

Date of Hearing: July 8, 2025

ASSEMBLY COMMITTEE ON HEALTH

Mia Bonta, Chair

SB 503 (Weber Pierson) – As Amended June 30, 2025

SENATE VOTE: 38-0

SUBJECT: Health care services: artificial intelligence.

SUMMARY: Requires developers and deployers of artificial intelligence (AI) systems in specified health care applications to take steps to identify, mitigate, and monitor biased impacts, as specified. Specifically, **this bill:**

- 1) Requires developers of AI systems, as defined, and health facilities, clinics, physician's offices, or offices of a group practice, as defined, to make reasonable efforts to identify AI systems used to support clinical decision-making and healthcare resource allocation that are known or have a reasonably foreseeable risk for biased impacts in the systems output resulting from use of the system in health programs or activities.
- 2) Requires developers and deployers to make reasonable efforts to mitigate the risk for biased impacts in the systems outputs resulting from the use of the system and health programs or activities.
- 3) Requires deployers to regularly monitor these AI systems and take reasonable and proportionate steps to mitigate bias that may occur.
- 4) Specifies a person, partnership, state or local government agency, or corporation may be both a developer and a deployer.
- 5) Defines "protected characteristics" by reference to current state antidiscrimination law and defines additional terms as follows:
 - a) "Biased impact" means an unintended adverse impact, including diminished access to healthcare, quality of care, or outcomes, on an individual based on their protected characteristics;
 - b) "Deployer" means a person, partnership, state or local governmental agency, corporation, or developer that uses an artificial intelligence system to support clinical decisionmaking and health care resource allocation; and,
 - c) "Developer" means a person, partnership, state or local governmental agency, corporation, or deployer that designs, codes, substantially modifies, or otherwise produces an AI system for commercial or public use to support clinical decisionmaking and health care resource allocation.
- 6) Does not supplant or replace any other applicable provision of state law regulating the use of AI or automated decision systems, and prohibits the use of compliance with its requirements from being used as a defense to a claim of unlawful discrimination.

EXISTING LAW:

Federal law:

- 1) Prohibits, pursuant to the federal Patient Protection and Affordable Care Act (ACA), an individual, on the grounds prohibited under various civil rights laws, including title VI of the Civil Rights Act of 1964, title IX of the Education Amendments of 1972, the Age Discrimination Act of 1975, or section 504 of the Rehabilitation Act of 1973, from being excluded from participation in, be denied the benefits of, or be subjected to discrimination under, any health program or activity, any part of which is receiving Federal financial assistance, including credits, subsidies, or contracts of insurance, or under any program or activity that is administered by an Executive Agency or other entity, as provided. [42 United States Code (U.S.C.) § 18116 (“Section 1557”)]
- 2) Provides that the enforcement mechanisms provided for and available under federal laws described in 1) above shall apply for purposes of violations of the above. Authorizes the Secretary of the federal Health and Human Services Agency (HHS) to promulgate regulations relevant thereto. [42 U.S.C. § 18116]
- 3) Provides the following regulatory guidelines with regards to 1) above:
 - a) A covered entity must not discriminate on the basis of race, color, national origin, sex, age, or disability in its health programs or activities through the use of patient care decision support tools;
 - b) A covered entity has an ongoing duty to make reasonable efforts to identify uses of patient care decision support tools in its health programs or activities that employ input variables or factors that measure race, color, national origin, sex, age, or disability; and,
 - c) For each patient care decision support tool identified in (b), a covered entity must make reasonable efforts to mitigate the risk of discrimination resulting from the tool's use in its health programs or activities. [45 Code of Federal Regulations (C.F.R.) § 92.210]

State law:

- 1) Establishes the Unruh Civil Rights Act (“Unruh”), which provides that all persons within the jurisdiction of this state are free and equal, and no matter what their sex, race, color, religion, ancestry, national origin, disability, medical condition, genetic information, marital status, sexual orientation, citizenship, primary language, or immigration status, are entitled to the full and equal accommodations, advantages, facilities, privileges, or services in all business establishments of every kind whatsoever. [Civil Code (CIV) § 51]
- 2) Requires a health facility, clinic, physician’s office, or office of a group practice that uses (GenAI) to generate written or verbal patient communications pertaining to patient clinical information to ensure that those communications include a disclaimer that indicates to the patient that the communication was generated by GenAI, and clear instructions on how to contact a human person. Exempts from this requirement a communication read and reviewed by a human licensed or certified health care provider. [Health and Safety Code (HSC) § 1339.75]

