

Miika Koskinen ja Hans Moen

Kieliteknologia terveydenhuollossa

Kieliteknologian kehittäminen ja integrointi terveydenhuollon tietojärjestelmiin ja asiakkaille suunnattuihin digitaalisiin palveluihin voi merkittävästi parantaa ohjelmistojen käytettävyyttä, nopeuttaa kirjaamista, tiedon löytämistä ja käsittelyä sekä edistää prosessien automatisointia. Terveydenhuollon jatkuvasti laajenevat tekstiaineistot ovat arvokas tietolähde ja sisältävät informaatiota, jota rakenteisessa datassa ei ole. Kieliteknologia avaa uusia mahdollisuuksia työkalujen kehittämiseen niin olemassa olevien massiivisten kuin myös pienten tekstiaineistojenkin hyödyntämiseen klinisen työn, tietojohdamisen ja tutkimuksen tarpeisiin. Teknologinen kehitys on hyvin nopeaa, moni menetelmä on kehitysvaiheessa ja terve kriittisyys on tarpeen. Teknistä suorituskkyä arvioitaessa tulee huomioida myös kieliteknologiaan kiinteästi liittyvät inhimilliset tekijät.

Kieliteknologia, toiselta nimeltään luonnollisen kielen käsittely (natural language processing, NLP), tarkoittaa mallipohjaisia menetelmiä ja työkaluja, jotka mahdollistavat tekstin ja puheen analysoinnin ja tuottamisen sekä vuorovaikutuksen koneiden kanssa luonnollisella kielellä (1). NLP on yksi tekoälyn kehittämisen keskeisistä alueista, jonka mahdollisuuksia terveydenhuollossa on hyödynnetty niukalti mutta jonka merkitys on ratkaiseva tekstiaineistojen hyödyntämisen kannalta. Merkittäviä nykyisiä sovelluskohteita ovat muun muassa ohjelmistojen ja digipalvelujen käytettävyyden kehittäminen ammattilaisen ja asiakkaan näkökulmasta, terveydenhuollon prosessien automatisointi sekä tiedon rakenteistaminen tietojohdamisen ja tutkimuksen tarpeisiin (2).

Terveydenhuollon vaatimukset

Tekstien ominaispiirteet. Terveydenhuollon teksteillä on ominaispiirteitä, jotka poikkeavat monista toimialoista. Terveydenhuollossa hyvin suuri osa viestinnästä tapahtuu ammattilaisten kirjoittamien läheteiden, lausuntojen ynnä muiden dokumenttien muodossa (3). Tekstityyppejä erikoisaloilla lienee tuhansittain, esimerkiksi radiologian lausunnoilla on

TAULUKKO. HUS:n digitaalisen tekstiaineiston koko käsittelemättömänä on arviolta 14,5 miljardia sanaa, mikä vastaa 250 000 kirjaa.

Tekstityyppi	Sanoja (miljoonaa)
Potilastekstit	13 200
Radiologian lausunnot	909
Patologian lausunnot	343
Ensihoito	50
Asiakaspalautteet	28

nimikekohtaisesti ominainen rakenne, sisältö ja sanasto.

Monista usein NLP-kehityksessä käytetyistä aineistoista poiketen asiakirjat eivät ole toisistaan riippumattomia, vaan tekstit ja rakenteinen tieto hoidon eri vaiheista muodostavat yhdessä hierarkkisen ja ajassa etenevän jatkokertomuksen potilaan tilasta, oireista, diagnooseista, sairauden etenemisestä, hoitopäätöksistä, toimenpiteistä ja niiden vaikuttavuudesta (**TAULUKKO**). Lääketieteen sanasto, ammattislangi ja kirjoitustyyli poikkeavat yleiskielestä. Myös suomen kieli jää helposti marginaaliin teknologiaa kehitettäessä.

Kun klinisen tekstin erityispiirteet ja alan laatuvaatimukset huomioidaan, terveydenhuollon aineistoilla opetetut kielimallit voivat toimia yleisesti käytettyjä aineistoja paremmin esimerkiksi ennustettaessa, ketkä potilaat pa-

laavat takaisin sairaalahoitoon (4). Toimialan erityisyys ja pieni kielialue lisäävät tarvetta kehittää nimenomaan suomalaisiin aineistoihin pohjautuvia teknologioita.

