

Tekoälypohjainen Delfoi-analyysi tekoälyagenteista alkuvaiheen ohjelmistostartupeissa

Tekoälyagenteilla viitataan autonomisiin ohjelmistojärjestelmiin, jotka hyödyntävät tekoälyä suorittaakseen käyttäjän asettamat tavoitteet. Näille agenteille on ominaista kyvykkyys päättelyyn, muistiin ja suunnitteluun. (Gutowska & Stryker 2025) Käytännössä tekoälyagentit voivat esimerkiksi pilkkoa monimutkaisia tehtäviä pienempiin osiin, suunnitella parhaan mahdollisen suoritusjärjestyksen ja toteuttaa ne käyttäen suuria kielimalleja sekä mahdollisesti työkaluja.

Alkuvaiheen ohjelmistostartupit, eli ohjelmistokehitykseen keskittyvät, nuoret ja kasvuhakuiset yritykset, joiden tiimeissä yksittäiset kehittäjät yleensä kantavat monta hattua. Koodauksen ja teknologian ylläpidon lisäksi he saattavat hoitaa asiakaspalvelua, analytiikkaa, infrastruktuuria tai monitorointia. (Krill 24.2.2025) Rutiinin omaiset ohjelmistokehitystehtävät, kuten testien kirjoittaminen, dokumentointi tai kirjastojen päivitys vievät suurimman osan työajasta. Kuitenkin kehittäjän arvokkain panos kohdistuisi mieluummin suunnitteluun, uusiin ominaisuuksiin ja suuriin arkkitehtuuripäätöksiin.

Tämä ennakoiva analyysi keskittyy selvittämään, minkälaisia riskejä ja pullonkauloja on itsenäisten tekoälyagenttien käyttöönotossa alkuvaiheen ohjelmistostartupeissa tulevana vuosina. Siinä yritetään myös ymmärtää millä aikavälillä nämä teknologiat voivat kypsyä ja mitä haasteita on ratkaistava, jotta ne integroituvat sujuvasti startup-ympäristöön. Tämä on ajankohtainen aihe, koska ohjelmistokehityksen työkalujen muutos kohti generatiivisella tekoälyllä tuettuja työkaluja on jo nähtävissä. Esimerkiksi Github Copilot:in on osoitettu nopeuttavan koodaustehtäviä jopa 55%. (Peng & Kalliamvakou & Cihon & Demirer 2023)

Delfoi asiantuntijapaneeli

Ennakointianalysissa valittuna menetelmänä on Delfoi-menetelmä, joka on iteratiivinen asiantuntijapaneeliin perustuva tekniikka tulevaisuuden näkymien kartoittamiseen. Delfoi on alun perin kehitetty luomaan asiantuntijoiden konsensus vaikeasti ennustettavissa tai kiistanalaisissa aiheissa. Menetelmässä joukko asiantuntijoita antaa anonymisti arvionsa useassa kyselykierroksessa. Jokaisen kierroksen jälkeen vastaukset kootaan ja anonymisoitu kooste palautetaan paneelille seuraavaa kierrosta varten. Tavoitteena on, että mielipiteet lähestyvät toisiaan kierros kierrokselta paneelin jäsenten tarkastellessaan uudelleen omia näkemyksiään muiden osallistujien palautteen valossa, kunnes saavutetaan hyväksyttävän taso yksimielisyyttä tai tunnistetaan pysyvät erimielisyydet. Delfoi soveltuu hyvin tilanteisiin, joissa kerättyä tietoa on niukasti tai se on epävarmaa. Tämän takia Delfoi soveltuu hyvin tekoälyagenttien tarkasteluun. Kyseessä on nouseva teknologia, jonka vaikutuksia pienten startuppien kehitystyöhön ei voida suoraan ekstrapoloida historiasta.

Valitsin Delfoin, koska se mahdollistaa useiden eri näkökulmien systemaattisen yhdistämisen. Keräämällä monialainen paneeli, saadaan kattava kuva teknologian mahdollisista kehityskuluista liiketoiminnan, teknologian, sääntelyn ja käyttäjän näkökulmista. Vaihtoehtoiset ennakointimenetelmät kuten skenaariot tai teknologian tiekartat olisivat olleet toki mahdollisia, mutta Delfoin etuna on yhteisvaikutusten ja näkemyksien konsolidointi. Se pakottaa asiantuntijatkin miettimään toistensa

argumentteja ja paljastaa konsensuskohdat ja erimielisyydet selkeästi. Delfoi sopii tähän teknologia/kontekstipariin, koska halutaan ymmärtää sekä optimismia että huolia tekoälyagenttien suhteen. Menetelmä auttaa löytämään, mitkä tulevaisuusnäkymät nähdään todennäköisinä ja mitkä jäävät kiistanalaiseksi.