- 3) Requires a health plan or insurer that uses an AI, algorithm, or other software tool, and subcontractors of those health plans and insurers as specified, to, among other requirements, ensure the use of the AI, algorithm, or other software tool does not discriminate, directly or indirectly, against enrollees in violation of state or federal law. [HSC § 1367.01, Insurance Code § 10123.135]
- 4) Requires licensure of health facilities, including clinics, by the Department of Public Health (DPH). [HSC § 1200 *et seq.*, § 1250 *et seq.*]
- 5) Defines “artificial intelligence” as an engineered or machine-based system that varies in its level of autonomy and that can, for explicit or implicit objectives, infer from the input it receives how to generate outputs that can influence physical or virtual environments. [Government Code § 11546.45.5]

FISCAL EFFECT: According to the Senate Committee on Appropriations, unknown ongoing costs, likely hundreds of thousands, for the DPH for state administration (Licensing and Certification Fund).

COMMENTS:

1) PURPOSE OF THIS BILL. According to the author, this bill is a crucial step towards ensuring fairness in health care by addressing the racial biases embedded in AI models and systems. This technology is becoming more prevalent in healthcare, yet research has shown that these systems can produce biased outputs that disproportionately affect communities of color. Without proper oversight, these biases can go unchecked, deepening existing disparities in our healthcare system. The author states that this bill will require collaboration between developers and healthcare facilities to identify AI tools used in the delivery of patient care and proactively work towards meaningfully reducing bias. By requiring identification, mitigation, and oversight, the author notes this bill will help promote safety, equity, and exceptional performance while protecting patients against avoidable harm. This bill was inspired by the 2023 California Reparations Task Force Report.

2) BACKGROUND.

- a) **AI.** AI is the mimicking of human intelligence by artificial systems. AI uses algorithms—sets of rules—to transform inputs into outputs. Inputs and outputs can be anything a computer can process: numbers, text, audio, video, or movement. AI is not fundamentally different from other computer functions; unlike other computer functions, however, AI is able to accomplish tasks that are normally performed by humans. Models trained on small, specific datasets in order to make recommendations and predictions are referred to as “predictive AI.” This differentiates them from GenAI, which are trained on massive datasets in order to produce detailed text, images, audio, and video.
- b) **AI in Health Care.** AI in health care is not new; AI models of varying degrees of sophistication have been developed and deployed in the health care setting for decades, and the use of both GenAI and predictive AI are growing. According to the National Academy of Medicine (NAM), GenAI and large language models (models designed for natural language processing tasks) have the potential to transform health and medicine as we know it: improving health care delivery, advancing medical research, and augmenting the capacity of clinicians to provide personalized care at an unprecedented scale.

However, NAM also notes that the potential for both breakthrough innovation and unintended consequences demands careful consideration. With the recent advancement of GenAI, particularly in natural language processing, interest in, use of, and hype over AI has grown rapidly and health care applications have proliferated. According to the market research firm Market.us, the global net value of GenAI in health care was approximately \$800 million in 2022, with projections to grow to \$17.2 billion by 2032.

c) Informational Hearing. On May 28, 2025, the Assembly Committees on Health and Privacy & Consumer Protection held an informational hearing titled, “*Generative Artificial Intelligence in Health Care: Opportunities, Challenges, and Policy Implications.*” The hearing examined:

- i) The history of AI in health care and current applications;
- ii) A range of challenges that pose barriers to responsible, effective adoption of AI, including bias, cognitive burden, safety and effectiveness, cost and resource equity, governance, “model drift” and the need for local validation, and transparency and explainability;
- iii) Other considerations, such as workforce, cost, reimbursement, liability, and data privacy; and,
- iv) Existing regulatory frameworks and best practices, including federal and state law and regulation and private efforts.