Vapaamuotoinen teksti ja rakenteinen data. Vapaamuotoinen teksti on kirjoittajan ja lukijan näkökulmasta luonteva mutta analytiikan, automatisoinnin ja suurten aineistojen käsittelyn kannalta haasteellinen ilmaisumuoto. Tietojohdaminen ja lääketieteellinen tutkimus tarvitsevat helposti käsiteltävää, rakenteellista dataa, ja tähän esimerkiksi diagnoosi- ja toimenpidekoodit soveltuvat hyvin. Luonnollinen kieli on kuitenkin rikasta, eivätkä ontologiat välttämättä taivu vastaavaan vivahteikkuteen, esimerkkinä potilaiden oireet, suunniteltu hoito tai leikkauskertomukset (5,6).

Potilastekstit sisältävät usein tietoa, jota rakenteisesta datasta ei löydy. Eräässä kansainvälisessä selvityksessä todettiin, että valittaessa koehenkilöitä kliniseen tutkimukseen, 38–89 % (mediaani 55 %) kelpoisuusvaatimuksista voitiin arvioida rakenteisen datan perusteella (7). Siten tekstikenttiin jäävän informaation määrä on vielä huomattava. Tilannetta on mahdollista parantaa rakenteisella kirjauksella, mutta myös hyödyntämällä kieliteknologian mahdollisuuksia informaation muuttamisessa rakenteiseen muotoon. Esimerkiksi HUS:ssa on käytössä reseptejä rakenteistava sovellus.

Rakenteisen datan tuottaminen ei kuitenkaan ole aina välttämätöntä. Koneoppimismallit kykenevät oppimaan opetusaineiston merkijonoista säännönmukaisuuksia, yhteyksiä, kieliooppia, jopa semantiikkaa (representation learning). Kielimallien tuottamaa abstraktia esitystapaa ja numeerista tulkintaa informaatioisällöstä voidaan hyödyntää jatkoanalyyseissa esimerkiksi vektorimuotoisena tai liittämällä loppusovelluksessa tarvittavat mallikomponentit suoraan osaksi kielimallia.

Työkaluja ja sovellusesimerkkejä

Säännölliset lausekkeet. Yleisesti hyödynnetty menetelmä sanojen ja lauseiden löytämiseen ja poimimiseen tekstistä on säännöllisten lausekkeiden (regular expressions) käyttö. Perusajatuksena on suoraviivaisesti vertailla et-

sittävää merkijonoa tai merkijonorakennetta tekstin kanssa. Menetelmä ei perustu koneoppimiseen tai kielimalleihin, vaan merkijonon etsintäsäännöt koodataan tarkasti tai etsittävät sanat listataan, jolloin taivutusmuodoilla, kirjoitusvirheillä ja vaihtoehtoisilla ilmaisutavoilla on merkitystä. Säännölliset lausekkeet soveltuvat yksinkertaisimpiin hakutehtäviin, esimerkiksi maininnat etäpesäkkeistä, ejektiofraktio, NYHA-luokka, lääkkeiden nimet tai tietyt oireet.

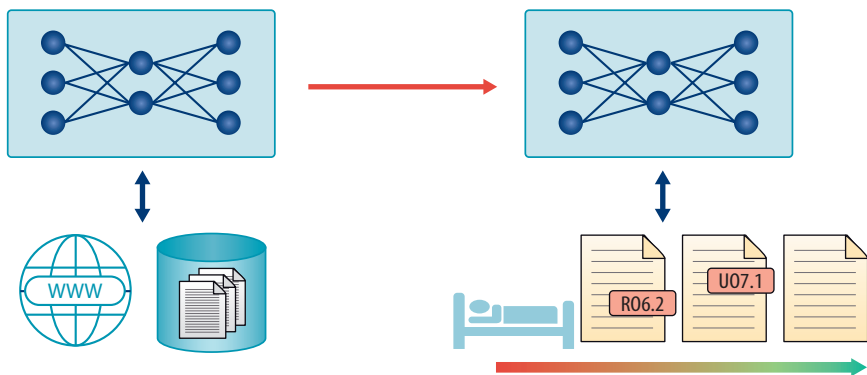
Luokittelu. Koneoppimista hyödyntävistä sovelluksista dokumenttien luokittelu lienee yleisin (8). Luokittelumalli opetetaan yhdistämään dokumentti haluttuihin kategorioihin opetusaineiston avulla, esimerkiksi tupakoiko potilas aktiivisesti, onko aiemmin polttanut, onko koskaan polttanut tai löytyykö tätä tietoa ollenkaan (9). Malli oppii, mikä tekstissä on merkittävää luokittelutehtävän kannalta.

