



Tekoäly haastaa korkeakoulut pohtimaan akateemista rehellisyyttä uudella tavalla

16.9.2026

Generatiivinen tekoäly on muutamassa vuodessa muuttanut perusteellisesti tapaa, jolla opiskelijat opiskelevat, kirjoittavat ja hakevat tietoa. Samalla se on tuonut korkeakouluille uudenlaisen haasteen: miten varmistaa akateeminen rehellisyys tilanteessa, jossa tekoäly kykenee tuottamaan laadukasta tekstiä, ratkaisemaan tehtäviä ja osallistumaan ongelmanratkaisuun lähes ihmisen tavoin? Viime vuosien tutkimukset osoittavat, että generatiivinen tekoäly haastaa perinteisiä käsityksiä oppimisesta, arvioinnista ja tekijyydestä (Bittle & ElGayar, 2025; Kasneci ym., 2023).

Tekoäly on tullut osaksi opiskelijoiden arkea. Sitä käytetään esimerkiksi ideointiin, tiedonhakuun, tekstien muokkaamiseen, kielentarkistukseen ja vaikeiden aiheiden ymmärtämiseen. Oikein käytettynä tekoäly voi toimia tehokkaana oppimisen tukena. Samalla se kuitenkin haastaa perinteiset käsitykset siitä, miten osaamista osoitetaan ja kuka on työn todellinen tekijä (Kasneci ym., 2023; Oravec, 2023).

Vilppi ei ole enää vain plagiointia

Perinteisesti akateemisella vilpillä on tarkoitettu esimerkiksi plagiointia eli toisen tekstin esittämistä omana työnä. Generatiivinen tekoäly muuttaa tilannetta merkittävästi. Opiskelija voi pyytää tekoälyä kirjoittamaan kokonaisen vastauksen, muokkaamaan olemassa olevan tekstin uudelleen tai ratkaisemaan tehtävän opiskelijan puolesta. Tällöin lopputulos voi olla teknisesti alkuperäinen, mutta sen tekijyys jää epäselväksi. Tutkimusten mukaan tekoäly onkin laajentanut akateemisen vilpin muotoja perinteisestä plagioinnista kohti

uudenlaisia tekijyyteen ja vastuuseen liittyviä kysymyksiä (Gray ym., 2025; Oravec, 2023).

Tutkimuksessa havaittiin, että tekoälyyn liittyvät vilppiepäilyt kohdistuvat erityisesti kahteen alueeseen: kirjallisiin oppimistehtäviin ja digitaalisiin tentteihin. Kirjallisissa tehtävissä epäilyksiä herättävät esimerkiksi opiskelijan aikaisemmasta tuotannosta poikkeava kirjoitustyyli, rakenne tai tekoälyn tuottamat virheelliset lähdeviitteet. Digitaalisissa tenteissä puolestaan haasteena ovat luvattomat laitteet ja tekoälyn hyödyntäminen arviointitilanteessa.

Korkeakoulujen näkökulmasta kysymys ei ole kuitenkaan vain vilpin tunnistamisesta. Yhä useammin joudutaan pohtimaan, missä kulkee hyväksyttävän ja ei-hyväksyttävän tekoälyn käytön raja. Jos opiskelija käyttää tekoälyä tekstin kieliasun parantamiseen, onko kyse vilpistä vai normaalista työkalun käytöstä? Entä jos tekoäly osallistuu ideointiin tai lähteiden jäsentelyyn? Nämä kysymykset liittyvät laajempaan keskusteluun akateemisesta integriteetistä ja vastuullisesta tekoälyn käytöstä (Eaton, 2019; Bretag, 2013).

Selkeät pelisäännöt tukevat vastuullista käyttöä

Tutkimuksen mukaan merkittävä osa tekoälyyn liittyvistä ongelmista johtuu epäselvistä ohjeista. Jos opiskelija ei tiedä, mitä tekoälyllä saa tehdä ja mitä ei, syntyy helposti väärinkäsityksiä ja epävarmuutta. Aiemmat tutkimukset ovat osoittaneet selkeän ohjeistuksen vähentävän eettistä epävarmuutta ja tukevan opiskelijoiden vastuullista toimintaa (Bretag, 2013; Sharma & Panja, 2025).

Suomalaisissa ammattikorkeakouluissa tähän haasteeseen on vastattu Arenen kehittämällä liikennevalomallilla (Arene, 2024). Mallissa vihreä tarkoittaa, että tekoälyn käyttö on sallittua, keltainen rajattua käyttöä ja punainen sitä, että tekoälyn käyttö on kielletty. Joissakin tilanteissa voidaan hyödyntää myös sinistä tasoa, jossa tekoälyn käyttö on osa oppimistehtävää ja arvioinnin kohde.

Liikennevalomallin etuna on sen yksinkertaisuus. Opiskelija näkee nopeasti, millaisia odotuksia tehtävään liittyy, ja opettajan on helpompaa viestiä omat vaatimuksensa. Selkeät pelisäännöt lisäävät läpinäkyvyyttä ja tukevat opiskelijoiden eettistä päätöksentekoa.

Tekoälyä hyödynnetään myös vilpin tunnistamisessa

Mielenkiintoista kyllä, tekoälyä käytetään nykyisin myös tekoälyn mahdollistaman vilpin tunnistamiseen. Eri korkeakouluissa kehitetään ratkaisuja, jotka hyödyntävät esimerkiksi tekstianalyysiä, käyttäytymisen seuranta ja objektintunnistusta.