d) Current Use Cases. Researchers, health care providers and facilities, health plans and others are deploying AI for a range of tasks in biomedical and health research, as well as various administrative and clinical use cases. Common administrative use cases include billing, claims processing for health care providers and health plans, prior authorization review, and appointment scheduling and other routine, nonclinical communication. Common clinical and clinical-adjacent use cases include:

- i) Ambient scribe technology, which can generate draft clinical notes and patient summaries;
- ii) Realistic conversational voicebot agents that offer case management, appointment preparation, and health education;
- iii) Assistance to synthesize, augment, and interpret medical images;
- iv) Mental health support bots;
- v) Predictive models that predict health trajectories or risks for inpatients, identify high-risk outpatients to inform follow-up care, monitor health, and recommend treatments; and,
- vi) Clinical decision support systems designed to aid physicians in diagnosing, managing, and treating patients.

e) Racial, Ethnic and Gender Bias in AI. There is a famous saying in computer science: “garbage in, garbage out.” The performance of an AI is directly impacted by the quality,

quantity, and relevance of the data used to train it. If the data used to train the AI is biased, the tool's outputs will be similarly biased and the results can be inaccurate when applied to populations not reflected in the training data. This applies to both predictive and GenAI.

In their work on mitigating bias in AI, the Berkeley Haas Center for Equity, Gender and Leadership (Center) tracks publicly available instances of bias in AI systems using machine learning. In their analysis of around 133 biased systems across industries from 1988 to the present day, the Center found that 44% (59 systems) demonstrate gender bias, with 26% (34 systems) exhibiting both gender and racial bias.

When automated decision systems are deployed in healthcare, biased historical data can lead to patients being recommended substandard care on the basis of their race or ethnicity. In 2007, an automated decision system was developed to help doctors estimate whether it was safe for people who had delivered previous children through cesarean section to deliver subsequent children vaginally— a procedure that carries some risk. The system considered relevant factors as it made its decision, such as the woman's age, her reason for the previous cesarean, and how long ago the cesarean had been performed. However, a 2017 study found that the system was biased; it predicted Black and Latino people were less likely to have a successful vaginal birth after a cesarean than similar non-Hispanic white women. As a result, doctors performed more cesareans on Black and Latino people than on white people, perpetuating historical racial and ethnic biases.

Similarly, in 2019, a study discovered harmful racial bias in an AI tool developed by the health care company Optum – a subsidiary of UnitedHealth Group – and used by providers across the country to offer care management services. The tool assigned Black patients lower likelihoods of adverse health outcomes than equally at-risk white patients. The authors found that this happened because the tool was designed to predict healthcare costs instead of needs. Because the healthcare system has historically spent less on care for Black patients than white patients for the same health conditions, the tool was issuing a prediction that mirrored and perpetuated past discrimination.

The University of California San Francisco also reported bias in an algorithm used to identify potential appointment no-shows to facilitate double-booking for appointments. The program was confirmed to result in low-resourced and marginalized populations being double-booked more often than others, reflecting underlying structural inequalities and highlighting how these tools, if not studied and corrected for bias, that can create feedback loops that worsen discrimination.

- f) ACA Anti-Discrimination Regulations.** Section 1557, the civil rights provision of the ACA, prohibits discrimination on the grounds of race, color, national origin, sex, age, or disability in certain health programs and activities. Section 1557(c) of the ACA authorizes the Health and Human Services (HHS) Secretary to promulgate regulations to implement the nondiscrimination requirements of Section 1557.

Section 1557 only applies to “covered entities,” that is, health programs and activities that receive federal financial assistance from HHS. Examples of types of covered entities under Section 1557 include hospitals, health clinics, physicians' practices, community health centers, nursing homes, rehabilitation centers, health insurance issuers, and state Medicaid agencies.

In 2024, HHS Office of Civil Rights (OCR) issued regulations with regard to “patient care decision support tools,” defined as any automated or non-automated tool, mechanism, method, technology, or combination thereof used by a covered entity to support clinical decision-making in its health programs or activities. The regulation provides:

- i) A covered entity must not discriminate on the basis of race, color, national origin, sex, age, or disability in its health programs or activities through the use of patient care decision support tools.
- ii) A covered entity has an ongoing duty to make reasonable efforts to identify uses of patient care decision support tools in its health programs or activities that employ input variables or factors that measure race, color, national origin, sex, age, or disability.
- iii) For each patient care decision support tool identified in ii) above, a covered entity must make reasonable efforts to mitigate the risk of discrimination resulting from the tool's use in its health programs or activities.

