
HRM304 · SESSION 10 OF 11

AI & Technology in Cross-Cultural HRM

*When AI enters HR, who governs the tools?
Where does a human stay in the loop?*

Lukas Wallrich · Birkbeck, University of London
Visiting Lecturer · SWUFE · June 2026

Today, in three parts

People and culture

How AI carries bias, and how its way of writing can change the way a message reads.

A policy choice

What an organisation must decide: adopt a tool, add safeguards, or refuse it — and that changes across borders.

See it, stay in

See if an AI agent can do a real task, then decide where we have to stay in charge.

THINK · PAIR ·
SHARE

思考·结对·
分享

03:00 · in pairs

YOU ARE THE APPLICANT, OR THE EMPLOYEE

A company you applied to uses AI to **screen your CV** and to **run your first video interview**. No person sees you until later, if at all.

TALK IT OVER

- ① How would that feel — and what would worry you most?
- ② Is there anything about it you would actually **prefer** to a human?

CHUNK 1 OF 3

What AI does to people and culture

*Where does bias get in, and what does
AI's way of writing do to a message?*

Algorithmic bias

算法偏见

An AI tool can treat one group unfairly even though no person decided to be unfair. It learns patterns from past data and from the choices its builders made — and those patterns can carry social bias with them.

A few ideas to keep straight

It predicts, it doesn't look up

A chat model writes the likely next words — not a database, so it can be confidently wrong.

Probabilistic, not fixed

Ask twice, get two answers. A trained classifier (a CV screener) is more fixed; a chat model varies.

Black box, not explainable

Often you cannot see why it decided as it did — it matters when it makes an HR decision.

Bias is not noise

Bias is a steady skew against a group; noise is random inconsistency. Both are unfair.

Where bias gets in

The training data. The examples it learns from: who is in them, and who is missing.

The labels. The "right answer" tied to each example. Treat "hired last time" as "good", and it copies the past.

The scoring rule. What the finished tool is told to reward. Score fluent English as "competence", and a strong non-native speaker is marked down.

A "bias-tested" label is not proof. Biased hiring tools also bring legal exposure, not only an ethics problem.

See where bias lives: inside the numbers

To work with words, an AI tool first turns each one into a list of numbers, learned from huge amounts of text. Words that share meaning end up with similar numbers — close together — so the tool can even do arithmetic: **king – man + woman** lands on **queen**.

Run the same arithmetic on jobs and it answers **doctor → nurse, engineer → man**. No one wrote that rule. The tool absorbed it from the text — so the bias is sitting inside the numbers, and any hiring tool built on them carries it without anyone noticing.



See it live: swufe-cchrm.pages.dev/s10embed

Cultural homogenisation

文化同质化

An AI learns language and cultural style together, fused. So its writing is fluent — and it quietly pulls everyone toward one dominant style, flattening how tone and feedback differ from one culture to another.

A cultural default you can steer

Western-built models (ChatGPT, Gemini) lean individualist, direct, self-promoting.

Chinese-built models (DeepSeek, Qwen) lean more toward hierarchy and social harmony — a different default, not a neutral one.

You can steer it. Ask a model to answer as someone from a given place and the default shifts — though it can swap one stereotype for another.

These are directions, not exact scores; even measuring cultural tilt is contested. ([Tao et al., 2024](#); [Dokic et al., 2025](#))

A person who knows the tilt is there can keep AI's writing on target.

AI can keep the facts and strip the politeness

AI can hold on to what a message says while removing how it says it: the politeness, the indirectness, the context a high-context reader leans on to understand it. ([Agarwal et al., 2025](#))

From Session 3. In a high-context setting, how you say something carries as much weight as what you say.

From the ad audit. The wording itself shapes who applies: gendered phrasing quietly narrows the pool. ([Gaucher et al., 2011](#))

Live: three intentions, one voice?

Same situation each time: a report came in late, with mistakes. Three managers, three intentions. I'll ask AI to **"turn each into a professional email."**

Encourage: "the report was a bit late and had a few small mistakes — could you fix them and resend by Friday? thanks so much!"

Warn: "report's late again and full of mistakes. fix it and get it back to me by Friday."

Show I'm unhappy: "late. several errors. resend by Friday."

