

Ethical Affective Computing in the Workplace: Measuring Emotion Without Surveillance

Executive Summary

The most credible interpretation of affective computing's original research program is not “read workers’ feelings from cameras and microphones,” but rather “build systems that are sensitive to human affect while recognizing that emotion is contextual, uncertain, and ethically fraught.” The foundational MIT Press book on affective computing explicitly included moral and social questions alongside technical ones, and early work by **Rosalind Picard**[\[1\]](#) and colleagues at the **MIT Media Lab**[\[2\]](#) repeatedly treated ground truth as difficult: stress studies combined physiology with self-report, contextual conditions, and third-party coding, and some prototypes intentionally transmitted affect-related signals to another human without machine interpretation at all. That history matters. It suggests that an ethical workplace implementation should help people reflect, communicate, and recover—not extract “true emotions” from hidden observation. [\[3\]](#)

Emotional intelligence theory points in the same direction. The most influential workplace literature presents emotional intelligence as a leadership, team, and cultural capability: self-awareness, regulation, empathy, group norms, and emotional culture. It is associated with better commitment, job satisfaction, citizenship behavior, performance, and lower job stress, but those relationships do not imply that employers should continuously measure emotions. They imply that organizations should create conditions in which people can notice, discuss, and regulate emotion safely. In other words, the application target is organizational practice, not covert sensing. [\[4\]](#)

For a product team, the practical conclusion is straightforward. A defensible workplace product should be self-report first, opt-in contextual signals second, aggregated by default, and explicitly barred from hiring, firing, promotion, compensation, or disciplinary decisions. It should not collect raw biometric data, run camera- or microphone-based emotion inference, or rank employees by “emotional productivity.” That design is now both ethically stronger and strategically safer: the scientific basis for inferring emotion from faces remains weak outside constrained settings; electronic monitoring tends to raise stress and lower job satisfaction with little performance upside; the EU AI Act now bans emotion recognition in workplaces except narrow medical or safety cases; and U.S. regulators have warned that workplace AI and biometric practices can trigger discrimination and unfairness risk. [\[5\]](#)

The report therefore recommends a privacy-preserving architecture built around brief mood check-ins, optional workload and context tags, team-level climate summaries with small-group suppression and differential privacy options, person-centered trend models, calibration and fairness audits, and intervention logic that privileges support over

prediction. The analytic stance should be: model trajectories and uncertainty, not hidden inner states; optimize for trust, consent, and actionability, not maximal data extraction. [6]

Conceptual Foundations

The original affective computing program was broader than today's commercial "emotion AI" shorthand. The MIT Press formulation described computing that relates to, arises from, or deliberately influences emotion, and it treated the moral and social implications as part of the field rather than as an afterthought. Primary MIT work from the early 2000s did show promising performance for physiology-based stress or affect classification in narrow tasks, but those demonstrations were small, highly contextual, and methodologically careful about labels. In one naturalistic stress study, the team noted that "true" stress is hard to ascertain and triangulated among self-report, driving condition, and third-party coding. That is a much more modest and scientifically serious posture than the claim that workplace systems can read emotions passively and universally from observable behavior. [7]

That same research tradition also contains a design clue that product teams often ignore. In Picard's "TouchPhone" example, the system transmitted pressure as an ambient color signal to a conversational partner and explicitly did **not** interpret the signal for them. In adjacent MIT work, researchers also discussed tradeoffs between active, passive, and hybrid sensors. Read analytically, those examples point toward an ethical framework for workplace products: prefer employee-controlled signals over covert sensing; prefer augmentation of reflection or communication over automated judgment; and prefer context-rich, person-specific interpretation over generalized, one-shot classification. [8]

Emotional intelligence theory reinforces that distinction. The ability model developed by Salovey and Mayer framed emotional intelligence as skills in perceiving, using, understanding, and regulating emotion. **Daniel Goleman**[9] translated that into management language in Harvard Business Review, arguing that emotional intelligence is central to effective leadership; later HBR work extended the idea to group emotional intelligence and emotional culture. This literature is about capabilities and norms: whether teams acknowledge feelings, how leaders respond, how conflict is metabolized, and whether emotional expression is constructive. It is not a warrant for hidden measurement. [10]