The Federal Register discusses how OCR interprets the term “patient care decision support tools,” which includes the following: automated decision systems; AI; flowcharts; formulas; equations; calculators; algorithms; utilization management applications; software as medical devices; software in medical devices; screening, risk assessment, and eligibility tools; and diagnostic and treatment guidance tools. The regulatory definition for “patient care decision support tool” also includes non-automated and evidence-based tools that rely on rules, assumptions, constraints, or thresholds, as these also have the potential to result in discrimination.

In addition, with regard to the definition of “patient care decision support tool,” OCR states in the Federal Register that its regulation:

“applies to tools used in clinical decision-making that affect the care that patients receive. This includes tools, used by covered entities such as hospitals, providers, and [health plans and insurers] in their health programs and activities for “screening, risk prediction, diagnosis, prognosis, clinical decision-making, treatment planning, health care operations, and allocation of resources” as applied to the patient. The OCR clarifies that tools used for these activities include tools used in covered entities' health programs and activities to assess health status, recommend care, provide disease management guidance, determine eligibility and conduct utilization review related to patient care that is directed by a provider, among other things, all of which impact clinical decision-making.” [89 Federal Register (FR) § 37522]

- g) **Alignment with Federal Section 1557 Regulations.** This bill is similar to federal Section 1557 regulations in that deployers of these tools are required to make “reasonable efforts to mitigate the risk of discrimination on the basis of a protected characteristic resulting from the tool’s use in its health programs or activities,” similar to the federal regulation.

However, the bill is different than federal regulations in key ways. First, this bill is narrower than the federal regulation in the types of tools it covers, as it is specific to AI

systems and only includes tools related to clinical decision-making and health care resource allocation, a subset of the types of applications intended covered by the federal regulation, as described in the Federal Register and discussed above.

Second, this bill applies to fewer *health care entities* than the federal regulation applies to—for instance, the bill does not apply to health care plans and insurers—but it extends its requirements to *developers*, which are not generally covered under Section 1557 regulations. Covered entities that develop AI systems or similar tools in-house, however, would be covered by the federal regulation. This bill also requires monitoring, which is not explicitly addressed in Section 1557 regulations, although the regulations do specify an “ongoing duty” to identify and mitigate bias.

Third, the bill defines “protected characteristic” as those laid out in state civil rights law (see 1) above under Existing State Law). State law includes many more factors than the federal regulation, which prohibits discrimination based on “race, color, national origin, sex, age, or disability.”

- h) How is Bias in AI Systems Identified, Mitigated, and Monitored?** There is widespread awareness that bias is a problem that needs attention from developers and deployers of AI, and there is ongoing work to develop ways to measure and address it. As a practical matter, at this time, identifying, monitoring, and mitigating for bias may require a deployer to be aware of the literature on what types of systems may be prone to bias, understand how a model’s training data compares to their patient population, conduct sensitivity analyses to see how calibrating a model in different ways effects the outputs of the model, and make technical adjustments to a model. For instance, a clinic might have to calibrate or adjust the model in a specific way to ensure it works effectively for their particular patient population. As industry standards continue to develop to support such monitoring, it is possible that more of this work to mitigate bias will be done on the front end in the development of AI solutions, and thus the work of addressing bias may become less burdensome on individual providers over time.

