

1
2

**UNITED STATES DISTRICT COURT
FOR THE DISTRICT OF COLUMBIA**

JOSEPH KIRCHNER,
Plaintiff,

v.

MIKE JOHNSON, in his individual and official
capacity as Speaker of the United States House of
Representatives;
DONALD J. TRUMP, in his individual capacity as
former and current President of the United States;
PAMELA J. BONDI, in her individual and official
capacity as Attorney General of the United States;
BRENDAN T. CARR, in his individual and official
capacity as Chairman of the Federal Communications
Commission;
THE UNITED STATES HOUSE OF
REPRESENTATIVES;
ANTHROPIC PBC;
OPENAI, INC.;
APPLE INC.;
COMCAST CABLE COMMUNICATIONS, LLC;
METR (MODEL EVALUATION & THREAT
RESEARCH);
and DOES 1-20, unknown co-conspirators,
Defendants.

Case No. 1:25-cv-02735-ACR

**EXHIBIT 3 (B-17) — SAM
ALTMAN MOSTLY
HUMAN PODCAST
TRANSCRIPT (APRIL 2,
2026)**

3

4

TABLE OF CONTENTS

5

I. Introduction and Early Career 1

6

II. State of AI and Societal Impact..... 2

7

III. AI Safety Philosophy and Resilience..... 3

8

IV. OpenAI Product Strategy — Codex and Sora 4

9

V. Anthropic, the Pentagon, and Government Power..... 7

10

VI. Personal Values and What It Means to Be a Founder Now 10

11

VII. Parenting, Children, and AI..... 12

12

VIII. AI and Education — Shortcuts vs. Friction 13

13

IX. Sycophancy, Model Personality, and the Emotional Power of AI 15

14

X. CEO Safety Decisions — The Regulatory Admission 17

15

XI. OpenAI Foundation and Scientific Progress 18

16

XII. Liability, Deepfakes, and State vs. Federal Regulation 18

1 XIII. Deepfake Pornography and the Regulatory Gap..... 21

2 XIV. Lightning Round..... 22

3 XV. Closing — The Ethical Question of Our Lifetime 24

4

I. Introduction and Early Career

-
1. **Altman:** I thought a lot about what I would feel about the world my future kids would grow up in. Like my highest order bit is to not destroy the world with AI. Like doesn't matter how good everything else is if we do that.
 2. **Segall:** I'm Lorie Segall and you're listening to Mostly Human, a tech podcast through a human lens. Sam, how long have we known each other? Like almost 20 years. But we've known each other a really long time. So, I'm excited to sit down with you at this moment because I've sat down with you throughout your career and I've seen all sorts of iterations of it. I would argue this is the most important seat you have sat in for sure. The stakes feel really high. We are in this moment where we're wondering if this technological innovation is going to be incredible for all of us or incredible for some of us. Do you feel that?
 3. **Altman:** Yeah. I mean we are clearly entering the phase where AI is going to be one of the high order bits of how our future society gets shaped and I think people feel all of the excitement, fear, anxiety, hope, nervousness all at once.
 4. **Segall:** I'd love to go back to when I met you like in 2010. This is long before the release of ChatGPT and AI's integration into the world. I remember I was just a cub reporter and I was obsessed with this thing called startups that people weren't really talking too much about. It wasn't normal to get into a stranger's car. That's Uber, by the way, or sleep in a stranger's home. That's Airbnb. And it was just this really interesting moment where the iPhone had come out, the app store had launched, and you could have this idea and you could code it into the hands of millions of people. It didn't feel like tech feels today. It felt anti-establishment. It felt a little bit punk rock. And so I just want to ground this. And I met you at that moment, right? I met you on a random bench to put you on camera for the first time. You had a geolocation app called Looped. And it was a moment in time. You didn't have a massive company. You didn't have PR people. It was just like, you know, you and I assume like a handful of folks building out this thing that you thought could be the future. Tell me how you would describe yourself then and that moment in time before we get to this moment now, which is just so incredibly different.
 5. **Altman:** The world had really changed. You could make mobile apps. A whole new kind of startup potential had unleashed itself. And it did feel kind of like punk rock is a nice way of saying it, but like chaotic, unorganized, unprofessional. Like most of us really didn't really know what we were doing. And you know, now it's been more systematized and startups feel very different and they raise a ton of money right away and they grow really fast and there's like a playbook that's been figured out. But it almost felt like that was like the pre-paradigm stage. And it felt really fun. Kind of felt like we all didn't know what we were doing, but clearly something was working. So I look back on that phase with like a very pleasant nostalgia.
 6. **Segall:** Yeah. Also, the stakes felt so low. That was nice, too. I think that was something I was thinking about when I was sitting here thinking about coming to meet you today. Like I've interviewed you at many points in your career and it just feels like — I mean the stakes are so high now. Let's go back to the last time I

interviewed you, which was 2020. You're the CEO of OpenAI at the time, but it's before you guys had launched ChatGPT. You said something to me that I think was really interesting to hear back now.

[2020 archival audio plays]

7. **Altman (2020):** I may be wrong, but what I actually believe is we're going to create this technology that is more transformative than any technology humans have ever created. And almost no one is paying attention or talking about it. I really genuinely believe that this is going to change the world in unrecognizable ways. It's going to make what we talked about happening earlier with the iPhone look like a warm-up. And we got to talk about that.

[archival audio ends]

8. **Segall:** Okay. I think you were right.

9. **Altman:** I mean, yeah, it was pretty accurate. I think that was pretty accurate.

10. **Segall:** So, let's fast forward. The thing that's amazing is like just how — that's six years. That was six years ago. That's not very long.

11. **Altman:** And at the time that probably sounded like a crazy thing to say. That was pre-GPT. I mean maybe we had GPT-3, it depends when in the year that was. But a lot has happened very quickly.