Luokittelua on hyödynnetty muun muassa fenotyyppien ja tapahtumien löytämisessä, oireiden vakavuuden luokittelussa ja lasten pahoinpitelystä kertovan tekstin tunnistamisessa (8,10). Viime aikoina mielenkiintoa on herättänyt myös tekstiperustainen ICD-10-koodaus, jonka tarkoituksena on nopeuttaa diagnoosin antamista sekä täydentää puuttuvia tai korjata virheellisiä merkintöjä (KUVA) (11).

Luokittimet ovat hyödyllisiä prosessien automatisoimisessa. Esimerkiksi HUS:ssa tekstiluokittimia käytetään lähetelajittelussa, jossa automatisoinnin tarkoituksena on vähentää käsittelyviivettä, vapauttaa henkilöresursseja ja säästää kustannuksia.

Käytettävyys. Kieliteknologian kehittämisen sekä menetelmien integrointi potilastietojärjestelmiin ja asiakkaille suunnattuihin digitaalisiin palveluihin voi parantaa merkittävästi ohjelmistojen käytettävyyttä, nopeuttaa kirjauksia, tiedon löytymistä ja käsittelyä. Järjestelmiä ja palveluita on syytä kehittää älykkäämmiksi (12).

Esimerkkejä. Tekstin oikoluku kirjoitettaessa on ollut arkipäivää jo vuosia. Nykyisin suosittu tekstin luettavuutta ja halutun tiedon löytymistä helpottava koneoppimismenetelmä on nimetyn kohteen tunnistus (named entity recognition, NER), joka mahdollistaa



Kielimallin itseohjautuvassa opetuksessa hyödynnetään suurta määrää eri lähteistä poimittuja tekstejä.

Esiopetettua mallia voidaan hienosäätää esimerkiksi luokittelu-tehtävää varten, kuten liittämään diagnoosikoodi tekstiin.

KUVA. Esiopetetut kielimallit sisältävät tietoa kielen rakenteesta ja riippuvuuksista. Malleja voidaan edelleen opettaa pienemmällä opetusaineistolla ohjatusti eri käyttötarkoituksia varten. Ohjattu opetus tapahtuu esimerkiksi, kuten annotoitujen tekstien avulla.

sanatasolla käsitteiden tunnistamisen tekstistä ja liittämisen ontologioihin: esimerkiksi eturauhanen on anatominen rakenne, asetyylisalisyylihappo on lääke ja 50 mm tarkoittaa mitaustulosta (3,8,13). Malli opetetaan ohjatusti opetusaineiston perusteella, toisin kuin esimerkiksi säännölliset lausekkeet. Menetelmää on sovellettu muun muassa diagnoosien, oireiden ja hoidon tunnistamisessa tekstistä, henkilötiedon poistamisessa ja negatiivisten tunnistamisessa lauseyhteydestä (8,14). Käsitteiden tunnistus on hyödyllistä myös tiedon struktuuroinnissa.

Lisääntyvä määrä kieliteknologian kehitystä liittyy dokumentointiin menevän ajan säästämiseen (15,16). Nopean yleiskuvan saamista potilaan tilanteesta on mahdollista tukea esimerkiksi tuottamalla yhteenveto teksteistä automatiikan avulla (17,18). Tekstiä kirjoitettaessa kone voi auttaa tekstin muotoilussa ja otsikoinnissa. Esimerkiksi Tyksissä hoitajat jakavat kirjoittamansa tekstin kappaleisiin ja antavat kappaleille standardimuotoiset otsikot, joita voi olla yli 600 erilaista (19). Kirjoitetusta vapaamuotoisesta tekstistä, kuten puheestakin, on mahdollista erotella lauseet, joille algoritmi ehdottaa otsikkoa. Lauseet voidaan ryhmitellä otsikoiden mukaan määrämuotoiseksi tekstiksi.