After the polish, can you still tell which is which?

Same note, change only the name

Now take the **warning** note and send it twice, changing only the recipient: **Bob** in our London office, then **Qiuhua** in our Chengdu office. Same words, same instruction.

Watch the greeting. "Hi Bob", or "Dear Qiuhua"?

Watch the cushioning. Short and direct, or softened with thanks and "review it carefully"?

Helpful adaptation, or a stereotype? You never asked it to choose.

IT ERASED YOUR INTENT, THEN INVENTED A CULTURE

AI flattened encourage, warn and displeasure into one voice. Then it gave Qiuhua a softer, more formal email than Bob, unasked. Helpful adaptation, or a stereotype?

THINK ABOUT

- ① A reader who can't tell a warning from encouragement: what do they do wrong?
- ② Guessing culture from a name: when does it help, when is it a stereotype? Who decides?
- ③ The one human step before any of this is sent.

So a person stays in the loop

Bias gets in through the data, the labels, and the scoring rule.

The writing has a steerable cultural default: it can flatten tone and miss context (especially but not only in translation).

So a person should stay in the loop on anything cultural or consequential — checking, adapting, deciding.

"Human in the loop" means a person reviews and decides before the output is used.

CHUNK 2 OF 3

Govern it, or it governs you

*When a firm runs AI at scale, who sets the rules —
and across which borders?*

Algorithmic management

算法管理

When software — not a human manager — hands out tasks, scores performance, and enforces rules, at a scale no manager could match.

A food-delivery app that sets each rider's route, delivery time, and rating is algorithmic management at work.

(Kellogg et al., 2020)

Same job, different algorithm

Chengdu • Meituan / Ele.me

The app sets routes, tight delivery times, and speed scores. Pressure is high, driven mainly by platform labour rules and market structure (not "culture").

London • Uber Eats

Similar app-led dispatch, but looser time pressure, shaped by different worker-classification rules and labour market.

Are there cultural differences beyond the institutions? Economic differences? Let's not overuse culture - but do not deny either.

What "watching work" can look like

Activity tracking. Software logs active time, takes occasional screenshots, and records location at clock-in.

Communication scanning. AI reads staff emails and chats to flag risk — fraud, corruption, or leaks. Common in finance compliance.

Productivity scoring. Keystrokes, app-switching and "focus time" turned into a single performance number.

Each one moves from a guess about who is working to a constant data feed.

Watching work: a real trade-off

AI can track behaviour directly: active time, screenshots, location at clock-in. That can make performance management more evidence-based. The cost is privacy. What counts as acceptable varies by culture and by law.

*A monitoring-led approach is not right or wrong
— it just will not be accepted everywhere.*

QUICK PAIR ·

03:00

Where would direct monitoring be accepted, and where not — and is that culture or law?

Moving HR data across borders

PIPL. China's Personal Information Protection Law (2021): the rules for personal data in China.
(National People's Congress of the People's Republic of China, 2021)

GDPR, and beyond. The EU's data-protection law. California's CCPA / CPRA does much the same, now for employee data too.

More than one at once. Cross a border and several apply to the same HR data.

Employee "consent" is legally weak, so firms rely on legal routes like SCCs (Standard Contractual Clauses), not a signature.

With new options come new restrictions: the EU AI Act

What it is. The European Union's law on AI systems (Regulation 2024/1689). It sorts AI uses by how risky they are. ([European Commission, 2024](#))

HR is high-risk. Using AI to hire, score, or monitor staff is classed "high-risk" — so the law requires human oversight and bias testing.

Friction when you expand. A tool a Chinese firm runs legally in Chengdu can be a compliance problem the moment it operates in Munich.

YOUR ROLE

You are the global HR steering committee of a Chinese tech firm opening offices in Europe. Five AI proposals are on your desk — weigh each one.

FOR EACH PROPOSAL

- ① Write the **key benefit(s)** and **key risk(s)**.
- ② Decide: **Adopt** · **Adopt with safeguards** (name the protection) · **Reject** (why).



Open the board: `swufe-cchrn.pages.dev/s10board`

REPORT-BACK

反馈

HOW WE'LL DO IT

Hand-vote each proposal — Adopt · Safeguard · Reject
— then we dig into some disagreements.