Myös SEAMKissa kehitetään parhaillaan tekoälyavusteista järjestelmää EXAM-tenttien videotallenteiden tarkastamiseen. Ratkaisu hyödyntää koneoppimista ja YOLO-objektintunnistusta tunnistaakseen tenttitilanteissa mahdollisesti kiellettyjä esineitä, kuten puhelimia, älykelloja tai muistiinpanoja. Tavoitteena ei

ole korvata ihmisen tekemää arviointia, vaan auttaa tarkastajia kohdentamaan huomionsa niihin tilanteisiin, joissa havaitaan poikkeamia.

Teknologisilla ratkaisuilla on kuitenkin myös rajoituksensa. Väärät hälytykset, yksityisyydensuojan kysymykset ja eettiset näkökulmat edellyttävät huolellista harkintaa. Tutkimuksen mukaan teknologia voi toimia tukena, mutta lopulliset päätökset tulee edelleen jättää ihmiselle.

Arviointi muuttuu tekoälyaikana

Tutkimuksen keskeinen johtopäätös on, ettei tulevaisuuden ratkaisu löydy pelkästä valvonnasta. Sen sijaan korkeakoulujen on kehitettävä arviointia siten, että opiskelijan oma ajattelu, osaamisen soveltaminen ja oppimisprosessi tulevat näkyviksi. Tämä tukee myös kansainvälisen tutkimuksen havaintoja autenttisten tehtävien, reflektiivisen oppimisen ja prosessiarvioinnin merkityksestä tekoälyaikana (Kasneci ym., 2023; Sharma & Panja, 2025).

Käytännössä tämä voi tarkoittaa esimerkiksi prosessipohjaista arviointia, autenttisia työelämälähtöisiä tehtäviä, suullisia esityksiä, reflektiota tai tehtäviä, joissa opiskelijan on perusteltava ja arvioitava omaa toimintaansa. Tällaisissa tilanteissa tekoäly voi toimia oppimisen tukena ilman, että se korvaa opiskelijan omaa osaamista.

Leppäsen ja Taijalan mukaan tavoitteena ei ole estää tekoälyn käyttöä, vaan rakentaa toimintakulttuuria, jossa sitä osataan hyödyntää vastuullisesti ja eettisesti. Tekoäly tulee olemaan osa tulevaisuuden työelämää, joten myös korkeakoulujen tehtävänä on opettaa sen tarkoituksenmukaista käyttöä. Samalla on huolehdittava siitä, että osaamisen arviointi säilyy luotettavana ja opiskelijoiden akateeminen rehellisyys vahvistuu myös tekoälyn aikakaudella.

Tämä artikkeli perustuu kirjoittajien heinäkuussa 2026 ICTEM-konferenssissa esittämään abstraktiin sekä sitä taustoittavaan tutkimustyöhön tekoälyn mahdollistamasta akateemisesta vilpistä ja sen ehkäisemisestä korkeakoulutuksessa. Aihetta tarkastellaan syvällisemmin kirjoittajien laatimassa SEAMK Journaliin tarjotussa tutkimusartikkelissa. Kim Leppänen

Lehtori

SEAMK

Beata Taijala

Lehtori

SEAMK

Kirjoittajat ovat pitkän kokemuksen omaavia korkeakoulupedagogiikan kehittäjiä. Viime vuosina he ovat julkaisseet useita tekoälyn hyödyntämistä, oppimista ja arviointia käsitteleviä artikkeleita. Heitä kiinnostavat erityisesti tekoälyn pedagoginen käyttö, opiskelijoiden tekoälyosaamisen kehittäminen sekä akateemiseen rehellisyyteen ja vastuulliseen tekoälyn käyttöön liittyvät kysymykset korkeakoulutuksessa.

Lähteet

Arene. (2024). *Arenen suositukset tekoälyn hyödyntämisestä ammattikorkeakouluille*. Ammattikorkeakoulujen rehtorineuvosto Arene ry.

Bittle, K., & ElGayar, O. (2025). Generative AI and academic integrity in higher education: A systematic review and research agenda. *Information*, 16(4), 296. <https://doi.org/10.3390/info16040296>

Bretag, T. (2013). *Challenges in addressing plagiarism in education*. *PLOS Medicine*, 10(12), e1001574. <https://doi.org/10.1371/journal.pmed.1001574>

Eaton, S. E. (2019). *Academic integrity in Canada: An enduring and essential challenge*. Springer.

Gray, S. L., Edsall, D., & Parapadakis, D. (2025). AI-based digital cheating at university and the case for new ethical pedagogies. *Journal of Academic Ethics*, 23, 2069-2086.

Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., ... & Kasneci, G. (2023). *ChatGPT for good? On opportunities and challenges of large language models for education*. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>

Oravec, J. A. (2023). Artificial intelligence implications for academic cheating: Expanding the dimensions of responsible human-AI collaboration with ChatGPT. *Journal of Interactive Learning Research*, 34(2), 213-237.

Sharma, R. C., & Panja, S. K. (2025). Addressing academic dishonesty in higher education: A systematic review of generative AI's impact. *Open Praxis*, 17(2), 251-269.