
READING & WRITING ASSIGNMENTS

Artificial Intelligence

One theme - four different reading options.

In this packet (suggested pairings):

- **AP Computer Science Principles**
The Computer Said It Was You (Facial Recognition)
- **AP Computer Science A**
Rejected by a Robot (AI Resume Screening)
- **Cybersecurity**
Seeing Is No Longer Believing (Deepfakes)
- **Independent Study**
Guilty Until Proven Human (AI Detectors)

Each assignment shares the same spine: a specific story, the broader concept, at least four of your own sources, tiered questions, and a prepared for/against debate.

The Computer Said It Was You

Facial Recognition, AI, and Wrongful Arrest

The Story That Started This

Imagine being arrested in your own driveway, in front of your two young daughters, for a crime you had nothing to do with - because a computer picked your face out of a database.

That happened to **Robert Williams**. In January 2020, Detroit police arrested him for stealing watches, based on a facial-recognition “match” from blurry store security footage. He spent about 30 hours in a crowded, dirty cell before anyone realized: the computer algorithm was wrong. His was the one of the first publicly known cases of a bad face-match leading to a wrongful arrest in the United States. It would not be the last.

In June 2024, the city settled his lawsuit: \$300,000, plus what the ACLU called the strongest police facial-recognition rules in the country - including a flat ban on arresting someone based solely on a face-match.

The Bigger Idea

Facial recognition is an AI system: you train a model on millions of faces, and it learns to turn a new photo into a set of numbers and find the closest match in a database. Like any model, it doesn't return truth - it returns a *probability*, a best guess. And its guesses are not equally good for everyone. Study after study has found these systems misidentify women and people of color at much higher rates, because of what was - and wasn't - in the data they learned from.

That's the core AI lesson hiding inside this story: a model is only as fair as its training data, and its output is a suggestion, not a verdict. The danger isn't just that the algorithm is imperfect - it's that people treat the output as the gospel-truth.

Why This Is Tricky

Police make a fair point: facial recognition can generate real leads and help solve real crimes, and nobody's asking officers to ignore useful tools. The question is where the line sits: is a face-match a starting clue that still needs human investigation - or is it being treated as the answer? When the cost of the model being wrong is an innocent person in a jail cell, how sure is sure enough? Does it depend on the situation?

Step 1 - Read

The first two tell Robert Williams' story and the settlement; the third zooms out to how states are trying to regulate this, so you see the range of positions.

- **The case (ACLU)** - Williams v. City of Detroit.
<https://www.aclu.org/cases/williams-v-city-of-detroit-face-recognition-false-arrest>
- **The settlement (Michigan Public)** - "It didn't make sense at all."
<https://www.michiganpublic.org/criminal-justice-legal-system/2024-06-28/it-didnt-make-sense-at-all-wrongful-facial-recognition-arrest-leads-to-landmark-settlement>
- **The bigger picture (Stateline)** - facial recognition is getting state-by-state guardrails.
<https://stateline.org/2025/02/04/facial-recognition-in-policing-is-getting-state-by-state-guardrails/>

Step 2 - Research on Your Own

Find **at least four more** reputable sources you dig up yourself - news reporting, the ACLU, research on facial-recognition accuracy (look up the "Gender Shades" study), or a police department's own policy. Random social-media posts don't count. You'll cite these in your paper and lean on them in the debate.

Step 3 - Written Response

Turn in a typed paper. Answer in complete sentences. Point values are shown; total is **25 points**.

Comprehension

1. In your own words, how does a facial-recognition system decide two faces are a "match"? (1 pts)
2. What went wrong in Robert Williams' case, and what did the 2024 settlement change? (1 pts)

Analysis

1. How has facial recognition benefited society? Give at least 2 examples. (2 pts)
2. Who benefits from police facial recognition, and who bears the risk? Name at least two groups on each side and explain each. (3 pts)
3. These systems misidentify women and people of color more often. Explain how a model's *training data* can produce that kind of bias. (3 pts)
4. A face-match is a probability, not a fact. Explain the difference between using it as an investigative *lead* versus as *proof* - and why that distinction can decide whether an innocent person goes to jail. (5 pts)

Position

Take a clear position on the resolution below and argue it - this is the required part.