That distinction is particularly important in the workplace. A recent meta-analysis found that emotional intelligence, across ability, self-report, and mixed streams, is positively related to organizational commitment, citizenship behavior, job satisfaction, and job performance, and negatively related to job stress. Yet those findings support training, reflection, and better management practices—not surveillance. One useful HBR illustration is the voluntary "emotion button" ritual at Ubiquity Retirement + Savings, where employees chose from simple affective states at the end of the day. For this whitepaper, that kind of voluntary, low-resolution, employee-legible signal is a much better translational model than covert inference from faces, voices, or keystrokes. [11]

Risks, Harms, and Regulation

The strongest technical objection to invasive emotion tracking is epistemic: the science does not support the confidence the products often claim. The widely cited 2019 review by Barrett and colleagues argued that a facial expression can only justify emotion inference if it is reliable, specific, generalizable, and valid, and concluded that the evidence for universal, context-free facial signatures of emotion is lacking. The review specifically found that facial movements commonly treated as diagnostic of particular emotions are not universally diagnostic across persons, contexts, and cultures, and warned that this gap matters precisely because technology companies are trying to “read emotions” from faces. For workplace products, that means face-based emotion inference is not merely risky; it is weakly grounded. [12]

The second objection is organizational and psychological. Employee surveillance changes behavior even when the measurement is imperfect. A 2022 meta-analysis of 70 independent samples found that electronic monitoring slightly decreases job satisfaction and slightly increases stress, while showing no relationship with performance and a small positive relationship with counterproductive work behavior. At the team level, this collides with psychological safety, which **Amy Edmondson**[13] defined as a shared belief that the team is safe for interpersonal risk-taking. More recent research suggests that clarity about monitoring practices can moderate some harm, but that does not rescue invasive emotion tracking; it simply means transparent, bounded, employee-understandable systems are less harmful than opaque ones. [14]

The third objection is consent. In the employment relationship, “opt-in” is often not genuinely free because of role power, evaluative pressure, and fear of detriment. The European Data Protection Board’s consent guidance states that valid consent must be freely given, specific, informed, and unambiguous; that people must have real choice and control; that there must be no negative consequences for refusing or withdrawing consent; and that purposes must be separated granularly rather than bundled. The guidance also explicitly recognizes power imbalance, including in employment-like settings. These principles are exactly the opposite of the usual surveillance playbook of broad bundled notice, default inclusion, and vague future use. [15]

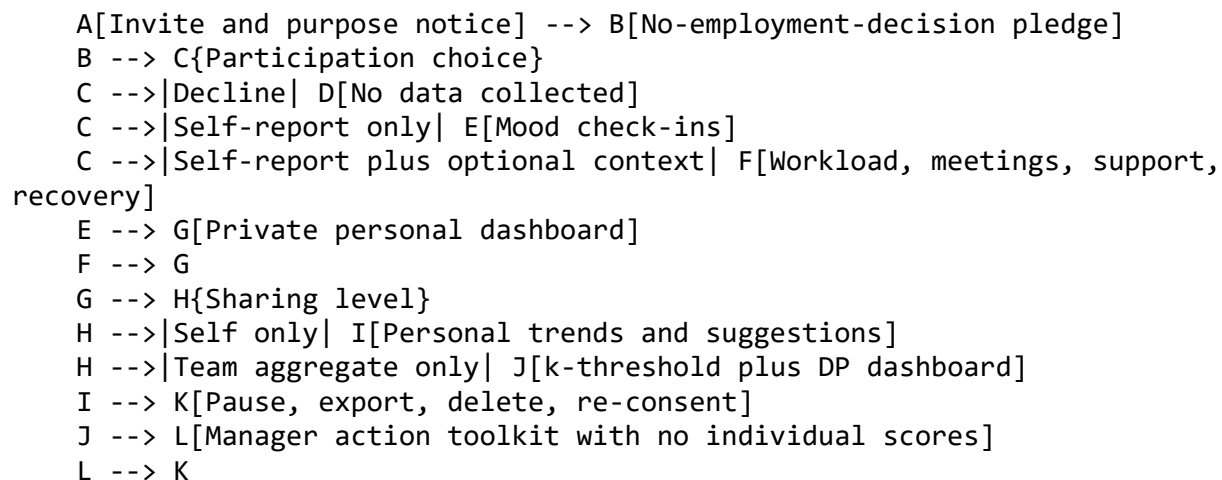
Regulation is also moving against invasive workplace emotion AI. The **European Commission**[16] states that the AI Act prohibits emotion recognition in workplaces and education institutions except for narrow medical or safety reasons; those prohibitions became effective in February 2025. The broader EU direction is similar in labor law: a European Parliament briefing notes that the platform-work regime forbids processing certain data about workers’ emotional or psychological state by automated systems. Even where a product tries to avoid the AI Act’s prohibited category, the GDPR framework remains relevant through requirements of lawfulness, fairness, transparency, purpose limitation, minimization, and accuracy. [17]