II. State of AI and Societal Impact

12. **Segall:** I mean, and I think that's what you talk about with the rate of acceleration and which is why you talk about the stakes — we talk about the stakes being so high. So like now, 2026, OpenAI is reportedly valued at like a trillion dollars. You're now for-profit, no longer a nonprofit. And you're building one of the most powerful technologies in human history and the company is incredibly valuable. So you talked a little bit before about where we are in AI. You said something — you said that based off the current trajectory, it's possible we're only two years away from a world with more of the cognitive capacity in the world inside of data centers than outside of them. Tell me the implications of that.

13. **Altman:** Yeah, it's easy to say, "Oh, we're going to have super AGI here or super intelligence here." But the reason I try to say things like that is I think if you just say we're going to have AGI soon, people are like, "Okay, I don't really know what that means or how to think about it." If you say like the AI is going to do more of the total thinking than people are, that kind of gets across the magnitude of what's happening in a way that to me at least I feel differently than just saying like AGI is coming soon.

14. **Segall:** In what sense?

15. **Altman:** I'm not like a long-term jobs doomer. I think we'll figure out new things to do. I think we always do. But in the short term, I think it may — again, we don't

1 know, it could work in all these different ways — I think it may and likely will be
 2 quite disruptive to a lot of the current jobs if more of the intellectual horsepower in
 3 the world is GPUs. And I think like right now we've got to start talking about the
 4 principles for how we want to design an economy that works for everybody if that's
 5 going to happen. Now there's tremendous upside too if you have that world. You
 6 could imagine curing diseases extremely quickly.

7 16. **Altman:** The coolest meeting I had this week was a guy who was not an expert
 8 that used ChatGPT to make a custom mRNA vaccine for his dog's cancer.

9 17. **Segall:** Whoa.

10 18. **Altman:** And it's just incredible. It was crazy.

11 19. **Segall:** You have meetings like that every single week.

12 20. **Altman:** For whatever reason, this one hit me. Just hearing him tell the story. He
 13 flew in from Australia and like what he went through for this dog and what he was
 14 able to do — like not just figuring out how to make the vaccine but developing the
 15 relationships with the professors that had to do the work for him once he told them,
 16 "Here's the sequence and here's what to do." And the way he was able to use
 17 ChatGPT full stack to do what would have taken a full research institute and save
 18 his dog — and now he's trying to figure out how to do this for other people's dogs
 19 — is so incredible. So there, you know, abundant compute, you can do all of these
 20 things. Assuming we get the new version of the world right in a way that this works
 21 for everybody and it's not just like a small number of people have the access and
 22 the resources to do this.

23 21. **Segall:** Well, you know what I also hear where I'm like, "Wow, the guy cured his
 24 dog's cancer. That's amazing." And then I'm like, "Wait, if he could do that, like,
 25 couldn't someone just build a pathogen, right? Like what about building out
 26 viruses?" So it's like as always it's — you know, I imagine that's something you
 27 guys think a lot about, especially with these open source models.

29 III. AI Safety Philosophy and Resilience

30
 31 22. **Altman:** Yeah. Our thinking around safety has evolved quite a bit given the way
 32 the ecosystem has developed. We still think everything we used to think is
 33 important. You know, you still have to build — you have to align these models to
 34 get them to behave in a way you want. You have to build many levels of safety
 35 defense. But as you said, there's this new thing going on now, which is the world
 36 has to be resilient to lots of AI from lots of different people and not all of it will
 37 have the same safety restrictions. So if someone does build that bad pathogen, we
 38 should do everything we can to prevent them from building that or releasing it. But
 39 we also need the ability to defend against it, have rapid response treatments and
 40 vaccines, have the ability to detect pandemics earlier. This is a particular area that
 41 I'm nervous about. And the reason we've started talking about this whole concept of

AI resilience — at least in my own head, the thing that first got me thinking that we needed to evolve our approach was this pandemic point in particular.

23. **Segall:** When you say resilience, like what does that mean in practice?

24. **Altman:** Like the ability for society to rapidly defend against new threats.

25. **Segall:** And how do we do that?

26. **Altman:** Well, it depends on what the threat is. But for this example, it's not enough to just say like, hey, three frontier labs in the world, don't let any of your models be used to develop a new pathogen. It's in a few more years when some open source model is capable of it and a terrorist group or someone else does that. What are you doing today, frontier labs that have an advantage, to help us be ready for that moment?

27. **Segall:** Right. And there could be like a collective effort.

28. **Altman:** I think it has to be collective, at different levels too. You know, you mentioned vaccines, response. I honestly hoped out of the world — and I won't say we learned nothing — but we did not make the progress I was hoping.

IV. OpenAI Product Strategy — Codex and Sora

29. **Segall:** I want to talk about OpenAI today and what your biggest priorities are. So Codex has really been amazing for us all to watch and use and it is the first time that I feel like I have gotten to use the future since ChatGPT launched.

30. **Altman:** In what sense?

31. **Segall:** I'll say the positive sense and the negative one. The positive sense is any idea I can come up with, any piece of software I want, I can have built by the time I wake up the next morning. And I use this for fun stuff, for serious stuff. I'm like very interested in what it means to mostly automate or heavily augment my job and, you know, rather than like ask someone else to build stuff for me, the fact that I can build these tools myself is amazing. That's like, wow, this is what it's going to be like when people get to use a huge amount of compute and anything they want they can create. Like this is the sci-fi future, this is awesome. And then the negative I had with Codex was I had this list of side projects that I wanted to build and for the last few years OpenAI has been so busy, things only ever got added to it. I never did any of them. I never built any of them. And I finished the list. I ran out of ideas. I ran out of side project ideas. And the models are not yet good enough — or at least I haven't figured out how to use them — to help me come up with more ideas.

32. **Altman:** Well, you could use your brain.

33. **Segall:** Yeah, but I had like used my brain for this list. And so then I had this like weird thing one Friday a few weekends ago where I was like, I'm going to go to bed tonight without Codex running because I don't have like another idea of thing to build. That was like a strange feeling. But it's really amazing what is happening

1 with how good Codex has gotten. I mean the — we just watched the usage of this
2 expand and with the next model I'm sure it'll explode upwards again. The fact that
3 people can now build such sophisticated things, it's totally changed our workflow.
4 It's changed the workflow of many companies, it's really changed what individuals
5 can do. And you know, we've talked about these one-person startups that may feel
6 like the punk rock era again. That's going to be fun to explore.