An example can demonstrate what bias mitigation looks like in the field. University of California (UC) Davis researchers have developed a 9-step framework called BE-FAIR (e Bias-reduction and Equity Framework for Assessing, Implementing, and Redesigning) for organizations to use to assess and correct for bias in health care predictive AI models in development and implementation. According to an article in the *Journal of General Internal Medicine* published in March 2025, “*Developing and Applying the BE-FAIR Equity Framework to a Population Health Predictive Model: A Retrospective Observational Cohort Study*,” applying this framework allowed them to identify appropriate outreach thresholds for a population-based intervention: specifically, which patients are likely to benefit from care management services to deal with health problems before they lead to emergency department visits or hospitalization. A team of researchers evaluated the model’s performance over 12 months, and identified that the model underpredicted the probability of hospitalizations and emergency department visits for African American and Hispanic groups. The team then “identified the ideal threshold percentile to reduce this underprediction by evaluating predictive model calibration.” The researchers believe these analytic methods can be easily applied in other health systems to assess for bias. While predictive model development is resource intensive, the authors indicate, such evaluations for bias are feasible for both internally or vendor-developed

models with a part-time statistical analyst. However, for some providers, this level of analytics may be difficult to muster.

- i) **Cost and Resource Equity in AI Deployment.** While private hospital systems and commercial insurance plans have the financial and administrative resources to deploy AI technologies that can alleviate burdens on their workforce and improve patient care, recent work from the California Health Care Foundation (CHCF) concludes California's health care safety net is at risk of being left behind in their ability to adopt beneficial AI technologies.

CHCF, in partnership with the California Health and Human Services Agency, convened 45 safety-net leaders from across the state in three focus group sessions conducted between August and October 2024 to discuss their views on AI. According to CHCF, these conversations confirmed that safety-net organizations face restrictive barriers to the safe and effective adoption of AI. Many said their organizations cannot afford to integrate new digital tools into their workflows. Participants also raised workforce limitations and concerns about liability as barriers to adoption. According to Stella Tran, senior program investment officer at the CHCF Innovation Fund, if safety-net institutions miss out on the potential of AI, it could widen persistent racial and ethnic health disparities in that population and create a "tale of two health systems." For instance, although survey data does not appear to be available reflecting the levels of current adoption of ambient scribe technology that assists in generation of clinical notes among California health care providers, adoption has anecdotally have been rapid in better-resourced systems, while adoption among safety net providers has been slow.

- j) **Private and Industry Efforts on Transparency, Disclosure, and Evaluation.** Transparency about AI models is critical to identifying the potential for bias, and deployers often do not receive standardized information on model development that is needed to meaningfully identify and address bias. An August 2022 survey by the Office of California Attorney General (AG) Rob Bonta examined how California hospitals are addressing racial and ethnic disparities in their utilization of commercially available decision-making technologies. The AG reported the survey demonstrated these types of decision-making tools are now regularly used by hospitals to make judgments about patients across many contexts, ranging from medical treatments to managing revenue. Yet, the AG found, many hospitals report they rely on the vendor's assessment that the tools they use are ethical and unbiased, and that they lack insight into vendors' data modeling.

To address this lack of transparency and improve the effective and appropriate deployment of AI technology in health care, the Coalition for Health AI (CHAI), a large national collaborative effort of health systems, public and private organizations, academia, patient advocacy groups, and AI experts, has released a draft template for an "applied model card." The model card would be published by a developer and describe key information about health AI models, in a manner somewhat similar to a "Nutrition Facts" label. It would include the developer's name and contact information; summary; uses and directions, including intended use and patient population; warnings related to limitations, biases, ethical considerations, and clinical risk; system facts; key metrics; and other resources. Other items that may assist in assessing for bias include: a description of outcomes, outputs, and data used in the development of the model; known biases or

ethical considerations; continuous monitoring; Transparency, Intelligibility, and Accountability mechanisms; bias mitigation approaches; key evaluation metrics of the model, including those related to fairness and equity; and stakeholders consulted during the design of the solution.

Though CHAI expects to finalize the applied model card template soon, the effort is still in the draft phase and compliance with the proposed transparency measures is voluntary, as it is an industry-led effort to establish standards. If the model card were widely adopted, a number of elements included therein would assist a deployer of an AI model in identifying, mitigating, and monitoring for bias.

In addition, CHAI has proposed a model, which has been discussed among stakeholders and in the academic literature, whereby CHAI would certify labs that would rigorously evaluate AI models across pre-deployment, implementation, and post-deployment and monitoring. According to CHAI, the labs would focus on ensuring that AI models meet high standards for accuracy, reliability, and safety before they are deployed in clinical settings and provide an independent assessment to verify that AI tools function as intended and do not pose risks to patients. If such a system is developed, lab certification could potentially help deployers apply AI models more appropriately without having to rigorously evaluate the AI models themselves.