Chatbotit ovat vakiinnuttaneet asemaansa asiakkaiden verkkopalveluissa. HUS:lla on käy-

tössä Neuvo-botti, joka antaa hoitoon pääsyyn, kuvantamiseen valmistautumiseen, yhteystietoihin, pysäköintiin ja maksuihin liittyvää tietoa, sekä virtuaaliapuri Milli, joka auttaa arvioimaan nuoren ahdistus- tai masennusoireiden vakavuutta ja ohjaamaan käyttäjää palvelujen piiriin (20,21). Keskustelujen määrä Millin kanssa vuositasona on ollut kymmeniä tuhansia, ja käyttäjäpalautte on ollut yllättävänkin positiivista.

Menetelmistä

Koneoppimismallien opetus. Koneoppimismallit opetetaan tehtävänsä esimerkkien avulla tarkoituksena löytää säännönmukaisuuksia ja riippuvuuksia opetusdatasta. Esimerkiksi yksinkertaisen suoran sovituksessa datapisteisiin regressiomallin tapauksessa mallin opetus on perimmiltään optimointiongelma, jonka tavoitteena on löytää sopivat malliparametrien lukuarvot niin, että mallin tekemä ennustevirhe minimoituu. Optimointiyhtälöön voidaan tuoda lisäksi esimerkiksi regularisaatiotermejä, joilla pyritään välttämään mallin ylisovittamista opetusaineistoon.

Koneoppimismallien opetus on iteratiivista, jolloin laskennasta tulee raskaampaa parametrien ja opetusdatan määrän lisääntyessä. Toisaalta malliparametrien suuri määrä voi olla

Ydinasiat

- ▶ Kieliteknologiaa voidaan hyödyntää käytettävyyden kehittämiseen, prosessien automatisointiin ja tiedon rakenteistamiseen.
- ▶ Yleisiä kielimallien sovelluksia ovat muun muassa dokumenttien luokittelutehtävät sekä käsitteiden tunnistaminen tekstistä ja liittäminen ontologioihin.
- ▶ Kielimallien tuottamaa abstraktia esitystapaa ja numeerista tulkintaa tekstin informaatioisällöstä voidaan hyödyntää jatkoanalyysissa.
- ▶ Pieni kielialue ja toimialan erityisyys lisäävät tarvetta kehittää suomalaisiin terveydenhuollon aineistoihin pohjautuvia teknologioita ja sovelluksia.

tarpeen, jotta malli kykenee oppimaan säännönmukaisuuksia riittävän yksityiskohtaisesti. Opetusaineiston kattavuus on keskeistä tosielämän sovelluksille. Myös mallin toimivuuteen datalla, jota ei ollut mukana opetusaineistossa, eli kykyyn yleistää sekä mahdollisiin vinoumiin aineistossa tulee kiinnittää huomiota.

Suuret kielimallit. NLP-sovellusten taustalla on iso kirjo teknologioita, mutta nykyisin vallalla ovat syviin hermoverkkomalleihin perustuvat, sana- ja merkkijonojen käsittelyyn soveltuvat menetelmät, erityisesti muuntajapohjaiseen (transformer) arkkitehtuuriin perustuvat mallit, kuten Bidirectional Encoder Representations from Transformers, BERT (22,23). Mallien ei tarvitse rajautua vain yhden kielen aineistoihin, vaan malli voi hyötyä usean kielen opetusaineistosta (multilingual-mallit) erityisesti konekäännössovelluksissa, mutta myös kun opetusaineisto yhdellä kielellä on niukka (esimerkiksi mT5-malli) (24).

Kielimallien parametreja voidaan optimoida itseohjautuvasti esimerkiksi jättämällä sanoja pois ja antamalla mallin ennustaa puuttuvat sanat (masked language modeling) tai ennakoida seuraava lause. Suuren kielimallin opetus on aikaa vievää sekä laskennallisesti raskasta ja kallista parametrien määrän ja opetuk-

seen tarvittavan aineiston suuren koon vuoksi. Raskaiden mallien opetuskustannukset voivat suurentua kymmeniin tuhansiin tai miljooniin euroihin. Jopa mallien opetuksen aiheuttamaan hiilijalanjälkeen on kiinnitetty huomiota (25).

Sovelluksia varten suuria kielimalleja ei tarvitse alusta saakka opettaa, vaan valmiiksi opetettua mallia voidaan käyttää lähtökohtana ja hienosäätää tapauskohtaisesti pienemmällä aineistolla toista tehtävää, kuten luokittelua varten (transfer learning) (KUVA). Moni valmiiksi opetettu malli on avoimesti saatavilla myös suomen kielellä ja lisenssin sallimin ehdoin räätölitävissä haluttuun tehtävään (26,27).