1. Where do you stand on the resolution? Argue your side using evidence from at least two of the

sources you found, then fairly state the single strongest argument against your view and respond to it. (10 pts)

(OPTIONAL) Prepare Both Sides

Your teacher may ask you to be ready to argue *either* side. If so, also prepare a short case for each:

1. The strongest case **FOR** the resolution, backed by evidence you found yourself.
2. The strongest case **AGAINST** the resolution, backed by evidence you found yourself.

Step 4 - (OPTIONAL) Class Debate

Resolved: *"Police should be banned from using facial recognition to identify criminal suspects."*

If your teacher runs the debate, you'll defend a position on this resolution, and you may be **assigned a side at random** - which is why the optional prep above asks you to be ready for both. You won't have your paper or notes with you, so know your evidence well enough to argue from memory.

Two rules make the debate work:

1. **Evidence only.** Every claim points to a specific source for evidence, not a feeling.
2. **Steelman first.** Before you rebut the other side, restate their strongest argument fairly enough that they'd nod along. Then tell your tale.

One More Thing

You're going to hear all about the latest advances in AI and the newest things it can do. This is a story about what happens when people forget what it *can't* do - and trust a confident number over a real human being. Remembering that difference may end up being one of the most important takeaways from this class.

Rejected by a Robot

AI Resume Screening and Algorithmic Hiring

The Story That Started This

This assignment started with an article that made me angry. When Stanford's medical center rolled out its first COVID-19 vaccines, the algorithm it used to decide who went first left out almost all of its front-line resident doctors - the very people treating COVID patients - while prioritizing higher-ups and staff working from home. When the residents protested, leadership blamed "a very complex algorithm." My wife is a front-line worker, so I took it personally. And it got me noticing how often an algorithm quietly decides something big about a person's life - and how often, when it goes wrong, nobody wants to own the decision. Hiring turned out to be full of examples.

Derek Mobley applied for more than 100 jobs. He was rejected from every single one - sometimes within minutes of hitting submit, at hours when no human was awake to read anything. He's Black, he's over 40, and he has a disability. The common thread across all those companies wasn't a person. It was a piece of software: **Workday**, whose AI screening tools help sort applicants for thousands of employers.

The Bigger Idea

An AI resume screener is a model trained to predict who's a "good" hire, usually by learning from a company's past hiring data. And that's the trap: if the people a company hired in the past skewed young, or male, or from a particular background, the model learns those patterns as the definition of "good" - and quietly reproduces them at scale. Amazon learned this the hard way years ago: it built an experimental hiring AI, discovered it was penalizing resumes that included the word "women's," and scrapped it.

Keep in mind: a model has no understanding of what's *fair*, only what's *frequent*. Train it on biased history and it will automate that bias, faster and more consistently than any human ever could - while looking perfectly objective because "the computer decided."

Why This Is Tricky

Companies aren't using these tools to be cruel (hopefully). A big employer can get tens of thousands of applications; no human can read them all, and manual screening by a human can introduce its own biases. AI screening is faster, cheaper, and at least consistent. So the question isn't "humans good, robots bad." It's whether a system that automates past patterns can ever be fair to the people those patterns left out - and who's accountable when an algorithm says no.

Step 1 - Read

The first two cover the Workday case and what it means for hiring; the third is an analysis that includes how companies defend these tools.

- **The lawsuit (CNN Business)** - Workday's hiring tech accused of screening out people over 40.
<https://www.cnn.com/2025/05/22/tech/workday-ai-hiring-discrimination-lawsuit>
- **What it means (SHRM)** - "The Workday AI Lawsuit Is a Wake-Up Call for HR."
<https://www.shrm.org/topics-tools/news/technology/workday-ai-lawsuit-wake-up-call-hr>
- **The bias problem (Forbes)** - What the Workday Lawsuit Reveals About AI Bias - And How Firms Respond.
<https://www.forbes.com/sites/janicegassam/2025/06/23/what-the-workday-lawsuit-reveals-about-ai-bias-and-how-to-prevent-it/>

Step 2 - Research on Your Own

Find **at least four more** reputable sources you dig up yourself - news reporting, the EEOC, a law like New York City's hiring-algorithm audit rule, or the story of Amazon's scrapped hiring AI. Random social-media posts don't count. You'll cite these in your paper and lean on them in the debate.