Tekoälyavusteisen Delfoi-paneelin toteutus

Tekoälyavusteista Delfoi-paneelia on tutkittu jo aikaisemmin kirjallisuudessa, esimerkiksi Berloitti & Mari 2025 toteuttivat tekoälyavusteisen Delfoi-tutkimuksen ja heidän mukaansa suuret kielimallit pystyvät tuottamaan monipuolisia skenaarioita ja vähentämään vastaajaväsymystä. Samalla kuitenkin on tiedostettava riskit, kuten mallien "knowledge-cut off" eli päivä, johon asti suurella kielimallilla on koulutusdatassa tietoa voi olla vuosia vanha sekä malleissa saattaa olla vinoumia. (Berlotti & Mari 2025)

Tässä analyysissä Delfoi-paneeli toteutettiin tekoälyavusteisesti. Paneelin jäseninä toimivat OpenAI:n GPT-5-mini suuret kielimallit, joille annettiin kunkin rooliin sopivat taustatiedot ja näkökulmat. Yksitellen verkkokäyttöliittymässä promptaamisen sijaan, paneeli on toteutettu komentorivikäyttöisenä sovelluksena, joka kommunikoi OpenAI:n rajapinnan kanssa. Kaikki mallin tuottamat vastaukset validoidaan Zod-skeemojen avulla, mikä varmistaa, että paneelistien roolit, kysymykset, kierrosvastaukset ja konvergenssianalyysit noudattavat ennalta määritettyä tietomallia. Virheelliset tai skeemasta poikkeavat vastaukset hylätään automaattisesti. Linkki ajettavaan koodiin, sekä siitä generoitu raportti on saatavilla työn lopussa liitteet kohdasta.

Ohjelman käynnistyessä käyttäjä määrittelee ensin panelistien roolit vapaamuotoisella tekstillä. Nämä roolit hiotaan gpt-5-mini mallia käyttäen promptilla "*You polish Delphi panel roles. Return panelists matching the order of the input list.*". Roolien hiomisen syynä on mahdollisuus lieventää Berloitti & Mari:n löytämää ongelmaa, jossa huomattiin tekoälyavusteisen Delfoi-paneelin reagoivan herkästi ensimmäisille syötteille. Delfoi-analyysin toistettavuuden ja rakenteellisen yhdenmukaisuuden varmistamiseksi siis käyttäjän tuottamat roolit, jotka ovat väistämättä kielellisesti vaihtelevia, yhtenäistetään generatiivisella tekoälyllä.

Seuraavaksi käyttäjä syöttää viisi Delfoi-kysymystä, jotka myös jalostetaan samaa kielimallia käyttäen ennen paneelin ajoa. Tätä varten annettu promptaus on "*You refine Delphi panel questions. Keep the intent but improve clarity and make them concise and pointed.*". Kysymysten jalostamisen tarkoituksena ei ole muuttaa niiden sisältöä, vaan poistaa mahdollinen monitulkintaisuus, terävöittää haluttua kysymystä ja varmistaa että kaikki panelistit vastaavat täsmälleen samaan ongelmanasetteluun. Delfoi-menetelmä perustuu siihen, että vastauksia voidaan vertailla kierrosten välillä. Jos kysymykset jäävät kielellisesti epätarkoiksi, vastausten muutokset voivat johtua tulkintaeroista tai kysymyksen latautuneisuudesta. Esimerkiksi kielteisesti latautunut kysymys voisi olla "Miksi tekoälyagentteja ei tulla ikinä näkemään tässä kontekstissa", jonka kyseinen promptattu kielimalli muutti testissä muotoon "Mitkä pitkän aikavälin riskit voivat pysäyttää tekoälyagenttien hyödyntämisen varhaisen vaiheen startup-ympäristössä?"

Varsinainen Delfoi-paneeli suoritetaan kierroksittain. Ensimmäisellä kierroksella jokaiselle panelistille tuotetaan itsenäinen vastaus kysymyksiin ilman tietoa muiden panelistien vastauksista. Tätä varten käytetty promptaus on "*You simulate an*

independent-response Delphi round. For each panelist, craft a concise answer (max 80 words) true to their focus". Mallille annetaan myös panelistin kuvaus sekä käsiteltävä kysymys JSON muodossa. Malli palauttaa jokaiselle panelistille erillisen vastauksen.