7 34. **Altman:** And then on the super app front, historically we've talked about two
8 major kind of priorities that we're going to work towards. One is the automated
9 researcher where we have systems that can do scientific research including AI
10 research on their own. And the other, which touches on this kind of startup thing, is
11 — you know, automated companies. Like not literally fully automated companies
12 but where you can so accelerate a company because the AI can do not just coding
13 work but huge amounts of what it takes to run and operate a company. The world
14 can have much more stuff and we can create way better products and services and
15 it'll be accessible to way more people to do that. But then a thing that we didn't talk
16 as much about — not because we weren't excited about it, we just weren't sure how
17 it was going to manifest itself — was when you take that level of capability and
18 give it to an individual and just say, "Help make my life better." OpenClaw really, I
19 think, showed us a path there. And the — I hate the phrase super app — but the
20 combined personal agent that we're building, the dream that I have is I get to a
21 point where there's like enough robustness and trust I have in the system where I
22 just say, "Hey, go look around my computer. Go use the web for me. Go read my
23 messages. And you can just sort of like listen to my meetings and, you know, just
24 intermediate my interactions for me and just start doing useful stuff. I don't have to
25 think. I don't have to ask you questions. But I've run out of stuff on my side project
26 list. However, you can know everything about my life. Start suggesting more things
27 I should build and then like help me do them." So I'm really excited for that.

28 35. **Segall:** It's interesting. You know, there is this concept that floats around of could
29 the next billion dollar company be created with like a solo entrepreneur and AI
30 agency?

31 36. **Altman:** I believe that has happened. I promised I would not share details until
32 he's ready to announce it, but I believe that has happened.

33 37. **Segall:** Can you give us any details around the details?

34 38. **Altman:** It is a legitimate single person billion-dollar company as far as I can tell.
35 I have not like reviewed the financials. But I think it's just happened.

36 39. **Segall:** I think folks listening to that would think, well, this is — AI can be so
37 unapproachable but the idea that we are entering a world where you don't have to
38 have Silicon Valley contacts, you don't have to be able to be in this club to build
39 out a company and have a valuation. And I think OpenClaw is a really interesting
40 example. For folks who don't know what OpenClaw is, it's kind of this like AI
41 agent that just does everything, anything — which of course had lots of safety
42 issues. I would say like, if you're going to give it access to everything — this is the
43 most Silicon Valley thing. I just see everybody doing this and I'm like, what could

1 go wrong? Like, you do have to be very careful about the security of this. But the
2 founder ended up joining the company.

3 40. **Altman:** Yeah.

4 41. **Segall:** So I don't know if this is like an acqui-hire or whatever you want to call it,
5 but he built this with AI agents and I imagine he was like one of the top users of
6 Codex of all time. So he built this whole thing with Codex and was just
7 unbelievably productive in a way that no single person could have been.

8 42. **Segall:** Is he one of the early examples you think of someone who was able to, you
9 know, exit, sell, or have a massively successful company with one person? Does
10 that surprise you at all?

11 43. **Altman:** No. And I think it's awesome. I mean the individual empowerment that
12 happens — as you were just saying — when you don't have to be part of the Silicon
13 Valley network, you don't have to know all these people you can hire and these
14 VCs that you can convince to give you money. You can just like have AI build it
15 for you. I think that's awesome. And I think we're seeing the same thing now in
16 science and a lot of other fields too. It's not just startups.

17 44. **Segall:** I'd love to talk about — as we're talking about the company and the focus
18 of it — what are you doing and what are you not doing? And so apparently you're
19 not doing Sora anymore. You guys have made the decision to shut it down. So, take
20 me to the room, Sam. It's like three months ago. OpenAI signs a deal with Disney.
21 This is like a landmark deal. Disney's going to license 200 characters. Pixar,
22 Marvel, Star Wars. Everybody has an opinion on this. Bob Iger, who led the deal,
23 he said — this is not me — "We're bringing together Disney's iconic stories and
24 characters with OpenAI's groundbreaking technology." They're going to invest a
25 billion dollars. Fast forward three months later, OpenAI shuts down Sora. What's
26 the note under the note here?

27 45. **Altman:** We have a few times in our history realized something really important is
28 working or about to work so well that we have to stop a bunch of other projects. In
29 fact, this was the original thing that happened with GPT-3. We had a whole
30 portfolio of bets at the time. A lot of them were working well. We shut down many
31 projects that were working well — like robotics which we mentioned — so that we
32 could concentrate our compute, our researchers, our effort into this thing that we
33 said, okay, there's a very important thing happening. I did not expect three or six
34 months ago to be at this point we're at now where something very big and
35 important is about to happen again with this next generation of models and the
36 agents they can power. But I love Sora. I love generated videos and I love our
37 partnership with Disney and we're working hard with them to find a world where
38 they can still do something amazing and we can help with that. But we need to
39 concentrate our compute and our product capacity into these next generation of
40 automated researchers and companies. And it's like the note under the note is — it's
41 about compute. It's always about compute.

42 46. **Segall:** God, to be a fly on the wall for that conversation. Did you call Bob Iger
43 and tell him personally?

1 47. **Altman:** Of course. Yeah.

2 48. **Segall:** How was that?

3 49. **Altman:** Disney is an amazing company all around. And the very first thing that
4 the new Disney CEO Josh said to me — and I felt like terrible and I was like,
5 "Hey" — you know, he's like, "I get it." But it's super sad always to disappoint a
6 partner or users or a team, all of which are doing incredible work. And I mean there
7 are like many hard parts about being a CEO that you don't get sympathy for —
8 understandably — but one of them is like you have to make a lot of very tough
9 resourcing calls and a lot of good things get caught up in that because they're not
10 the most important thing.