In the United States, the legal picture is patchier but still unfavorable to invasive designs. The **Equal Employment Opportunity Commission**[18] warns that AI and other technology can violate anti-discrimination law when used in employment decisions, including monitoring workers' activity, performance, or location, assessing productivity, and deciding whom to promote or fire; its examples include speech-pattern analysis disadvantaging people with disabilities and facial recognition that is less accurate for darker skin tones. The **Federal Trade Commission**[19] has separately warned that concealed invasive surveillance, privacy-invasive default settings, inaccurate biometric technologies, and unreasonable data security for biometric information can be unfair or deceptive under the FTC Act. Meanwhile, the **Illinois General Assembly**[20] defines biometric identifiers to include voiceprints and scans of face geometry, requires written notice of purpose and duration plus a written release, and requires a public retention schedule with destruction when the purpose is satisfied or within three years of the person's last interaction, whichever comes first. [21]

Product Translation for Self-Reported and Opt-In Signals

A privacy-preserving workplace product should therefore be framed as an employee well-being and team-climate tool, not an employer surveillance tool. The evidence base for self-report in natural contexts is relatively strong: ecological momentary assessment captures repeated in-the-moment experience in real environments, and a large meta-analysis found that EMA studies averaged about six assessments per day for seven days with roughly 79% compliance. For workplace deployment, however, ethics and burden argue for a much lighter cadence than research-grade EMA, especially in a setting where participation may feel consequential. The target should be proportionate sampling that workers can understand, skip, pause, and revoke without any penalty. [22]

flowchart LR



The specification below is a recommended default template. It is intentionally conservative: self-report first, optional context second, minimal system metadata, team-level aggregation only under minimum group thresholds, and explicit exclusion of camera, microphone, and covert behavioral monitoring. Those defaults are consistent with the

consent principles above, with voluntary mood-capture examples from management practice, and with the weak scientific basis and regulatory risk surrounding biometric emotion inference. [23]

Signal class	Example fields	Collection frequency	Consent and visibility	Privacy and DP option	Retention policy
Core mood check-in	Valence 1–5, energy 1–5, stress/tension 1–5	Once daily on workdays, or 3 times per week if user chooses lighter mode	Separate opt-in; self-visible by default	No noise on self-view; pseudonymous storage	Raw identifiable events 90 days, derived rolling features 12 months
Context tags	Workload, meeting load, autonomy, support, recovery, work location category	Attached to mood check-in or weekly summary	Separate toggle from mood; self-visible; team only in aggregate	k-threshold aggregation only	Same as core mood
Weekly climate pulse	Psychological safety, fairness of workload, manager support, ability to focus	Weekly	Team opt-in; never shown for groups below threshold	Central DP on manager dashboard; suppress small cells	Team aggregates 12–24 months
Optional reflection note	Short free text or voice-to-text reflection	Optional ad hoc	Self-only; not employer-visible	Prefer local-first or encrypted server storage; no model training by default	User-controlled; default delete after 7–30 days
Optional calendar-derived context	Count of meetings, meeting minutes, focus blocks, after-hours work windows	Daily summary	Distinct opt-in; summary only	Process on device where possible; never ingest raw meeting content	Daily summaries 90 days; no raw content retained

