

Het schrijfstijl-plan: mails die klinken als de medewerker zelf

Dit document beantwoordt de moeilijkste van jullie vier vragen: hoe zorg je dat conceptmails herkenbaar klinken als de medewerker zelf, zonder AI-gladheid? We hebben hiervoor diepgaand onderzocht hoe de toonaangevende producten (Shortwave, Grammarly, Superhuman) dit oplossen én wat recent wetenschappelijk onderzoek (2024-2026) erover zegt. De uitkomsten zijn verrassend eenduidig — en goed nieuws: de beste aanpak is precies degene die een klein team zelf kan bouwen.

Waarom dit het moeilijkste onderdeel is

Persoonlijke schrijfstijl is impliciet. Je kunt hem bijna niet in regels vangen ("schrijf kort en direct") — en dat is meetbaar bewezen: een AI die alleen een stijlbeschrijving krijgt, wordt in maar 26 procent van de gevallen herkend als de bedoelde schrijver. Het resultaat is dan precies de gladde AI-mail die jullie in de instructie verbieden. Dezelfde AI met vijf echte voorbeeldmails erbij: 69 procent herkenning. Dit verklaart ook waarom de schrijfassistent in onze demo nog niet overtuigt — die draait op drie verzonnen voorbeeldmails, en verzonnen voorbeelden zijn precies wat niet werkt.

Wat de markt doet (en wat níet)

De drie toonaangevende e-mailproducten gebruiken alle drie dezelfde kern, en geen van hen traint een apart AI-model per gebruiker:

Shortwave bouwt per gebruiker een vooraf berekend, tekstueel stijlprofiel én stopt daarnaast relevante echte voorbeeldmails in de prompt. Ze hebben model-training per gebruiker geprobeerd en verworpen als onwerkbaar. Grammarly destilleert automatisch een "voice profile" uit je schrijfgedrag, werkt dat continu bij, en houdt aparte profielen bij per tekstsoort (berichten versus documenten). Superhuman geeft de gebruiker een zelf aanpasbaar personalisatie-profiel (aanhef, afsluiting, over-mij) en adviseert expliciet om voorbeeldmails mee te sturen.

Model-training per medewerker (fine-tuning) haalt in het lab de hoogste stijlgetrouwheid, maar vereist honderden tot duizenden mails per persoon, maakt het model slechter in het opvolgen van instructies, en is voor Hamza een beheerlast zonder einde. Bij tien tot vijftig mails per medewerker — jullie realiteit — is het aantoonbaar de verkeerde route.

Het recept voor GreenFox, in vijf stappen

Stap 1 — verzamel per medewerker 10 tot 15 echte verzonden mails. Onderzoek op zakelijke e-mail laat zien dat circa vijf goed gekozen voorbeelden al 95 procent stijlherkenning geeft; meer dan tien voegt weinig toe. Belangrijk: laat de medewerker de mails zélf kiezen ("waar ben je tevreden over?") en splits in twee stapels: mails aan opdrachtgevers en mails aan kandidaten — stijl verschilt per publiek, dus elk krijgt zijn eigen setje. Deze zelf-selectie is meteen jullie filter tegen het meekopiëren van slechte gewoontes.

Stap 2 — destilleer per medewerker een stijlprofiel, en laat het corrigeren. De AI vat de voorbeeldmails eenmalig samen in een leesbaar profiel: vaste aanhef en afsluiting, gemiddelde zinslengte, terugkerende frasen, wat iemand juist nooit doet. De medewerker leest dat profiel en streept door wat niet klopt. Dat is het Superhuman-model, en het lost drie dingen tegelijk op: het profiel wordt beter, de medewerker houdt regie (belangrijk tegen het gevoel van controle dat in jullie organisatie gevoelig ligt), en het is transparant richting AVG.

Stap 3 — bouw elke conceptmail uit drie lagen. Laag één: de vaste GreenFox-huisstijl als systeeminstructie (geen streepjes, niets dikgedrukt, geen emoji's, snel ter zake, concrete vervolgstap). Dit soort expliciete regels werkt wél als instructie — het is de impliciete stijl die voorbeelden nodig heeft. Laag twee: het gecorrigeerde stijlprofiel van de medewerker. Laag drie: de drie tot vijf meest relevante échte mails uit diens setje, automatisch gekozen op gelijkenis met de huidige situatie (een afwijzing haalt eerdere afwijzingen op, een offerte-opvolging eerdere offertes). Onderzoek naar gepersonaliseerde e-mailgeneratie vond dat ophalen van circa vier relevante eerdere teksten optimaal is.

Stap 4 — meet "klinkt als mij" op de juiste manier. Niet door het aan de AI zelf te vragen (aantoonbaar onbetrouwbaar), maar simpel en menselijk: leg de medewerker en een collega blind twee mails voor — één echte oude, één AI-concept — en kijk of ze het verschil zien. Wil Hamza het technisch: er bestaan stijl-vergelijkingsmodellen die dit automatisch scoren, als

aanvulling. Spreek vooraf af wat goed genoeg is, bijvoorbeeld: de collega gokt niet beter dan kans.

Stap 5 — onderhoud zonder inteelt. Ververs het setje voorbeeldmails elk kwartaal met recente, door de medewerker zelf geschreven mails. Voer goedgekeurde AI-concepten nooit terug als nieuwe voorbeelden — dan traint de assistent op zichzelf en glijdt de stijl langzaam terug naar generiek.

Eerlijke kanttekeningen

Vrijwel al het onderzoek is gedaan op Engelstalige zakelijke e-mail; voor Nederlands zijn de cijfers aannemelijk maar onbewezen — precies daarom is de blind-test uit stap 4 geen luxe maar noodzaak. En verwacht geen perfecte kloon: meetbaar blijft AI-tekst iets verder van de echte schrijver af dan diens eigen mails. Het doel is dan ook niet onzichtbaarheid, maar een concept waar de medewerker in één keer "ja, dit ben ik" bij voelt en hooguit een woord in aanpast. Alles blijft bovendien concept — versturen doet altijd een mens, zoals jullie zelf als harde eis stelden.

Wat dit betekent voor de demo en voor fase 1

De schrijfassistent-demo op deze site laat het twee-lagen-principe zien met verzonnen medewerkers — het bewijs van het mechanisme, niet van de kwaliteit. De kwaliteit komt uit jullie echte mails, en die stap is klein: tien mails aanleveren per pilot-medewerker is een half uur werk. Concreet voorstel: neem dit recept mee in fase 1 voor de eerste twee gebruikers, en maak van stap 4 (de blind-test) het meetpunt dat bepaalt of de schrijfassistent door mag naar de rest van het team.

Onderliggende bronnen

Producten: Shortwave engineering-blog over hun AI-assistent (shortwave.com/blog/deep-dive-into-worlds-smartest-email-ai), Grammarly voice-features documentatie (support.grammarly.com), Superhuman Write-with-AI documentatie (help.superhuman.com).

Onderzoek: stijllimitatie-benchmark op zakelijke e-mail, EMNLP 2025

(arxiv.org/abs/2509.14543 — 5 voorbeelden: 95-96 procent herkenning; kale instructies falen), gepersonaliseerde lange-tekstgeneratie LongLaMP (arxiv.org/abs/2407.11016 — retrieval van 4 eerdere teksten optimaal), vergelijking instructie/few-shot/fine-tuning (arxiv.org/abs/2409.04574 — 26/69/88 procent), evaluatie van stijlgelijkenis (arxiv.org/abs/2508.06374 en arxiv.org/pdf/2510.13898 — combineer meetmethodes, vertrouw niet op één AI-beoordelaar).