Step 3 - Written Response

Turn in a typed paper. Answer in complete sentences. Point values are shown; total is **25 points**.

Comprehension

1. How does an AI resume screener typically "learn" who counts as a good candidate? (1 pts)
2. What is Derek Mobley alleging about Workday, and what did the court allow in 2025? (1 pts)

Analysis

1. Who benefits from algorithmic hiring, and who bears the risk? Name at least two groups on each side and explain each. (2 pts)
2. Explain, in technical terms, how training a model on a company's past hires can turn historical bias into automated bias - even with no one intending it. (3 pts)
3. When an AI hiring tool makes a harmful call, who should be accountable - the company that built it, the company that used it, the people whose data trained it, no one at all, or some other option? Take a position and support it with evidence. (3 pts)
4. Workday says its tool only assists and that humans stay in control. Does "a human is technically in the loop" solve the fairness problem? Argue your view with evidence. (5 pts)

Position

Take a clear position on the resolution below and argue it - this is the required part.

1. Where do you stand on the resolution? Argue your side using evidence from at least two of the sources you found, then fairly state the single strongest argument against your view and respond to it. (10 pts)

(OPTIONAL) Prepare Both Sides

Your teacher may ask you to be ready to argue *either* side. If so, also prepare a short case for each:

1. The strongest case **FOR** the resolution, backed by evidence you found yourself.
2. The strongest case **AGAINST** the resolution, backed by evidence you found yourself.

Step 4 - (OPTIONAL) Class Debate

Resolved: *"Companies should be required to disclose and independently audit any AI that screens job applicants."*

If your teacher runs the debate, you'll defend a position on this resolution, and you may be **assigned a side at random** - which is why the optional prep above asks you to be ready for both. You won't have your paper or notes with you, so know your evidence well enough to argue from memory.

Two rules make the debate work:

1. **Evidence only.** Every claim points to a specific source for evidence, not a feeling.
2. **Steelman first.** Before you rebut the other side, restate their strongest argument fairly enough that they'd nod along. Then tell your tale.

One More Thing

You are a couple of years from submitting resumes into exactly these systems. And someday soon, you might be the one writing the screening code. Both roles come with a stake in getting this right - so it's worth wrestling with now, while it's someone else's job on the line.

Seeing Is No Longer Believing

Deepfakes as an Attack Vector

The Story That Started This

A finance worker at the engineering firm **Arup** got an email from the company's CFO asking for a confidential transfer. It smelled like a scam, and the worker almost ignored it. So the fraudsters set up a video call. On the call were the CFO and several familiar colleagues - faces the worker recognized, voices that sounded right. Reassured, the worker made **15 transfers totaling \$25 million** to accounts in Hong Kong.

Every single person on that video call was a deepfake. The attackers had built AI-generated recreations of the CFO and coworkers from real footage of company meetings posted online. The only real human on the call was the victim. Hong Kong police reported the attack in early 2024; Arup confirmed it was the target a few months later. This is the part of AI that quietly overturns something humans have always taken for granted - that seeing is believing.

The Bigger Idea

A deepfake uses AI - typically a model trained on images or audio of a real person - to generate convincing fake video or voice of that person saying and doing things they never did. For your whole life, seeing or hearing someone was decent proof they were real. Deepfakes quietly break that assumption, and security is built on assumptions about what counts as proof of identity.

Deepfakes supercharge classic attacks: phishing becomes a video call from your "boss"; the grandparent scam becomes your grandchild's actual cloned voice begging for help; disinformation becomes a politician "caught on tape" saying something they never said. The "cyber exploit" has evolved from a bug in a computer to a bug in human trust - and the patch has to be new habits, not just new software.

Why This Is Tricky

The same technology has genuinely good uses: film and dubbing, restoring a voice for someone who lost theirs, satire and art that have always been protected speech. So "ban deepfakes" is harder than it sounds - you'd be banning a tool, not a harm. That's why the debate below targets a narrower question: making a deepfake of a *real person without their consent*. Even there, you'll run straight into free-speech and enforcement problems.

Step 1 - Read

The first two report the Arup attack and how it worked; the third is a security breakdown of the techniques, useful for your own analysis.