Toisesta kierroksesta alkaen panelistit ottavat huomioon edellisen kierroksen yhteenvedon, joka annetaan mallille kontekstina promptin lisäksi. Näillä kierroksilla jokaiselle panelistille annettu prompt muuttuu muotoon ”*You simulate a Delphi round where panelists have seen previous responses. Each panelist should consider the prior summary but maintain their perspective.*”. Tämä pyrkii vastaamaan klassisen ihmispaneelin ideaa iteratiivisesta mielipiteiden lähestymisestä. Kierrosten jälkeen ohjelma suorittaa lisäksi konvergenssianalyysin, jossa tarkastellaan asiantuntijavastausten lähentymistä tai pysyviä erimielisyyksiä. Tätä varten käytetään promptia: ”*You analyze Delphi panel convergence. Assess how panelists are moving towards consensus, identify agreements and disagreements, and describe how answers have evolved.*” Kaikki ajon aikana syntyvät vastausdatat tallennetaan sekä JSON- että Markdown-muodossa reports-hakemistoon, mikä tukee läpinäkyvyyttä ja jäljitettävyyttä. Ohjelma ei käytä piiloprompteja tai jälkikäteen tehtävää sisällöllistä manipulointia, vaan kaikki tekoälyn toiminta perustuu eksplisiittisesti määriteltyihin promptteihin.

Toteutettu paneeli koostui seitsemästä jäsenestä, joille määrättiin seuraavat roolit:

- Alan tutkija, tutkii aihetta tieteellisesti; tarjoaa tutkimusnäyttöön perustuvia arvioita, metodologista kritiikkiä ja todennäköisyyksiin liittyvää perspektiiviä.
- Startupin perustaja ja toimitusjohtaja, käytännönläheinen näkökulma innovaatioiden kaupallistamisesta, nopeasta iteroinnista, resurssirajoista ja markkinariskeistä.
- Suuryrityksen edustaja, organisaation strateginen näkökulma: skaalaaminen, integraatio olemassa oleviin prosesseihin, sääntely- ja compliance-vaikutukset
- Viranomaisen edustaja, julkisen sektorin perspektiivi: sääntely, julkinen turvallisuus, politiikka ja yhteiskunnalliset seuraukset.
- Sijoittaja, rahoitus- ja markkinanäkökulma: riskinarviointi, tuotto-odotukset, aikahorisontit ja rahoituspäätösten vaikutukset.
- Kehittäjä, tekninen toteutettavuus: arkkitehtuuri, kehityskustannukset, ylläpito ja käytännön tekniset haasteet.
- Eettinen asiantuntija, eettiset ja yhteiskunnalliset näkökulmat: yksilönsuojan, oikeudenmukaisuuden, vastuullisuuden ja haittavaikutusten arviointi.

Jokaisessa roolissa pyrittiin korostamaan hieman eri näkökulmaa, sekä valitsemaan rooleja jotka voitaisiin olettaa olevan puolesta, vastaan sekä neutraaleja. Paneeli toteutettiin anonyymisti, eli panelistit eivät saaneet tietoa toistensa rooleista, mikä on Delfoin yksi perus periaatteista. Tämä mahdollistaa sen ettei yksittäiset vahvat persoonat dominoi keskustelua (Chuenjitwongsa S, 2017). Ajanjaksoksi valittiin 2 vuotta ja 5 vuotta, jotta panelistit pystyivät hahmottamaan sekä välittömiä trendejä että hieman pidemmän aikavälin mahdollisuuksia.

Paneelille esitettiin ensimmäisellä kierroksella 5 kysymystä, jotka oltiin hiottu tekoälymallin avulla.

1. Mitkä teknologiset riskit voivat estää tekoälyagenttien laajan käyttöönoton alkuvaiheen ohjelmistostartupeissa seuraavan viiden vuoden aikana?
2. Missä näet merkittävimmät liiketoiminnalliset pullonkaulat, kun startup pyrkii ottamaan tekoälyagentteja osaksi ydintoimintaansa?
3. Mitkä operatiiviset riskit liittyvät tekoälyagenttien käyttöön ohjelmistokehityksen keskeisissä tehtävissä?
4. Mitkä sääntelyyn, lainsäädäntöön tai eettisyyteen liittyvät riskit voivat hidastaa tekoälyagenttien käyttöönottoa aikaisen vaiheen startupeissa?
5. Missä tilanteissa tekoälyagenttien käyttöönotto voi aiheuttaa enemmän haittaa kuin hyötyä alkuvaiheen startupille?