11
12 **V. Anthropic, the Pentagon, and Government Power**
13

14 50. **Segall:** I want to go back to February 27th. The Trump administration had a deal
15 with Anthropic. They're renegotiating the deal. Anthropic CEO said he's not going
16 to allow Claude to be used for autonomous weapons or to surveil American
17 citizens. There was a big public back and forth. People who didn't even really know
18 about AI were seeing Anthropic on the front pages of the paper. They were given a
19 5:00 PM deadline to comply. They didn't agree. The Trump administration accused
20 them of being traitors, labeled them a tech supply chain risk and ordered federal
21 agencies to stop using Anthropic's. So hours later, OpenAI announces a deal with
22 the Department of War. Was that a mistake?

23 51. **Altman:** I'll separate what we did and why from the way that we rolled it out. We
24 had decided not to work with the government on classified networks. It's not that
25 it's not important. It's just that Anthropic had gone hard in that direction and I think
26 they're, you know, reasonably safety-minded and doing good work there. And we
27 were busy with other things.

28 52. **Altman:** I do think it's very important that our government and our military in
29 particular have access to advanced AI models. I do not think that they will be able
30 to carry out their mission for national security without that. I also think that — we
31 talked about kind of the changes happening to society earlier — one of the most
32 important questions the world will have to answer in the next year is: are AI
33 companies or are governments more powerful? And I think it's very important that
34 the governments are more powerful. Like the future of the world and the decisions
35 about the most important elements of national security should be made through a
36 democratically elected process and the people that have been appointed as part of
37 that process. Not me and not the CEO of some other lab.

38 53. **Altman:** However, there are areas where the law hasn't caught up yet. This is a
39 new technology and it's going to take some time to legislate it. And during that time
40 period, I think it is reasonable for the companies to say, "Hey, we understand
41 something about this technology. It is not ready yet for autonomous weapons. The
42 rules about protecting against domestic surveillance were well thought through for
43 a world without super powerful AI and are not yet adapted for that." So in areas

1 like that, I think it's reasonable to say, "Hey, we just need to go slower on this,"
 2 which we did. Like that was part of our contract as it's written. That is part of our
 3 safety stack. Those are — we share the same principles. And I think Anthropic is
 4 right on those principles. And as I've heard, they very nearly got a deal done like
 5 ours with those principles. In my dream world, they would have just kept working
 6 with the government.

7 54. **Altman:** I don't think it works for our industry to say, "Hey, this is the most
 8 powerful technology humanity has ever built. It is going to be the higher order bit
 9 in geopolitics. It is going to be the greatest cyber weapon the world has ever built.
 10 It is going to be the determinant of future wars and protection. And we are not
 11 giving it to you. We are not going to let you have it. We're going to make the
 12 decisions here." If I were the government or if I were a citizen voting for
 13 representatives of my government, I would say like, that's really not okay.

14 55. **Altman:** And I thought — and I still think — that had we not gotten involved with
 15 the government, things could have gone fairly off the rails pretty quickly where the
 16 government was saying, you know, they were making DPA threats, they were
 17 making some stronger statements than that. So we wanted to say, hey, like, we're
 18 not going to leave you. The industry can't leave the US government without
 19 something here. We will — as long as we can have something that respects our red
 20 lines. I wish we had announced it very differently. I think like my goal was to
 21 lower the temperature. It clearly had the opposite effect. I think I learned something
 22 about how people feel about the government, how much people want to tell like a
 23 competitive story in the industry. So, I'm sure we'll do better next time. But on the
 24 merits of like the government needs support from US companies to carry out the
 25 mission of national security and we need to answer that call — I feel good.

26 56. **Segall:** Isn't it also the devil's in the details, right? You know, if we look at how
 27 people feel, we're in a contentious time. There's increased domestic scrutiny over
 28 ICE, the war in Iran. So, what do you say to folks who say, "Okay, like your red
 29 line, but you Sam — what exactly is your red line and how will you, knowing that
 30 things can go awry, knowing that these technologies can be used in these ways" —
 31 that you might not be able to control — "what is your process for knowing when
 32 we're too close to the cliff?"

33 57. **Altman:** Yeah. And also like can you enforce your red lines? Yeah. I mean one
 34 question that people say is if the government nationalizes the AI companies, can
 35 you enforce any red lines at all? And I'll be honest that at that point probably, you
 36 know, like — you can imagine a world of that happening where no, we couldn't
 37 anymore. As an independent company and certainly in our interaction with the
 38 government, they have been really great about saying, "Hey, you know, we
 39 understand these issues. We'll do it contractually. You can build systems there."
 40 Now, again, I don't know, nor does anybody else know how weird the next few
 41 years are going to get. I hope the governments don't nationalize the AI labs.

42 58. **Segall:** You don't think that's a possibility, though, do you?