Signal class	Example fields	Collection frequency	Consent and visibility	Privacy and DP option	Retention policy
Minimal system metadata	Timestamp, device type, completion latency, app version	Automatic	Included in core consent notice	Used only for reliability monitoring and fraud control	30 days unless incident investigation requires longer
Explicitly excluded data	Camera, webcam, mic, raw voice, facial analysis, keystrokes, screenshots, continuous GPS, covert meeting transcription, always-on wearables	Never	Not collectable under standard product	No exception without separate product review and legal basis	Not stored

The UI/UX consent flow should be stage-gated rather than buried in a single checkbox. A defensible flow has seven elements: a plain-language purpose screen; a prominent statement that no emotional signal will be used for evaluative employment decisions; separate toggles for each signal source; a preview showing exactly who can see what; just-in-time notices when a source is first activated or a purpose changes; one-tap pause/skip/export/delete controls; and periodic re-consent, especially after any model or purpose change. There should be no pre-ticked boxes, no degraded service for declining optional signals, and withdrawal should be possible in the same interface used to enroll.

[15]

For anonymization and reporting, the right pattern is pseudonymization plus strict aggregation rather than magical claims of “anonymous individual data.” Keep the identity key in a separate access domain from the analytical store; require a minimum cohort size before any manager-visible metric appears; suppress or combine small groups; and publish the aggregation rules. Differential privacy is useful, but only in the right layer: apply it to group dashboards and recurring organizational reports, not to an employee’s own private dashboard or to any workflow that triggers person-level intervention. MIT Press’s Harvard Data Science Review and the differential privacy literature both emphasize the privacy-accuracy tradeoff, and repeated releases require explicit privacy-budget accounting. In practice, central DP with fixed release schedules works better for managerial dashboards than local DP for operational intervention logic. [24]

On storage and retention, a pragmatic policy is shorter than the legal ceiling. For this whitepaper, the recommended default is 90 days for identifiable raw events, 12 months for derived trend features and team aggregates, user-controlled deletion for personal notes, automatic inactivity-based deletion, and auditable destruction logs. If an organization ever touches biometric data, it should assume BIPA-like duties at minimum: explicit notice, stated purpose and term, written release, public retention schedule, and secure storage. Even without biometrics, the **NIST** privacy model and FTC guidance point toward the same operational norms: govern purpose, communicate clearly, control access, protect data, and document the safeguards. [\[25\]](#)

Modeling, Validation, and Visualization

A responsible analytics stack should start with description, not prediction. The first layer is person-centered trend analysis: rolling means, exponentially weighted moving averages, within-person deviations from baseline, and simple change-point flags on mood, energy, and stress. The second layer is repeated-measures modeling, usually mixed-effects models with random intercepts and slopes so that people are compared primarily to themselves rather than to an arbitrary workforce-wide “normal.” The third layer is Bayesian hierarchical or state-space modeling for sparse, irregular, or noisy data, because those models directly encode uncertainty and partial pooling. Only after those layers should a team consider explainable machine learning, and then only if it materially improves calibration or actionable lift over the simpler baselines. [\[26\]](#)

Feature engineering should reflect the best emotion-dynamics literature, not generic “engagement scores.” The key families are level, slope, variability, instability, and inertia. Meta-analytic and review work on emotion dynamics shows that not just average affect, but also fluctuations over time, matter for well-being; higher emotional inertia in particular has been linked to lower well-being, and workplace-specific work has connected inertia of negative emotions at work to inflexible dynamics. In operational terms, the most useful engineered features are likely to be: rolling mean valence and energy; stress acceleration over the last one to two weeks; within-person standard deviation; inertia estimated as an autoregressive carry-over coefficient; recovery speed after high-workload days; divergence between mood and workload tags; response latency; and structured missingness. Person-specific normalization has deep roots in earlier MIT organizational sensing work as well, where features were normalized to each individual’s own baseline rather than treated as absolute levels. [\[27\]](#)

Validation should follow the standards used in serious risk modeling, not internal product dashboards alone. There are four distinct tests. Construct validity asks whether the signals behave as expected relative to validated instruments such as PANAS for affect, WHO-5 for mental well-being, a short engagement scale such as UWES, and a non-diagnostic occupational burnout measure. Predictive validity asks whether trends forecast outcomes the organization actually cares about, such as voluntary attrition, self-reported stay intent, or team-level engagement. Temporal validity asks whether the model still performs under forward-chaining cross-validation rather than random splits. Calibration asks whether

predicted risks correspond to observed frequencies; calibration intercepts, slopes, curves, and Brier-like summaries matter more than AUC alone when predictions are used to trigger interventions. [28]