Tekstin tuottaminen. Vastaavaa teknologiaa voidaan käyttää tekstin tuottamiseen. Nämä ovat tyypillisesti merkkijonosta merkkijonoon tehtäviä, kuten konekäännös, jossa alkuperäinen teksti muunnetaan (koodataan) abstraktimpaan esitysmuotoon, jota käytetään edelleen lähtökohtana tulostekstin tuottamisessa (dekoodauksessa). Tekstin luomiseen käytetään mallin dekooderiosaa ja siementekstiä uuden tekstin tuottamiseksi.

Tämä lähestymistapa tuli tunnetuksi erityisesti generative pre-trained transformer (GPT)-mallien myötä, joista esimerkiksi ChatGPT:n jatkokehityksen pohjana oli GPT-3, jossa on 175 miljardia parametria (28). Tämän kaltaisia malleja kutsutaan suuriksi kielimalleiksi (large language models, LLM). Uusin GPT-4-malli voi hyödyntää jopa 32 000 merkin mittaista kontekstia päätöksessä, mikä sana kirjoitetaan seuraavaksi. Malli kykenee hyödyntämään myös kuvainformaatiota, mikä tekee siitä multimodaalisen. GPT-4-mallin parametrien määrä ei ole julkinen tieto, arviot liikkuvat sadoissa miljardeissa. Opetusdatan määrää on myös lisätty aiempaan verrattuna.

Osa tekstiä generoivista malleista on saatavilla eri kielillä, myös suomeksi, ja myös niitä voidaan hienosäätää suhteellisen helposti haluttua tehtävää varten (26,27).

Tekstiä tuottavia kielimalleja käytetään terveydenhuollossa chatboteissa sekä kirjoittamisen apuvälineenä, yhteenvetojen tuottamiseksi automaattisesti ja synteettisen tekstin luomisessa (3,29–31).

Ongelmia ratkottavaksi. Vaikka suuret kielimallit ovat viime aikoina olleet usein otsikoissa, niiden käyttö ei aina ole välttämätöntä. Hyvin moni toiminto voidaan käytännössä toteuttaa pienemmillä malleilla, eri mallityypeillä ja kohtuullisilla laskentaresursseilla.

Hermoverkkopohjaisia malleja kutsutaan usein mustiksi laatikoiksi, koska sen tulkitseminen, miten malli on päätenyt tiettyyn tulokseen, voi olla vaikeaa. Esimerkiksi diagnostisen tuen sovelluksissa ymmärrys päättelyketjusta on oleellista luotettavuuden arvioinnin kannalta. Läpinäkyvyys on selitettävissä olevan tekoälytutkimuksen (explainable artificial intelligence, XAI) keskeinen motiivi (32).

Esimerkiksi eräässä tutkimuksessa arvioitiin ensihoidon kuljettamattajättämisspäättöksen seurauksia (33). Ensihoidon dokumentaatiota ambulanssikäynneistä muutamien potilaan lisätietojen kanssa käytettiin ennakoimaan joko uutta yhteydenottoa hälytyskeskukseen, perusterveydenhuollon käyntiä tai sairaalahoitoa 48 tunnin kuluessa taikka kuolleisuutta 28 päivän kuluessa. XAI-menetelmällä pyrittiin tuomaan teksteistä esiin tekijöitä, sanoja ja lauseita, joihin luokittelumalli eniten kiinnitti huomiota, jotta voitiin arvioida mallin luotettavuutta ja sitä, vastaako mallin tekemä luokittelupäätös asiantuntijan tekemää.

Koneoppimiseen liittyy aina jonkin verran epävarmuutta. Automaattinen lähetelajittelija ei aina kykene päättämään läheteelle osoitetta vaan palauttaa epävarmat manuaalisesti arvioitaviksi. Mallien onnistumiseen liittyy kuitenkin usein inhimillisiä tekijöitä, kuten tekstien laatu, ilmaisun selkeys tai se, miten opetusaineiston

manuaalinen annotointi on onnistunut. Jopa asiantuntijan voi toisinaan olla vaikeaa tulkita lausunnosta, onko potilaalla etäpesäkettä vai ei.