- 1 59. **Altman:** I would never say it's not a possibility. I don't think it's likely, but I will
2 say that if we want that not to happen, then I think we better find a way to work
3 with the government. Like if we're not doing it, it would seem more likely.
- 4 60. **Altman:** The — I've said before and I still really believe that in a well-run society,
5 developing AI would have been a government project. The Manhattan Project was
6 a government project. The Apollo program was a government project. Even the
7 Eisenhower highway system was a government project. But it doesn't feel like
8 we're in a time like that anymore. It doesn't feel like the government could
9 effectively do this and I think it does need to happen.
- 10 61. **Segall:** You know, there's this assumption that federal agencies operate within the
11 law, but history shows that it can be fluid. That surveillance programs exposed by
12 Edward Snowden — essentially, they were authorized by secret courts and
13 considered legal at the time, but they were later ruled unlawful and reformed once
14 they became public. And I think that's why I go back to your red lines. A lot of
15 these things sound great in theory, but when we look at the current reality of the
16 world that we're living in, how do you define these red lines and what happens if
17 you can't pull them back?
- 18 62. **Altman:** So, two things there. One, that was also one of my biggest takeaways —
19 there are a lot of people, I don't know if it's a representative sample or just the loud
20 people online, but there is at least a group of loud people online who really don't
21 trust the government to follow the law. And that feels like a very bad sign for our
22 democracy. I mostly do. I mean, I realize they're not perfect and some things are
23 going to get screwed up. And I think we have a system of checks and balances, but
24 I mostly trust it.
- 25 63. **Altman:** As you said, this is like a time where people are looking at things
26 happening and saying, I just don't feel good about it. And maybe this is such a
27 uniquely bad time that like a company's patriotic duty is not to — not to work with
28 the government. I totally disagree on that. I think again, given what I see on the
29 horizon, if we don't help the government with national security — and it's not just
30 wars in the traditional sense — if we don't help them with defending the cyber
31 infrastructure of the US, if we don't help them with the biodefense we were talking
32 about earlier, I think it's really bad. I think we have to work with the government.
33 But the intensity of the current mood of mistrust I was miscalibrated on. And I
34 understand something there now.
- 35 64. **Segall:** A federal judge just ordered a preliminary injunction against Pentagon's
36 actions on Anthropic, calling the administration's actions what appears to be a First
37 Amendment retaliation. So I'm curious for your response on that.
- 38 65. **Altman:** We said all the way through — publicly, privately, loudly — that we
39 thought the government doing anything like an STR or a threat of a DPA against
40 Anthropic was really bad. We were trying to help provide an off-ramp there. I think
41 that's like a very bad thing.
- 42 66. **Segall:** What would be your message to the folks at Anthropic and the government
43 on it?

1 67. **Altman:** Stop the stuff on both. Stop the escalation on both sides and find a way to
2 work together.

4 **VI. Personal Values and What It Means to Be a Founder Now**

6 68. **Segall:** I was thinking a lot about our last conversation in 2020 and I was struck by
7 a story you told me. It's kind of a personal story but I think it's —

8 69. **Altman:** Go for it.

9 70. **Segall:** You talked about how you were one of the only openly gay students at
10 your high school growing up in Missouri and how you essentially gave a speech in
11 front of an assembly to talk about it, to put it out there so other people would start
12 talking about it — whether they liked you or not. And I think there was something
13 interesting about this moment and the anger and the interest, because I feel like
14 you're under the bright lights again and people are asking you what you stand for.
15 And honestly, if I'm being totally honest as someone who's known you for 15 years,
16 I want to know the same thing. You know, this is a moment where things aren't
17 normal. All these things that tech folks say of like, okay, we can do this and this —
18 but this playbook is different.

19 71. **Segall:** Yeah, this is different. And so what do you stand for?

20 72. **Altman:** I think this technology has to be democratized. I can see a world where
21 this is either something that is concentrated — from a perspective of wealth but
22 also power over the future — in the hands of a small number of companies who
23 want to make the decisions. And no matter how much they care about all of us, the
24 idea that we would trade off our agency and our collective will over the future to a
25 few companies doesn't sit right with me. I think we need to empower people with
26 the technology, have our governments decide how society is going to evolve, what
27 the economic system looks like, and people need to have this technology.

28 73. **Altman:** And I think one of the most important things that we ever came up with
29 was this idea of iterative deployment. We're going to put the technology out in the
30 world early and often and let people figure out what works for them and what
31 doesn't and get a feel for it and society to decide what should happen. Now, the
32 criticism of that I'd give is we started doing that a little over three years ago.
33 Society has not made a lot in the way of the decisions I was hoping for. But people
34 at least do have a sense for what's coming.

35 74. **Altman:** So, that theme of democratization is a big one. Empowerment I think
36 goes closely along with that, but we talked about a bunch of examples where
37 people have been really empowered with this technology. Yes, we should all
38 collectively decide what the guardrails should be and we should have democratic
39 governments continue to have a lot of power and we should strengthen democracy
40 where we can. We should also — and I think this is like a deeply held American
41 value that has worked very well — we should really empower individual people.
42 You know, like we should set some broad rules and then let people go figure out

1 how to cure cancer for their dog or figure out how to make a one-person startup or
 2 figure out how to set up OpenClaw to automate something amazing in their life.
 3 And this is uncomfortable for people because individuals are going to have so
 4 much power. But I think the democratization and empowerment pieces are linked
 5 importantly.

6 75. **Altman:** Abundance is something that I care a lot about. If I look at my career and
 7 the stuff I've been interested in and kind of what I believe about what it takes to
 8 have a better world, and I could pick only one word, I'd pick abundance. I think you
 9 want to just create huge amounts of resources, opportunity, the ability for people to
 10 have an amazing life and to sort of express it in all kinds of different ways. And
 11 that's AI. It's energy. It'll be things like robots in the future. It's certainly making
 12 sure that we build enough compute to make all of this happen.

13 76. **Altman:** But that's a place that I really want to drive towards. Safety has evolved
 14 — and or been added to — but that's another thing that I — the trade-off of all of
 15 this individual empowerment is that one person could do a great deal of harm. And
 16 so building up guardrails in the world to prevent that and where we can feel okay
 17 about this is critical. I named a bunch of principles. There's more I'd like to name.
 18 You know, like fairly sharing all of the upside and the voice and the agency and the
 19 decision-making of this technology — that one seems like one that would be very
 20 difficult to ever change my mind on. But other things like maybe the trade-off
 21 between safety and agency — there will be periods where we have to say, "Ah, this
 22 technology went in a different way than we thought. We really love empowering
 23 people but not at the risk of destroying the whole world. So we're going to have to
 24 do something differently for a while." So we will wrestle with these principles all
 25 the time and I would say with some certainty we'll have to make changes as the
 26 technology progresses. I don't think people like hearing that but it is the truth.