WHAT TO SHARE

- ① Hands up for each proposal in turn.
- ② For the big splits, some groups explain their perspectives.

One global system keeps getting harder to run

Law differs by place. PIPL, GDPR, the EU AI Act — each market sets its own rules.

Acceptance differs too. What staff will tolerate, from monitoring to data transfer, varies by culture and law.

So firms localise. Increasingly they run a localised stack per region rather than one worldwide tool.

This is another reason for localisation - alongside culture.

CHUNK 3 OF 3

See it work, and stay in charge

The tools you can actually use, what "agentic" means, and where you keep a person in charge.

Chat, or agentic?

Tools you know

DeepSeek · Qwen

Kimi · Doubao

ERNIE · Yuanbao

MiniMax ...

Many more, changing fast.

The difference that matters
is not which one, but how
you use it.

Chat. You ask, it answers in the box. Good for a quick question or a first draft, but it forgets, so you re-explain each time.

Agentic. You give it files and a multi-step job. It works the folder — reading, drafting, checking — and saves results as it goes, so the steps stay separate and you do not repeat yourself.

THINK · PAIR ·
SHARE

思考·结对·
分享

08:00 · in pairs

HOW YOU ACTUALLY USE AI

Think about how you really use AI — for study, work, or job-hunting. What is **most helpful**? And what does it **get wrong or miss**, especially about your own situation?

HOW WE'LL DO IT

- ① Think on your own for a minute.
- ② Tell your partner your example.
- ③ We'll hear a few before the demo.

Live demo: an AI agent applies for a job

I'll run a file-based agent on my own CV and a real SWUFE faculty post. It reads the job ad and the CV, interviews me, then drafts a cover letter and a tailored CV.

Capability

Multi-step,
working real files,
asking me
questions — not
one chat reply.

The human step

I will need to
check - here, very
briefly.

Culture

A Chinese CV may
carry a photo; a
UK-EU one omits
them. Which does
it pick?

Where a person stays in

AI drafts, a person decides. The output is a starting point, not the final word.

It only knows the context you give it. Missing context is missing judgement.

Check for overreach, bias, and cultural fit — the writing will be ok (though it might "sound like AI").

Never send sensitive HR text without safeguards, and know the law before HR data crosses a border.

On the horizon

Agents that act across systems. Booking, buying, filing across many tools at once, with less step-by-step human input.

Synthetic candidates. AI-generated faces and voices in remote hiring. How to make sure the applicant is real and speaking to you?

Next: bringing it together

Next session

Block 11: exam preparation.
We turn the whole module into one toolkit and practise scenarios together.

Your reflection

Today's share and the demo are your material: where AI helped, what it missed, and where you keep a person in. Full prompt on the assessment page.

References (1/2)

- Agarwal, D., Naaman, M., & Vashistha, A. (2025). AI suggestions homogenize writing toward Western styles and diminish cultural nuances. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1–21. <https://doi.org/10.1145/3706598.3713564>
- Dokic, K., Pisker, B., & Radisic, B. (2025). Mirroring cultural dominance: Disclosing large language models' social values, attitudes and stereotypes. *Societies*, 15(5), 142. <https://doi.org/10.3390/soc15050142>
- European Commission (2024). Regulatory framework on artificial intelligence (Regulation (EU) 2024/1689, the AI Act): High-risk AI systems. European Commission. <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
- Gaucher, D., Friesen, J., & Kay, A. C. (2011). Evidence that gendered wording in job advertisements exists and sustains gender inequality. *Journal of Personality and Social Psychology*, 101(1), 109–128. <https://doi.org/10.1037/a0022530>
- Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>

References (2/2)

National People's Congress of the People's Republic of China (2021). Personal Information Protection Law of the People's Republic of China (English translation, DigiChina, Stanford University). Stanford University (DigiChina). <https://digichina.stanford.edu/work/translation-personal-information-protection-law-of-the-peoples-republic-of-china-effective-nov-1-2021/>

Tao, Y., Viberg, O., Baker, R. S., & Kizilcec, R. F. (2024). Cultural bias and cultural alignment of large language models. *PNAS Nexus*, 3(9), pgae346. <https://doi.org/10.1093/pnasnexus/pgae346>