- **The attack (CNN)** - Arup revealed as victim of a \$25M deepfake scam.
<https://www.cnn.com/2024/05/16/tech/arup-deepfake-scam-loss-hong-kong-intl-hnk>
- **The first report (CNN)** - finance worker pays out after a call with a deepfake “CFO.”
<https://www.cnn.com/2024/02/04/asia/deepfake-cfo-scam-hong-kong-intl-hnk>
- **How it worked (Eftsure)** - a security walkthrough of the deepfake scam.
<https://www.eftsure.com/blog/cyber-crime/finance-worker-loses-39-million-to-deepfake/>

Step 2 - Research on Your Own

Find **at least four more** reputable sources you dig up yourself - news reporting, a cybersecurity firm’s guidance on detecting deepfakes, a voice-cloning scam case, or a law regulating deepfakes. Random social-media posts don’t count. You’ll cite these in your paper and lean on them in the debate.

Step 3 - Written Response

Turn in a typed paper. Answer in complete sentences. Point values are shown; total is **25 points**.

Comprehension

1. In your own words, what is a deepfake, and how did attackers build the fake video call at Arup? (2 pts)
2. Why did the deepfake succeed where the initial email alone failed? (2 pts)

Analysis

1. Who is put at risk as deepfakes spread, and who benefits from the underlying technology? Name at least two groups on each side. (3 pts)
2. Deepfakes attack “human trust” rather than a machine. What concrete verification steps could have stopped the Arup transfer - and why do technical defenses alone fall short? (3 pts)
3. The same tech enables both a \$25M fraud and legitimate film work. Explain why that dual-use nature makes deepfakes so hard to regulate. (5 pts)

Position

Take a clear position on the resolution below and argue it - this is the required part.

1. Where do you stand on the resolution? Argue your side using evidence from at least two of the sources you found, then fairly state the single strongest argument against your view and respond to it. (10 pts)

(OPTIONAL) Prepare Both Sides

Your teacher may ask you to be ready to argue *either* side. If so, also prepare a short case for each:

1. The strongest case **FOR** the resolution, backed by evidence you found yourself.
2. The strongest case **AGAINST** the resolution, backed by evidence you found yourself.

Step 4 - (OPTIONAL) Class Debate

Resolved: *“Creating a deepfake of a real person without their consent should be a crime.”*

If your teacher runs the debate, you’ll defend a position on this resolution, and you may be **assigned a side at random** - which is why the optional prep above asks you to be ready for both. You won’t have your paper or notes with you, so know your evidence well enough to argue from memory.

Two rules make the debate work:

1. **Evidence only.** Every claim points to a specific source for evidence, not a feeling.
2. **Steelman first.** Before you rebut the other side, restate their strongest argument fairly enough that they’d nod along. Then tell your tale.

One More Thing

So much of modern life runs over the internet, and it all rests on a deceptively simple question: how do you prove someone is who they say they are? Deepfakes just made that question a lot harder. The people who understand this technology are the ones who can protect their communities from it - and give lawmakers the nuance to write rules that actually work. That could be you.

Guilty Until Proven Human

AI Writing Detectors and the Cost of a False Positive

The Story That Started This

I'm a graduate student myself, working toward a Master's in Computer Science - so I see this from both sides of the desk. Most of my professors hand out a syllabus spelling out what AI use is allowed (though, more often than not, the policy is no AI use). And Reddit is full of students asking the same anxious question: what do you do when you get flagged for using AI on an essay or a programming assignment? But here's the part that should worry all of them: some of those flagged students never touched AI at all.

Picture writing an essay entirely yourself - your own late nights, your own words - and then being called into an office and accused of cheating, because a piece of software decided your writing "looks like AI." You can't prove a negative. The tool won't show its work. And the burden is suddenly on you to demonstrate that your own thoughts are yours.

This is happening to real students. As schools rushed to adopt AI-detection tools like Turnitin's, reports piled up of students falsely accused. Then researchers put numbers on it. A Stanford-led study found that popular detectors flagged writing by **non-native English speakers** as AI-generated at wildly high rates - one analysis put the false-positive rate for that group around 61%, roughly five times higher than for native speakers. Several universities, including Vanderbilt, looked at the evidence and simply **turned the detectors off**. The detector's job was to catch cheating. Instead it was punishing students for writing in plainer English.