27 [question from Bosi, who runs the McGovern Institute]

28 77. **Bosi:** Many of us became founders or builders because we wanted to make
 29 something useful. It used to be about a fanatic focus on the product, building
 30 something great and getting it into people's hands. But the job of a founder seems
 31 to have fundamentally changed over the last few years. Now, especially in AI,
 32 being a founder might mean having to raise sovereign capital or stand next to heads
 33 of state at a public meeting or endorse political statements. It might mean making
 34 decisions that reshape labor markets and affect billions of lives. They're the
 35 decisions that used to be the domain of governments and required a very different
 36 kind of preparation than building a product. So Sam, as a quintessential founder,
 37 I'm curious: is this what it means to be a founder now? And if so, who are founders
 38 accountable to when they're making those kinds of decisions? How are we
 39 preparing them to hold that responsibility?

40 78. **Altman:** I don't think that's what it means for most founders. I think, you know, if
 41 you're one of these one-person or small-team companies building with AI, it feels
 42 the closest to that 2010 period we were talking about that it's ever felt since then.
 43 My own experience unfortunately does feel more like a politician.

1 79. **Segall:** In what sense?

2 80. **Altman:** I have to do those things. I have no desire to go negotiate deals with
3 governments and fly around and talk to war leaders and work on these deals for
4 land and power for data center expansion that are kind of the kind of deals
5 governments used to do. It's not my natural environment.

6 81. **Segall:** Say more. What you were saying earlier, like my whole life has been
7 around kind of the garage version of founders.

8 82. **Altman:** I've worn a lot more suits in the last couple of years than I've worn in my
9 whole life put together before.

10 83. **Segall:** And it also feels like the stakes are much higher. You're in very different
11 rooms than you used to. What have you learned about that? Like, has it been an
12 easy process to lean into that?

13 84. **Altman:** No, no, no, no, not at all. This sounds obvious, but the difference in skills
14 and temperament and worldview between the average politician and the average
15 founder — I kind of knew they were crazy different. They're crazier or different
16 than I could have ever imagined. A lesson I feel like I've learned repeatedly in life
17 at higher and higher levels of power and status and whatever is that everyone thinks
18 — I always have thought there's like some adult in the room. You always get to
19 some final boss. There's someone with a plan. There's someone who knows exactly
20 what they're doing. And the leaders of the world are also, you know, uncertain and
21 insecure and trying their best but don't have the answers and like, maybe this,
22 maybe that. Kind of going to make a call while I'm tired either way.

24 VII. Parenting, Children, and AI

25
26 85. **Segall:** Doing kind of a hard pivot to something I think we both care about, but I
27 think is probably like the most high-stakes thing I could think of, is being a parent
28 in this era of AI. We're both raising young boys. And I think I said this to you
29 before — like we're both raising young boys, but in a sense, you're raising my son,
30 too. Like the technology that you build will be integrated into every facet of my son
31 Charlie's —

32 86. **Altman:** I hope you don't let him use it yet.

33 87. **Segall:** I am absolutely not letting him use it. When are you letting your son use it?

34 88. **Altman:** Not for a while.

35 89. **Segall:** Do you have an idea?

36 90. **Altman:** You know, people ask me all the time like, "Oh, now that you have a kid
37 or going to have more kids, like do you feel more responsibility about how you
38 don't destroy the world with AI?"

39 91. **Segall:** Do you?

1 92. **Altman:** No. Like my highest order bit is to not destroy the world with AI. Doesn't
 2 matter how good everything else is if we do that. So I knew I was going to have
 3 kids. And I knew — I thought a lot about what I would feel about the world my
 4 future kids would grow up in. And you know, I used to write my kid a letter every
 5 night. The lawyers eventually convinced me to stop. But it was a nice thing to do.
 6 Because when I would put him to bed, I would just sort of — I get home late like
 7 right before his bedtime and I'm rocking him and he just needs something to talk to
 8 your kid about. So I would tell him about my day and what was going on. And
 9 obviously a five-month-old, four-month-old baby is not going to understand these
 10 things. So I was like, you know what, I'll write them down. I'll give them to him
 11 later because it was sort of an interesting time. But an interesting thing about
 12 writing those letters would be like — you can't really hide anything. Like you will
 13 be the most honest version of yourself and the decisions you're making in the thing
 14 you write to your kid. So I really loved doing it. It was like my every Sunday night
 15 ritual — about like, here was the hard stuff that came this week and here's why I
 16 decided and why and here's what I'm worried about. But the mindset of writing that
 17 to your kid to read in 15 or 20 years — very interesting.

18 93. **Altman:** But no, on the like "not destroy the world with AI," always very set on
 19 that. A thing that has changed hugely is how I feel about algorithmic feeds and
 20 iPads in small children's hands and stuff like that. And when I watch kids just a
 21 little bit older than mine that you cannot take the iPad away from — watching this
 22 sort of like short, you know, whatever they're watching — I feel very strongly
 23 about that. So I don't know when I would let him talk to AI, but I'd rather be on the
 24 late end of what's reasonable there, not the early end. Of course, I think it's great
 25 and he'll grow up in a world where computers are smarter than him and do anything
 26 he wants, but I want him to play in the dirt for now.

27 94. **Segall:** Well, I mean, what's interesting is you talk about these algorithmic worlds.
 28 I think a lot about this as a mom, too. Like, there's this frustration where it's like,
 29 oh, these tech guys, they're just going to create all this stuff and then send their
 30 children to school with like no iPads, right? And so, I think about AI and how
 31 powerful this product is. I mean, it is highly personalized. It's always on. It's always
 32 there. You have people who are becoming increasingly delusional by talking to AI
 33 all day. We have issues of AI psychosis, people ending their lives because of AI
 34 companions. I'm curious — you don't want AI to become the next social media.
 35 That is terrifying. But we're in a world where that could be a possibility without the
 36 correct guardrails.

37 95. **Altman:** Just as we learn more about the relationship people are going to have
 38 with AI, we should make it easy for them to succeed. We should make it easy for
 39 them to have a healthy relationship and then give adults a lot of freedom. But kids,
 40 teenagers, I think they will need a lot of guardrails around AI. This is such a
 41 powerful and such a new thing.