Suomen mahdollisuudet

Tekstien arkaluonteisuuden vuoksi kieliteknologian tutkimus ja kehitys terveydenhuollon alalla on verraten rajoittunutta (5). Suuressa osassa alan julkaistuista hyödynnetään yhdysvaltalaisista Medical Information Mart for Intensive Care (MIMIC) -aineistoa. Suomessa on kuitenkin valtavasti terveydenhuollon dataa, tutkimuksen mahdollistava lainsäädäntö ja tietoturvallisia, tehokastaan soveltuvia käyttöympäristöjä, kuten HUS Acamedic tai Aurian analytiikkaympäristö, jotka mahdollistavat terveydenhuollon tekoälytutkimuksen ja antavat kansainvälisestäkin verrattuna erinomaiset edellytykset kehittää menetelmiä ja räätälöidä ratkaisuja suomalaisen terveydenhuoltoon sopiviksi.

Lopuksi

Tietojohtamisessa ja tutkimuksessa käytetään voittopuolisesti rakenteista dataa, kun taas sairaalan päivittäistoiminnassa tuotetaan runsaasti tekstiä. Kieliteknologian tutkimus, kehitys ja käyttöönotto avaavat uusia, vielä hyödyntämättömiä mahdollisuuksia sekä olemassa olevan tekstimateriaalin tehokkaampaan hyödyntämiseen että kirjausten ja tiedonhaun tehostamiseen, ohjelmistojen käytettävyyden parantamiseen ja automatisaation lisäämiseen päivittäistoiminnassa. ■

MIIKA KOSKINEN, TkT, dosentti
HUS Helsingin yliopistollinen sairaala, tietohallinto

HANS MOEN, PhD
Aalto-yliopisto, tietotekniikan laitos

VASTUUTOIMITTAJA
Tuomas Mirtti

SIDONNAISUUDET
Miika Koskinen: Tutkimusrahoitus (Bayer Oy), muut sidonnaisuudet (Aplagon Oy)
Hans Moen: Ei sidonnaisuuksia

