzoekers aan de Universiteit van de Chinese Academie van Wetenschappen in Beijing, UCAS, hebben in 2024 het model AccidentGPT (in concept) ontwikkeld. Dit LLM fungeert als een soort *veiligheidsexpert* die multimediale sensordata omzet in voorspellingen over ongevallen, waarschuwingen genereert en via dialoog preventieve maatregelen voorstelt.

Een andere mogelijkheid is om met behulp van LLM’s *ongevallenrapporten* te analyseren op risicofactoren. Ze kunnen onderbelichte oorzaken identificeren, zoals alcoholgebruik, vermoeidheid of afleiding tijdens het rijden.

Wat verkeersmanagement betreft zijn LLM’s ook waardevol voor het analyseren en voorspellen van *verkeersdynamiek* op macroniveau. Denk aan het ontrafelen van complexe afhankelijkheden binnen uiteenlopende datasets om daarmee verbanden te modelleren tussen bijvoorbeeld infrastructuur, weersomstandigheden en gebruikersgedrag. De te verwachten winsten zijn betere verkeersvoorspellingen, geoptimaliseerd infrastructuurbeheer en meer inzicht in hoe diverse factoren verkeersdynamiek beïnvloeden.

Uitdagingen van LLM’s

Bovenstaand overzicht is kort en onvolledig. De voorbeelden laten hoe dan ook goed zien, wat voor *game changer* LLM in de mobiliteit kan zijn. Bij die vaststelling hoort wel een disclaimer: de route van onderzoek naar ‘toepassing op straat’ is geen gemakkelijke. Eerst moeten we nog enkele (forse) technische en ethische hordes nemen.

Technische uitdagingen

Wat techniek betreft is er om te beginnen de horde van *training*. Een generiek taalmodel moet worden afgestemd op zeer specifieke



Foto: Mikki Orso

data – dienstregelingen, sensordata, beleidsdocumenten – zodat het de ‘vervoerstal’ spreekt. Daarbij mag het zijn algemene kennis natuurlijk niet verliezen. Dit vereist grote gelabelde datasets, duizenden GPU-uren en gespecialiseerde medewerkers. Die horde is voor kleinere organisaties vaak al te groot.

Ten tweede: *integratie*. Bestaande ITS-systemen, die vaak over decennia zijn opgebouwd, moeten we voorzien van nieuwe interfaces zodat ze LLM live-data kunnen lezen en bruikbare instructies terug kunnen geven. Dat is al een operatie op zich. Maar dan is er nog de uitdaging ‘hallucinaties’ te voorkomen: antwoorden die plausibel klinken maar onjuist zijn. Bij ChatGPT komen die nog in circa 3 procent van de antwoorden voor. Bij sommige concurrenten ligt het hallucinatiepercentage zelfs op 27 procent!

Die hallucinaties zouden al een probleem zijn bij bijvoorbeeld beleids-ondersteuning, maar wat te denken van een LLM in een autonoom voertuig? Daar moet de kans op hallucinaties echt tot nul worden gereduceerd. Dat kan wellicht met hybride AI-benaderingen. Bij het Amerikaanse bedrijf Nuro bijvoorbeeld draaien regelgebaseerde en AI-gebaseerde planners parallel, een interessante combinatie van de uitlegbaarheid van regels en de aanpasbaarheid van taalmodellen. Tegelijkertijd geldt dat in autonome voertuigen ook snelheid een vereiste is. Beslissingen moeten vaak binnen milliseconden genomen worden. Dat kan met de huidige generatie LLM’s nog niet betrouwbaar genoeg, laat staan als de uitkomsten gecombineerd moeten worden met die van regelgebaseerde planners.

Een derde horde is de *ict-infrastructuur*. Duizenden reizigersvragen of routeringsverzoeken per minuut kunnen zorgen voor hoge cloudkosten of vertragingen van enkele seconden. Dit dwingt ontwerpers soms tot zogeheten *edge deployments*: lichtere modellen die lokaal draaien. Die zijn goedkoper en sneller, maar ook veel minder krachtig.

Een laatste punt: *energieverbruik*. Alleen al het trainen van GPT-3 kostte zo’n 1287 MWh aan energie, goed voor een uitstoot van 552 ton CO₂. En elke live-query houdt stroom slurpende GPU’s actief – net op een moment dat overheden en vervoersinstanties hun emissies proberen te verlagen.

Ethische uitdagingen

Naast de technische hordes zijn er ook nog de ethische vraagstukken rond AI en LLM’s. Zo mogen we in ons enthousiasme niet vergeten dat LLM’s wel intelligent zijn, maar geen werkelijk begrip van de fysieke realiteit hebben. Ze blinken uit in redeneren en reageren op

natuurlijke taal, maar werken daarbij met benaderingen van betekenis, opgebouwd uit taalrelaties, niet vanuit echte wereldkennis. *Menselijke autonomie* is daarom ook bij de inzet van LLM’s een basisvoorwaarde. Het is aan ingenieurs en modelbouwers om ervoor te zorgen dat de modellen voldoende afgestemd zijn op de werkelijkheid – maar de mens is en blijft verantwoordelijk.

Denk ook aan *veiligheid*. We noemden al het probleem van hallucineren. Maar slechte trainingen en verkeerde implementaties kunnen net zo goed tot problemen leiden: ‘garbage in garbage out’.

Andere ethische kwesties zijn *rechtvaardigheid* en *privacy*. LLM’s kunnen onbedoeld bestaande ongelijkheden versterken als ze getraind zijn op datasets waarin die vooroordelen verstopt zitten. Datasets kunnen bovendien gevoelige vervoersdata bevatten.

Al deze technische en ethische punten moeten aangepakt en waterdicht opgelost worden voor er sprake kan zijn van ‘toepassingen op straat’.

Conclusie

We leven in het tijdperk van AI, waarin de innovaties elkaar in razend tempo opvolgen. De opkomst van *Large Language Models* verandert momenteel hele industrieën.

Gelet op technische implementatiehobbels en de ethische uitdagingen kunnen we echter stellen dat we ons nog in de beginfase van de ontwikkeling bevinden. LLM’s gaan ook ons vakgebied mobiliteit beïnvloeden en versterken – die uitspraak durven we wel aan. Maar menselijk oordeel, kritisch denken en domeinexpertise blijven onmisbaar voor hun veilige en verantwoorde inzet. ●

De auteurs

Ting Gao, Mahsa Movaghar, Theivaprakasham Hari en Alex Roocroft zijn onderzoekers van het DAIMOND Lab van TU Delft. Zij worden begeleid door dr. ir. Marco Rinaldi en dr. Yanan Xin, co-directeurs van het Lab, en prof. dr. ir. Serge Hoogendoorn, adviseur van het Lab.