- k) Federal Transparency Requirements for Electronic Health Records Vendors.** The Assistant Secretary for Technology Policy/Office of the National Coordinator for Health IT (ONC) certifies electronic health records and other health information technology (IT) software. According to ONC, ONC-certified health IT supports the care delivered by more than 96% of hospitals and 78% of office-based physicians nationwide.

In their issuance of new regulations on EHRs in 2023, ONC also notes that bias in the predictions of predictive AI models can result in consequential adverse events. Starting in 2025, ONC's regulations subject EHRs vendors to a range of transparency requirements about AI and other predictive algorithms used as part of certified health IT. ONC's stated goal in issuing this regulation is to promote responsible AI and make it possible for clinical users to access a consistent, baseline set of information about the algorithms they use to support their decision making and to assess such algorithms for fairness, appropriateness, validity, effectiveness, and safety (often shortened to FAVES). These transparency requirements are significant, but they only apply to tools that are developed as modules within an EHR that is certified by ONC.

- l) Enforcement.** DPH licenses and oversees clinics and health facilities. Similar to requirements related to AB 3030 (Calderon), Chapter 848, Statutes of 2024, which put into place disclosure requirements for AI-generated patient communications, it is expected compliance with this bill would be enforced for clinics and health facilities by DPH licensing staff. In the Fiscal Year 2025-26 Budget Change Proposal implementing AB 3030, DPH notes there are currently 15,000 open health facilities and clinics with active licenses under CDPH's purview.

Other deployers to whom this bill applies, including physician's offices and offices of a group practice, are operated under the auspices of physician licensure and are not specifically regulated in California, outside of financial solvency requirements for risk-bearing organizations (often, large physician groups that take on financial risk for

delivering health care to a patient population). Developers, similarly, are not subject to any specific state oversight regime that would monitor compliance.

- 3) **SUPPORT.** Oakland Privacy, a citizen's coalition that works regionally to defend the right to privacy and enhance public transparency, writes in support that it is sensible to place into California statute a corresponding version of federal regulations that they believe support California's goals and priorities. Oakland Privacy notes, given the interesting definitions of "DEI" and "discrimination" being promulgated by some federal agencies, it makes even more sense to capture the intent of federal regulations before they are twisted into shapes quite unlike their original intent. The California Hospital Association and the California Medical Association also support this bill.
- 4) **SUPPORT IF AMENDED.** Students for Patient Advocacy Nationwide (SPAN), a national, student-led coalition focused on advocating for patient-centered healthcare reform, writes that this bill appropriately focuses on bias related to protected characteristics defined in Civil Code § 51(b), but it excludes patients impacted by socioeconomic inequity. SPAN seeks an amendment to include socioeconomic status as a protected characteristic. In support of this request, SPAN cites research showing that asthma-predicting AI models are less accurate for patients with low socioeconomic status, increasing risk of harm or missed care. SPAN asserts many AI tools rely on inputs like ZIP code or insurance type that function as proxies for income or neighborhood inequality, which can lead to biased care decisions even without violating civil rights statutes.

5) **RELATED LEGISLATION.**

- a) SB 243 (Padilla and Becker) would impose a number of obligations on operators of "companion chatbot platforms" in order to safeguard users. AB 243 is pending in the Assembly Privacy & Consumer Protection Committee.
- b) SB 420 (Padilla) would regulate the use of "high-risk automated decision systems (ADS)," including requirements on developers and deployers to perform impact assessments on their systems. SB 420 would establish the right of individuals to know when an ADS has been used, details about the systems, and an opportunity to appeal ADS decisions, where technically feasible, related to decisions that materially impact access to, or approval for, health care services, among numerous other services and opportunities. SB 420 is pending in the Assembly Privacy & Consumer Protection Committee.
- c) AB 1018 (Bauer-Kahan) would create a comprehensive regime designed to ensure human oversight over ADS that are used in "consequential decisions" – those that materially impact an individual's rights, opportunities, or access to critical resources or services – in order to mitigate bias and unreliability in these systems. AB 1018 is pending in the Senate Judiciary Committee.