KIRJALLISUUTTA

1. Koskeniemi K, Linden K, Carlson, L, ym. (2012). Suomen kieli digitaalisella aikakaudella. Valkoiset kirjat, META-NET. Berliini: Springer-Verlag 2012. <http://www.meta-net.eu/whitepapers/e-book/finnish.pdf>.
2. Friedman C, Rindflesch TC, Corn M. Natural language processing: state of the art and prospects for significant progress, a workshop sponsored by the National Library of Medicine. *J Biomed Inform* 2013;46:765–73.
3. Li J, Pan J, Goldwasser J, ym. Neural natural language processing for unstructured data in electronic health records: a review. *Comput Sci Rev* 2022;46:100511.
4. Huang K, Altosaar J, Ranganath R. ClinicalBERT: modeling clinical notes and predicting hospital readmission. *arXiv* 2019. DOI:10.48550/arXiv.1904.05342.
5. Velupillai S, Suominen H, Liakata M, ym. Using clinical natural language processing for health outcomes research: overview and actionable suggestions for future advances. *J Biomed Inform* 2018;88:11–9.
6. Kolec TA, Tatonetti NP, Bakken S, ym. Identifying symptom information in clinical notes using natural language processing. *Nurs Res* 2021;70:173–83.
7. Claerhout B, Kalra D, Mueller C, ym. Federated electronic health records research technology to support clinical trial protocol optimization: evidence from EHR4CR and the InSite platform. *J Biomed Inform* 2019;90:103090.
8. Wu S, Roberts K, Datta S, ym. Deep learning in clinical natural language processing: a methodical review. *J Am Med Inform Assoc* 2020;27:457–70.
9. Karlsson A, Ellonen A, Irla J, ym. Impact of deep learning-determined smoking status on mortality of cancer patients: never too late to quit. *ESMO Open* 2021;6:100175.
10. Annapragada AV, Donaruma-Kwoh MM, Annapragada AV, ym. A natural language processing and deep learning approach to identify child abuse from pediatric electronic medical records. *PLoS One*, julkaisu verkossa 26.2.2021. DOI:10.1371/journal.pone.0247404.
11. Dong H, Falis M, Whiteley W, ym. Automated clinical coding: what, why, and where we are? *NPJ Digit Med* 2022;5:159.
12. Zhou B, Yang G, Shi Z, ym. Natural language processing for smart healthcare. *IEEE Rev Biomed Eng* 2022. DOI:10.1109/RBME.2022.3210270.
13. Noor K, Roguski L, Bai X, ym. Deployment of a free-text analytics platform at a UK national health service research hospital: CogStack at University College London hospitals. *JMIR Med Inform* 24.8.2022. DOI:10.2196/38122
14. Chowdhury S, Dong X, Qian L, ym. A multitask bi-directional RNN model for named entity recognition on Chinese electronic medical records. *BMC Bioinformatics* 2018;19:499.
15. Yee T, Needleman J, Pearson M, ym. The influence of integrated electronic medical records and computerized Nursing notes on nurses' time spent in documentation. *Comput Inform Nurs* 2012;30:287–92.
16. von Gerich H, Moen H, Block LJ, ym. Artificial Intelligence-based technologies in nursing: a scoping literature review of the evidence. *Int J Nurs Stud* 2022; 127:104153.
17. Wang M, Wang M, Yu F, ym. A systematic review of automatic text summarization for biomedical literature and EHRs. *J Am Med Inform Assoc* 2021;28:287–97.
18. Searle T, Ibrahim Z, Teo J, ym. Discharge summary hospital course summarisation of inpatient Electronic Health Record text with clinical concept guided deep pre-trained Transformer models. *J Biomed Inform* 2023;141:104358.
19. Moen H, Hakala K, Peltonen LM, ym. Supporting the use of standardized nursing terminologies with automatic subject heading prediction: a comparison of sentence-level text classification methods. *J Am Med Inform Assoc* 2020;27:81–8.
20. Laboratorio ja kuvantaminen. HUS Helsingin yliopistollinen sairaala. www.hus.fi/potilaalle/hoidot-ja-tutkimukset/laboratorio-ja-kuvantaminen.
21. Virtuaaliapuri Milli on tukenut jo tuhansia suomalaisia – apuna vuorokauden ympäri. HUS Helsingin yliopistollinen sairaala. STT Info, tiedote 24.8.2020. www.sttinfo.fi/tiedote/virtuaaliapuri-milli-on-tukenut-jo-tuhansia-suomalaisia-apuna-vuorokauden-ympari?publisherId=23980819&releaseId=69886875.
22. Vaswani A, Shazeer N, Parmar N, ym. Attention is all you need. *Advances in neural information processing systems* 30 (NIPS 2017) 2017:6000–10.
23. Devlin J, Chang MW, Lee K, ym. BERT: pre-training of deep bidirectional transformers for language understanding. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies* 2019:4171–86.
24. Xue L, Constant N, Roberts A, ym. mT5: a massively multilingual pre-trained text-to-text transformer. *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies* 2021:483–98.
25. Quach K. AI me to the moon... carbon footprint for 'training GPT-3' same as driving to our natural satellite and back. The Register 4.11.2020. www.theregister.com/2020/11/04/gpt3_carbon_footprint_estimate/.
26. TurkuNLP, Turun yliopisto. Hugging Face. <https://huggingface.co/TurkuNLP>.
27. Finnish-NLP. Hugging Face. <https://huggingface.co/FinNish-NLP>.
28. Brown T, Mann B, Ryder N, ym. Language models are few-shot learners. *NIPS'20: Proceedings of the 34th International Conference on Neural Information Processing Systems* 2020:1877–901.
29. Sirriani J, Sezgin E, Claman D, ym. Medical text prediction and suggestion using generative pretrained transformer models with dental medical notes. *Methods Inf Med* 2022;61:195–200.
30. Liang S, Kades K, Fink M, ym. Fine-tuning BERT models for summarizing German radiology findings. *Proceedings of the 4th Clinical Natural Language Processing Workshop* 2022:30–40.
31. Chintagunta B, Katariya N, Amatriain X, ym. Medically aware gpt-3 as a data generator for medical dialogue summarization. *Machine Learning for Healthcare Conference* 2021:354–72.
32. Danilevsky M, Qian K, Aharonov R, ym. A survey of the state of explainable AI for natural language processing. *Proceedings of the 1st Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics and the 10th International Joint Conference on Natural Language Processing* 2020:447–59.
33. Paulin J, Reunamo A, Kurola J, ym. Using machine learning to predict subsequent events after EMS non-conveyance decisions. *BMC Med Inform Decis Mak* 2022; 22:166.