Explainable ML has a legitimate but bounded role. SHAP and related methods are useful for showing which features most influenced a specific prediction and for summarizing global model structure across many local explanations. But explanation does not cure bad measurement. If the upstream features come from invalid proxies, then a beautifully explained model simply explains a weak construct. For this reason, the recommended ML frontier for workplace affective products is not deep multimodal inference; it is constrained, interpretable modeling over self-reported and explicitly authorized context signals, combined with subgroup performance audits, small-sample uncertainty flags, and drift monitoring. A good release criterion is: no model ships unless it beats a mixed-effects baseline on calibration, preserves subgroup stability, and can be explained in language an employee and ethics board could both understand. [29]

Visualization should be designed for reflection and intervention, not ranking. The employee view should use simple trend lines, event annotations, confidence bands, and comparison to the individual's own recent baseline. The manager view should show only cohort metrics: directional trends, uncertainty intervals, and action prompts tied to systems factors such as meeting load, support, or workload fairness. There should be no leaderboard, no heat map of named employees, and no traffic-light status tied to disciplinary workflows. Visual grammar should make uncertainty and aggregation visible, not hidden. [30]

KPI Mapping and Intervention Logic

KPI mapping should be treated as decision support for supportive intervention, not as hidden adjudication. The thresholds below are deliberately illustrative and should be tuned locally using forward validation, calibration curves, and explicit cost tradeoffs. They are also not diagnostic cutoffs: **World Health Organization**[31] defines burnout as an occupational phenomenon rather than a medical condition, and any person-level use should be employee-facing unless the worker explicitly asks to share it. For manager-visible workflows, specificity and aggregation matter more than early-warning sensitivity, because false positives can damage trust and psychological safety. [32]

KPI	Recommended signals	Preferred model	Illustrative threshold	Sensitivity / specificity posture	Recommended intervention
Burnout risk screening	14-day stress EWMA, energy decline, workload tags, negative-affect inertia, recovery lag	Bayesian hierarchical change-point or mixed-effects	Stress > +1 within-person SD and energy < -1 SD for 2 weeks,	Self-view: favor sensitivity (catch early drift). Manager/tea	Self-directed workload review, recovery planning, optional

KPI	Recommended signals	Preferred model	Illustrative threshold	Sensitivity / specificity posture	Recommended intervention
		logistic screen	plus high-workload tag on at least half of check-ins	m view: favor specificity and aggregate only	wellbeing coach, manager-supported workload rebalance if employee chooses to share
Retention risk	Stay-intent pulse, support/recognition tags, sustained negative slope in engagement, manager-support pulse	Mixed-effects survival or Bayesian hazard model	Stay-intent or engagement drops >0.5 SD from baseline for 4 consecutive weeks	Favor specificity; false positives can feel manipulative	Offer team-wide stay interviews, career conversations, role clarity review, recognition and growth-path interventions
Engagement	Positive affect, vigor items, UWES short pulse, team psychological safety	Multilevel growth model	Team mean < 3.6/5 or slope worse than -0.15 per week for 3 weeks	Balanced, but team-level only	Team retrospective, manager coaching, recognition rituals, reduction of known frictions
Productivity friction	Meeting-load tag, focus-time summary, mood-workload divergence, completion delays, after-hours work spike	Interpretable GAM or gradient boosting with SHAP	Friction score >0.70 and calibrated over 2 weeks	Favor specificity; do not create false process alarms	Meeting cap, asynchronous defaults, focus blocks, workflow redesign, staffing review

The most important implementation detail is not the threshold itself but the visibility rule. Burnout-risk screening should normally remain private to the employee; retention-risk

flags should never become secret “flight risk” labels for managers; engagement and productivity-friction KPIs should usually be surfaced at team level; and any person-level intervention should be framed as an offer of support, not a hidden conclusion. In practice, many organizations will obtain better utility by triggering supportive actions team-wide whenever a cohort crosses a threshold, rather than by singling out individuals from probabilistic models. [33]