6) **PREVIOUS LEGISLATION.**

- a) AB 2013 (Irwin), Chapter 817, Statutes of 2024, requires a developer of a GenAI system or service to publicly disclose specific information related to the system or service's training data, except as provided.

- b) AB 2930 (Bauer-Kahan) of 2024 would have regulated the use of ADSs in order to prevent “algorithmic discrimination.” This included requirements on developers and deployers that make and use these tools to make “consequential decisions” to perform impact assessments on ADSs. Would have established the right of individuals to know when an ADS is being used, the right to opt out of its use, and an explanation of how it is used. AB 2930 died on the Assembly Floor.
 - c) AB 3030 (Calderon), Chapter 848, Statutes of 2024 requires a health facility, clinic, physician’s office, or office of a group practice that uses GenAI to generate written or verbal patient communications pertaining to patient clinical information to ensure that those communications include a disclaimer that indicates to the patient that the communication was generated by GenAI and clear instructions on how the patient may contact a human person.
 - d) AB 2885 (Bauer-Kahan), Chapter 843, Statutes of 2024, established a uniform definition for “artificial intelligence,” which is used in this bill.
 - e) SB 1120 (Becker), Chapter 897, Statutes of 2024, requires a health plan or insurer that uses an AI, algorithm, or other software tool, and subcontractors of those health plans and insurers as specified, to, among other requirements, ensure the use of the AI, algorithm, or other software tool does not discriminate, directly or indirectly, against enrollees in violation of state or federal law.
- 7) **AMENDMENTS.** The author and committee have agreed to minor technical and clarifying amendments. Amendments will:
- a) Clarify, consistent with the author’s intent, that the bill applies to systems used to support clinical decisionmaking or health care resource allocation;
 - b) Incorporate, by reference to the Government Code § 11546.45.5, the definition of AI, rather than recreating the definition in this bill

This bill also lacks consistency in the application of its requirements—one requirement applies to specified health care entities, while the other two requirements apply to “deployers,” which is defined more generally. The author is encouraged to address this issue in subsequent amendments to the bill.

8) **POLICY COMMENTS.**

- a) **Beneficial Protection for California.** As noted above, bias is a known and real concern in the application of AI in health care. This bill appears to generally align with the intent of federal regulations prohibiting bias in the application of patient care decision support tools, although it is different in key ways, as discussed in g) above. As Section 1557 has been the subject of numerous regulatory changes based on changing political winds in Washington, D.C., codifying similar requirements in state law could provide a backstop to ensure bias in health care AI systems is addressed.
- b) **Increasing Compliance Requirements May Pose Challenges for Some Providers.** On the other hand, there is a risk that additional state compliance requirements for deploying AI could increase the division between the “haves and have-nots.” The complexity, cost,

and numerous implementation considerations of deploying AI in the health care system discourages adoption. Safety net patients may be harmed from bias, and may also be harmed from failure to deploy beneficial technology. It is unclear how these factors should be weighed against each other; indeed, it may be impossible to know the relative risks in order to make fully informed decisions. But given the challenges they already face, it should be acknowledged that implementing additional state licensing requirements may further discourage adoption by providers who have less administrative and technical capacity. Complimentary efforts, such as compliance support for smaller or safety net providers, or offering opportunities for technical assistance prior to a citation, could be helpful to reduce the potential impact of additional requirements.

- c) **Transparency from Developers.** Deployers would be better equipped to comply with this bill if all developers provided standardized information about their AI models, such as the information covered in the draft “applied model card.” Currently, developers are not required to provide information that deployers may need in order to understand, monitor, and mitigate the risk of bias. Smaller, less resourced providers may also be disadvantaged in their ability to seek favorable contract terms that provide the needed level of transparency. The author is encouraged to consider specifying required disclosures from developers to deployers to support compliance with this bill.

REGISTERED SUPPORT / OPPOSITION:

Support

California Hospital Association
California Medical Association (CMA)
Kaiser Permanente
Oakland Privacy

Opposition

None on file

Analysis Prepared by: Lisa Murawski / HEALTH / (916) 319-2097