The Bigger Idea

An AI-text detector is itself an AI model that estimates the probability that text was machine-generated. It doesn't "know" anything - it looks for statistical patterns: predictable word choices, uniform sentence structure, low "surprise." The problem is that plenty of humans write that way, especially people still learning the language who have been taught to write a certain way. To the detector, competence and formulaic writing can look identical to a machine.

Here's the hard part: a detector can be very "accurate" overall yet still be catastrophic when actually put in practice. If it's right 98% of the time but the 2% it's wrong about are real students accused of cheating, the cost of each error is enormous and lands on the person least able to fight back. This is a false-positive problem - the same math behind medical tests and spam filters - with a student's reputation as the stake.

Why This Is Tricky

Teachers have a real problem: some students genuinely do hand in AI-written work, and that's not fair to everyone else. Detectors feel like a defense. Turnitin, for its part, argues its detector is accurate and disputes claims of bias. So this isn't "detectors evil, students innocent." It's a question about proof and consequences: when a probabilistic tool can't explain itself, should it ever be enough to discipline a student?

Step 1 - Read

The first two lay out the false-positive problem and why a major university disabled its detector; the third is Turnitin's own defense of its tool - read it critically.

- **Why we turned it off (Vanderbilt)** - guidance on disabling Turnitin's AI detector.
<https://www.vanderbilt.edu/brightspace/2023/08/16/guidance-on-ai-detection-and-why-were-disabling-turnitins-ai-detector/>
- **The problem (Univ. of San Diego guide)** - false positives and false negatives in AI detectors.
<https://lawlibguides.sandiego.edu/c.php?g=1443311&p=10721367>
- **The vendor's defense (Turnitin)** - Turnitin says its detector shows no significant bias against language learners.
<https://www.turnitin.com/blog/new-research-turnitin-s-ai-detector-shows-no-statistically-significant-bias-against-english-language-learners>

Step 2 - Research on Your Own

Find **at least four more** reputable sources you dig up yourself - the Stanford study on detector bias, university academic-integrity policies, reporting on falsely accused students, or research on false-positive tradeoffs. Random social-media posts don't count. You'll cite these in your paper and lean on them in the debate.

Step 3 - Written Response

Turn in a typed paper. Answer in complete sentences. Point values are shown; total is **25 points**.

Comprehension

1. How does an AI-writing detector decide that text was probably machine-generated? (2 pts)
2. What did the research find about detectors and non-native English writers, and how did some universities respond? (2 pts)

Analysis

1. Who benefits from AI detectors, and who bears the risk? Name at least two groups on each side and explain each. (3 pts)
2. Explain why a detector that is “accurate” overall can still be deeply unfair - use the idea of a false positive and who pays its cost. (3 pts)
3. The detector outputs a probability but can’t explain itself. Should an unexplainable score ever be enough to discipline a student? What if additional evidence was included? Argue your position with evidence. (5 pts)

Position

Take a clear position on the resolution below and argue it - this is the required part.

1. Where do you stand on the resolution? Argue your side using evidence from at least two of the sources you found, then fairly state the single strongest argument against your view and respond to it. (10 pts)

(OPTIONAL) Prepare Both Sides

Your teacher may ask you to be ready to argue *either* side. If so, also prepare a short case for each:

1. The strongest case **FOR** the resolution, backed by evidence you found yourself.
2. The strongest case **AGAINST** the resolution, backed by evidence you found yourself.

Step 4 - (OPTIONAL) Class Debate

Resolved: *“Schools should be allowed to discipline students for cheating based on AI-detection software.”*

If your teacher runs the debate, you’ll defend a position on this resolution, and you may be **assigned a side at random** - which is why the optional prep above asks you to be ready for both. You won’t have your paper or notes with you, so know your evidence well enough to argue from memory.

Two rules make the debate work:

- 1. Evidence only.** Every claim points to a specific source for evidence, not a feeling.
- 2. Steelman first.** Before you rebut the other side, restate their strongest argument fairly enough that they’d nod along. Then tell your tale.

One More Thing

These assignments are independent work, which runs on trust - your teacher trusting that your ideas are yours. This assignment asks you to think hard about what should count as evidence when that trust is questioned by a machine.