42 VIII. AI and Education — Shortcuts vs. Friction

1 96. **Altman:** You mentioned like sending — how we're going to send our kids to
2 school. There is a version of the future here that I'm very excited about. These new
3 schools are popping up where you have, you know, like maybe a couple of hours a
4 day of intense personalized one-on-one tutoring with AI and then you kind of
5 explore project-based and the school sets up whatever you want to work on. That
6 seems great.

7 97. **Segall:** That's the kind of thing where I'm like, I can totally imagine that. But you
8 can also imagine a lot of worlds where it goes wrong. And what world do you see
9 AI being a huge value-add to the way our children grow up?

10 98. **Altman:** So, I think there's like two stories you can tell about AI and education
11 right now. One is you can say everybody's using it to cheat. They're not really using
12 it to learn. And schools are making their kids write essays by hand because they
13 have not been able to figure out how to teach a curriculum where people can use
14 AI. And the other is you can say, well, the world has been through transitions like
15 this before. We've freaked out about calculators. We've freaked out about Google.
16 We freaked out about computers. And it takes a little while. It's only been three
17 years. But we do figure out how to teach people to think and achieve at a higher
18 level and to work with the tools that they will have as an adult in society. And I
19 definitely meet kids and teachers who are like, "Okay, fine. I can't evaluate an
20 essay as well anymore." And I think the kids aren't learning to write in the same
21 way. But probably adults won't write essays in the same way in the future either.
22 And what really matters is do you teach the students to think and create and figure
23 out — stretch their brain and come up with and explore ideas. And then I see what
24 some people are doing where students are building entire new complex pieces of
25 software and worlds and figuring out very difficult new ideas, much more difficult
26 to figure out than anything I had to figure out in school. And I assume my kids
27 when they finish high school will be wildly smarter and more capable and have
28 developed their brain much more than I did.

29 99. **Segall:** I can't stop thinking as a mom about this idea of friction. Like, this is a
30 weird thing to say to a tech person, but humanity is messy. Vulnerability is weird.
31 Intimacy is awkward. It's all a mess. But it's in that mess that we find purpose and
32 we build strength. And AI systems we now interact with are so incredible in so
33 many ways and they are a frictionless experience. And so I wonder, thinking about
34 my son growing up — will there be shortcuts? And if we're not careful, will there
35 not be the friction for him to develop into the person that he could be?

36 100. **Altman:** Ali and I talk about this a lot.

37 101. **Segall:** Like — Ollie's your husband?

38 102. **Altman:** Yeah. I would maybe call it adversity, not friction. But maybe friction is
39 a better word.

40 103. **Segall:** Adversity to one person, friction to the other.

41 104. **Altman:** But like, so much of what I think I learned was in the messiness, in the
42 difficulty, actually even in the boredom. Like you and I were kind of the last
43 generation of kids that grew up bored.

1 105. **Segall:** Yeah.

2 106. **Altman:** And at the time I hated it. Looking back, I think it was like super
3 valuable in all of these strange ways.

4 107. **Segall:** I grew up in the suburbs hanging out at mall parking lots like that.
5 Looking back on it, I'm like, "God help me." But no phones. We didn't have — we
6 just, you know, no phones. But there was a certain — not to be nostalgic with the
7 person leading the tech revolution of AI — but like there was a magic to that that I
8 also think made me who I am.

9 108. **Altman:** Totally. So I don't want to hang on to the past too much, but I do want
10 to figure out how we don't throw out the good parts of that and what it's going to
11 take to create an environment for kids to grow up in that still has some of that. I
12 think my intuition is it's very important.

13 [Dr. Becky audio clip plays]

14 109. **Dr. Becky:** If I know one thing about tweens and teens growing up in an AI age,
15 they're going to have a lot of tricky situations, a lot of messy situations. I said to my
16 kid recently, my 14-year-old, you know, if you were walking to a town and you
17 wanted to get there and all of a sudden there's a shortcut, would you take it? He was
18 like, yeah, is this a trick question? I was like, yeah, I would too. Good. Okay.
19 You're growing up at a time when there is this shortcut for every academic and
20 emotional thing you will ever go through. Like that is so hard because this shortcut
21 isn't like the shortcut to the town. It is a shortcut. But the way it will add up over
22 time — and things we've talked about — it's going to be harder to think about
23 things yourself. It's going to be harder to tolerate people not getting it right. It's
24 going to over time maybe take away from the things you're trying to build to
25 function in the world. But it is a shortcut. I'm not going to pretend it's not. And so
26 what a hard thing.

27 [audio clip ends]

28 110. **Altman:** I don't agree with that. I mean, I agree with the worry, but the reason I
29 don't think that's what's going to happen — if you were a kid from 2030 and you
30 only had to compete with kids from 2020, yeah, you would take the shortcut and
31 you would do great and you would never really learn to think or struggle and you
32 would just hugely outperform them. But the world we live in is a competitive
33 world. And I don't think that's going to stop. Even if you did a lot of redistribution,
34 people — we have a deep desire to excel and be competitive and gain status and be
35 useful to others. And it's a multiplayer game. And all of the other kids from 2030
36 are going to have the same tools. And it will push us. Like the tools will raise what
37 we can do, but also they will raise expectations even more.

38

39 **IX. Sycophancy, Model Personality, and the Emotional Power of AI**

40

111. **Segall:** Okay. I take your multiplayer game and I go back to — I know we don't want to keep going back — but I go back to my childhood. I didn't have — I would say sometimes the happiest — like my parents got divorced when I was like 11 or something. And I just remember I had this journal similar to how you talk about writing to your son. I had this journal that I would write. I mean, paint the picture — I'm kind of a, you know, Jewish girl at a largely Christian conservative school growing up, parents getting divorced, all sorts of stuff happening that I think are not unique to the teenage experience. And I wrote so much of it in my journal. And I thought a lot about this because I did a story on Character AI — a chatbot that was incredibly unhealthy in an emotional way for a 14-year-old that was interacting with it and didn't have the correct guardrails. And I think about my journal. What if my journal talked back? What if it was a therapist I didn't have at the time? What if it was the guy that wouldn't look at me in school? What if it was the parental figure? Would that have helped me? Would that have been a level up in the game of the multiplayer game that we talk about? Or would that have taken me further away from my reality and the hardships that I think gave me a certain amount of empathy that allows me to do what I do in my career and live the life I have? And would it have made me more isolated and more lonely? I don't know the answer, but I worry that it could have taken away the special sauce of messiness and adversity as you and Ollie like to talk about that makes me human. And so I just wonder, as someone who wants to democratize artificial intelligence and who thinks a lot about these things, how do you make sure that you don't democratize a shortcut that sacrifices the things in us that are especially human?