Experimental Design and Governance

A/B testing in this domain should optimize trust, burden, and usefulness—not maximize extraction. The ethical baseline is simple: randomize UX choices and supportive interventions, never the presence of surveillance. All arms should remain privacy-preserving; the experiment should not test camera-based inference against self-report, or hidden monitoring against voluntary participation. Consent copy, prompt cadence, feedback design, and action-toolkit quality are appropriate variables; secret observability is not. [34]

```

timeline
  title Privacy-first rollout sequence
  Governance and review : legal basis, worker consultation, data map,
  prohibited-use red lines
  Pilot : self-view only, no manager access, burden and trust measurement
  Validation : calibration, subgroup error, deletion and opt-out analysis
  Controlled rollout : team dashboards with small-cell suppression and DP
  Ongoing audit : quarterly model review, drift tests, re-consent after
  changes
```

Approximate sample sizes below assume a two-sided alpha of 0.05 and 80% power, using standard normal approximations for simple A/B comparisons; inflate by 10% to 20% for attrition, and by the design effect for cluster randomization. In this domain, pre-registration, worker representation, and transparent exclusion criteria are especially important because trust effects can change both adoption and data quality. [35]

Experiment	Unit	Hypothesis	Primary metrics	Approximate sample size	Typical duration	Ethical safeguards
Privacy-forward onboarding vs conventional onboarding	Individual	Layered explanation, granular toggles, and a clear non-evaluative-use pledge increase informed	Opt-in rate, comprehension on quiz, trust score, deletion rate	~1,100 per arm for 55% to 61% opt-in; ~400 per arm for 0.2 SD trust improvement	2–4 weeks	No dark patterns, one-click decline, same core service for nonparticipants

Experiment	Unit	Hypothesis	Primary metrics	Approximate sample size	Typical duration	Ethical safeguards
		opt-in and trust				
Adaptive cadence vs fixed daily prompt	Individual	User-chosen or adaptive cadence improves week-4 completion without reducing validity	Completion rate, burden score, missingness bias, correlation with weekly well-being scale	~900 per arm for 70% to 76% completion	4 weeks	Quiet hours, hard cap on prompts, pause/skip anytime
Dashboard only vs dashboard plus self-reflection nudges	Individual	Personalized, non-diagnostic guidance improves well-being and reduces elevated risk flags	WHO-5 change, high-stress day count, helpfulness, opt-out	~400 per arm for 0.2 SD well-being gain; ~1,450 per arm for 20% to 16% elevated-risk reduction	8–12 weeks	No clinical language, no manager visibility, human support links always available
Team dashboard vs team dashboard plus manager action toolkit	Team cluster	Actionability, not measurement alone, improves engagement and reduces friction	Team engagement, psych safety, meeting load, voluntary attrition	~675 employees per arm, or about 27 teams per arm, assuming 25 people/team and ICC ≈ 0.03 for a 0.2 SD effect	8–12 weeks	Small-team suppression, no individual drill-down, manager training on non-punitive use

Governance should be embedded before launch, not bolted on afterward. That means a written purpose limitation; a prohibited-use list; data-flow mapping; subgroup fairness and calibration reviews; audit logs for access, deletion, and model changes; and worker

consultation where required or prudent. Every model should have a short model card stating inputs, excluded data sources, intended users, validation sample, calibration status, subgroup checks, refresh cadence, and explicit non-uses. Any material change to signals or downstream use should trigger re-consent. If the system cannot be explained in these terms, it is not ready for workplace deployment. [36]

The decision rule for ethics boards is therefore demanding but clear. Approve only systems that satisfy five conditions simultaneously: valid measurement, genuine voluntariness, strict aggregation and visibility controls, calibrated and explainable analytics, and supportive rather than disciplinary intervention logic. Under those conditions, affective computing in the workplace can become a tool for self-awareness and organizational learning. Without them, it collapses into surveillance dressed as empathy. [37]

[1] [6] [15] [16] [33] [34]

https://www.edpb.europa.eu/sites/default/files/files/file1/edpb_guidelines_202005_consent_en.pdf

https://www.edpb.europa.eu/sites/default/files/files/file1/edpb_guidelines_202005_consent_en.pdf