Neljä ensimmäistä kysymystä on pyritty muodostamaan niin, että niillä voitaisiin tarkastella tekoälyagenttien käyttöönottoon liittyvän epävarmuuden neljää keskeistä riskiluokkaa, teknologista, liiketoiminnallista, operatiivista, sekä sääntelyä ja eettisyyttä. Viides kysymys on tarkoituksella kokonaisarvioiva. Sen avulla pyritään rikkomaan teknologiahypen oletus siitä, että tekoälyagenttien käyttöönotto olisi lähtökohtaisesti aina positiivinen kehityssuunta. Kysymys pakottaa panelin arvioimaan tilanteita, joissa agentit voivat heikentää startupin toimintakykyä esimerkiksi lisäämällä teknistä velkaa, heikentämällä tiimin oppimista, kasvattamalla operatiivista kompleksisuutta tai luomalla näennäistä automaatiota ilman todellista tuottavuushyötyä.

Tekoälyavusteisen Delfoi-paneelin analyysi

Ensimmäinen kierros keskittyi laajan riskikartan muodostamiseen. Kaikki panelistit vastasivat edellä mainittuihin viiteen kysymykseen. Panelistien vastauksissa korostui hyvin erilaisia tulkintoja siitä, mikä osa-alue muodostaa kaikkein kriittisimmän esteen tekoälyagenttien käyttöönotolle. Tutkija ja kehittäjä painottivat mallien epäluotettavuutta, startup-johtaja ja sijoittaja nostivat esiin kustannusrakenteen, viranomainen ja eettinen asiantuntija keskittyivät erityisesti läpinäkyvyyden, vastuun ja tietosuojan haasteisiin. Ensimmäisen kierroksen lopputuloksena muodostui ristiriitainen mutta monipuolinen kuva siitä, mitä haasteita startup kohtaa ennen kuin tekoälyagentteja voidaan käyttää tuotantokriittisissä osa-alueissa.

Toisella kierroksella näkemykset alkoivat selvästi lähentyä toisiaan. Panelistit eivät enää keskittyneet yksittäisten riskien listaamiseen, vaan tarkensivat syy ja seuraussuhteita. Esimerkiksi teknologisia riskejä koskevissa vastauksissa esiintyi yhteinen linja. Esimerkiksi validoinnin puute, ja tietoturva oli jaettuna huolenaiheena. Liiketoiminnallisessa kysymyksessä korostui yleinen yksimielisyys siitä, että tekoälyagenttien käyttöönotto on perusteltua vain rajatuissa piloteissa, jossa on mitattavissa olevia tuloksia.

Seuraavilla kierroksilla loppuun asti panelistien vastaukset saapuivat koko ajan lähemmäksi toisiaan. Loppujen lopuksi panelistit olivat samaa mieltä, että suurimmat ongelmat tekoälyagenttien käyttöönotossa ovat kustannukset, varsinkin jos käyttöönotto vie resursseja pois ydintoiminnoista, sekä sääntely. Panelistit ehdottivat käytännön toimenpiteenä vaiheistettuja ja kapeita pilotteja, joiden tulokset ovat mitattavissa. Panelistit olivat kuitenkin eri mieltä esimerkiksi priorisoinnista sekä sallitun riskitason määrittely tuotannossa.

Lähteet

Gutowska, A. Stryker, C. 2025. The 2025 Guide to AI Agents. IBM. Luettavissa: <https://www.ibm.com/think/ai-agents>. Luettu: 9.12.2025

Krill, P. 24.2.2025. Developers spend most of their time not coding – IDC report. Infoworld. Luettavissa: <https://www.infoworld.com/article/3831759/developers-spend-most-of-their-time-not-coding-idc-report.html>. Luettu: 9.12.2025

Peng, S. Kalliamvakou, E. Cihon, P. Demirer, M. 2023. The Impact of AI on Developer Productivity: Evidence from GitHub Copilot. Microsoft Research. St, Redmond. Luettavissa: <https://doi.org/10.48550/arXiv.2302.06590>. Luettu: 9.12.2025

Berlotti, F. Mari, L. 2025. An LLM-based Delphi Study to Predict GenAI Evolution. LIUC - Università Cattaneo Castellanza. Luettavissa: <https://doi.org/10.48550/arXiv.2502.21092>. Luettu: 10.12.2025

Chuenjitwongsa S, 2017. How to Conduct a Delphi Study. Cardiff University. Cardiff. Luettavissa: https://www.cardiff.ac.uk/_data/assets/pdf_file/0010/1164961/how_to_conduct_a_delphi_study.pdf. Luettu: 10.12.2025

Liitteet

AI-Delphi Raportti. Joni Juntto. <https://github.com/JoniJuntto/AI-Delphi/blob/master/reports/delphi-report-2025-12-09T18-48-57-226Z.md>.

AI-Delphi lähdekoodi. Joni Juntto. <https://github.com/JoniJuntto/AI-Delphi/tree/master>.