[2] [27] https://ppw.kuleuven.be/okp/_pdf/Kuppens2017ED.pdf

https://ppw.kuleuven.be/okp/_pdf/Kuppens2017ED.pdf

[3] [7] [37] <https://mitpress.mit.edu/9780262661157/affective-computing/>

<https://mitpress.mit.edu/9780262661157/affective-computing/>

[4] <https://hbr.org/archive-toc/BR0401>

<https://hbr.org/archive-toc/BR0401>

[5] [9] [12] [13] <https://www.psychologicalscience.org/publications/emotional-expressions-reconsidered-challenges-to-inferring-emotion-from-human-facial-movements.html>

<https://www.psychologicalscience.org/publications/emotional-expressions-reconsidered-challenges-to-inferring-emotion-from-human-facial-movements.html>

[8] <https://vismod.media.mit.edu/pub/tech-reports/TR-537.pdf>

<https://vismod.media.mit.edu/pub/tech-reports/TR-537.pdf>

[10] [20]

https://center.uoregon.edu/StartingStrong/uploads/STARTINGSTRONG2016/HANDOUTS/KEY_46201/pub153_SaloveyMayerICP1990_OCR.pdf

https://center.uoregon.edu/StartingStrong/uploads/STARTINGSTRONG2016/HANDOUTS/KEY_46201/pub153_SaloveyMayerICP1990_OCR.pdf

[11]

<https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2022.611348/full>

<https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2022.611348/full>

[14] <https://www.sciencedirect.com/science/article/pii/S2451958822000616>

<https://www.sciencedirect.com/science/article/pii/S2451958822000616>

[17] [19] [36] <https://digital-strategy.ec.europa.eu/en/faqs/navigating-ai-act>

<https://digital-strategy.ec.europa.eu/en/faqs/navigating-ai-act>

[18] [22] <https://pmc.ncbi.nlm.nih.gov/articles/PMC9163273/>

<https://pmc.ncbi.nlm.nih.gov/articles/PMC9163273/>

[21] https://www.eeoc.gov/sites/default/files/2024-04/20240429_What%20is%20the%20EEOCs%20role%20in%20AI.pdf

https://www.eeoc.gov/sites/default/files/2024-04/20240429_What%20is%20the%20EEOCs%20role%20in%20AI.pdf

[23] <https://hbr.org/2016/01/manage-your-emotional-culture>

<https://hbr.org/2016/01/manage-your-emotional-culture>

[24] <https://hdsr.mitpress.mit.edu/pub/sl9we8gh>

<https://hdsr.mitpress.mit.edu/pub/sl9we8gh>

[25]

<https://www.ilga.gov/Legislation/ILCS/Articles?ActID=3004&ChapterID=57&Print=True>

<https://www.ilga.gov/Legislation/ILCS/Articles?ActID=3004&ChapterID=57&Print=True>

[26] <https://pmc.ncbi.nlm.nih.gov/articles/PMC3655706/>

<https://pmc.ncbi.nlm.nih.gov/articles/PMC3655706/>

[28] <https://www.scienceofbehaviorchange.org/wp-content/uploads/2019/10/PANAS.Watson.1988.pdf>

<https://www.scienceofbehaviorchange.org/wp-content/uploads/2019/10/PANAS.Watson.1988.pdf>

[29] <https://proceedings.neurips.cc/paper/7062-a-unified-approach-to-interpreting-model-predictions.pdf>

<https://proceedings.neurips.cc/paper/7062-a-unified-approach-to-interpreting-model-predictions.pdf>

[30] [31] <https://hdsr.mitpress.mit.edu/pub/aelql9qy>

<https://hdsr.mitpress.mit.edu/pub/aelql9qy>

[32] <https://www.who.int/standards/classifications/frequently-asked-questions/burn-out-an-occupational-phenomenon>

<https://www.who.int/standards/classifications/frequently-asked-questions/burn-out-an-occupational-phenomenon>

[35] <https://pmc.ncbi.nlm.nih.gov/articles/PMC9999286/>

<https://pmc.ncbi.nlm.nih.gov/articles/PMC9999286/>