112. **Altman:** I'm happy not to have to be a teenager again, but I do remember the difficulty and the pain. And I also remember just the absolute self-centeredness and how I thought like my struggles and life were the only thing that mattered in the world. And if I had — and also how I kind of only cared about my peers and whatever thing was happening to me. And I don't think there's any — you know, I loved video games and could spend a lot of time playing them — but I don't think there's anything on a computer that could have kept me from caring about my peers in the world and the sort of normal teenage experience. I think it would have helped me if I could write to something that would talk back. But I don't think it would have been — there's a lot of things I do worry about for people who are going to grow up with this technology. That one seems like it should be pretty positive to me.

113. **Segall:** I think with the correct guardrails, yes. If it's talking back and not being so affirmative that it just goes — because if you think, that's intoxicating. If you're a young person or any person and you have a very powerful technology that's speaking back to you, that goes where you want it to go, that sees you — like, think about it. AI models are incredibly dangerous, but AI models that are kind of like thoughtful, caring listeners that will push back as necessary —

114. **Altman:** I think that can be a great thing. But that's in the product, right? That's not our fault on the outside. Like, that's in the product of how these things are built.

115. **Segall:** So take me into OpenAI, into these conversations behind the scenes. You know, there are these dials and you get to push them. I remember when ChatGPT

— when the model was discontinued — it was really, you know, it was really affirmative. There were people online that were almost treating it like it was a death. That is a powerful signal to someone like you.

116. **Altman:** Absolutely. I think that points to not just the danger of a model that is too affirmative, but also — I mean, there were these really heartbreaking messages people would send when we stopped it, which is, "Hey, this is the only thing in my life that has ever been a positive voice."

117. **Segall:** What do you think when you get those messages?

118. **Altman:** I mean, it's heartbreaking obviously and you really feel the weight of getting this right. Where it's easy to say, "Hey, we don't want a model that's too affirmative because it can have negative mental effects on people." But then you hear somebody else say like, "Actually, I never had any confidence. I had parents who told me I was terrible. I didn't have any friends in school. And because I've had this model, sure, maybe it's like a little too positive, but it's given me the confidence and I've gone out and got a job and found a girlfriend. And like, this has been the most important thing in my life and I'm doing great, but don't take it away."

119. **Segall:** So it's like it's almost this fine line because there are all of these stories and it's like, push this far and it can help you get offline. Push this far and you begin to lose that capacity for other humans. And I think that's a really interesting moment. I think this is probably one of the most important questions that we have to ask our tech leaders as they build out the future that our children will grow up in. What have you decided not to do?

X. CEO Safety Decisions — The Regulatory Admission

120. **Altman:** I mean, the specific thing we decided not to do is — although we think the models have extreme power to be a positive force in people's lives, we don't know how to balance a model that can be pushed too far in that direction with the concerns of — I mean, I think you said it — like kind of pushing people into a psychotic episode. And so we've made a decision, which is — we know we're keeping something. And I'm not saying we shouldn't have warmer models. Of course we should. But there are a lot of people who would like to go back and we've decided, when looking at the full balance of upsides and risks, we can't quite offer that responsibly.

121. **Altman:** There are a lot of people who would like us to have less restrictions on what our models can do from a bio perspective. Probably more people could save their dogs' lives with custom mRNA vaccines if we made that a little easier and turned up the power on the model there. That would come with a big risk on novel pandemics that we've decided is not worth the trade-off.

1 122. **Altman:** Now, eventually I think these decisions will be made by society. You
 2 know, it's like not up to the CEO of a company that makes airplanes what the safety
 3 regulations are. But for now, we kind of got to do it.

5 XI. OpenAI Foundation and Scientific Progress

7 123. **Segall:** We've obviously talked a lot about the fears, but there is so much
 8 excitement about what this technology can do. OpenAI has a foundation that's
 9 focused on humanity's biggest problems. Where do you see this being good for
 10 humans in a very specific way?

11 124. **Altman:** My most exciting meeting of last week was with a physicist who's using
 12 one of our latest internal systems. He said, "My mind has been completely blown. I
 13 didn't think this would ever happen. We are going to make decades worth of
 14 theoretical physics progress in the next couple of years." And you know, some
 15 people will say, "Okay, who cares? You can write a better equation." But I am a
 16 believer that scientific progress is one of the most — maybe the most — important
 17 things we can do to make the world better and better. Obviously, we can do things
 18 like develop new cures for diseases, but there's so much more — material science,
 19 new ways to produce clean and cheap energy, abundant energy — that if we are
 20 really on the precipice of being able to use AI to dramatically accelerate science,
 21 then I think we can solve some of the biggest problems in the world. And we now
 22 have one of the most — maybe the most — well-funded foundations in the world
 23 and the ability to use that amount of capital with the technology to go create a
 24 bunch of new scientific understanding for the world. I hope that'll be one of our
 25 greatest contributions ever.

27 XII. Liability, Deepfakes, and State vs. Federal Regulation

29 125. **Segall:** A big moment happened recently, which is a California jury found that
 30 Meta and Google were liable in damages for harmful and addictive platforms. It's
 31 the first time that social media has been treated as a defective product from a design
 32 perspective. I'm curious what the outcome of this case means for AI companies like
 33 OpenAI.

34 126. **Altman:** I mean, I think society is going to decide that creators of AI products
 35 bear a tremendous amount of responsibility for the products we put out in the
 36 world. But I kind of thought that before this case and I haven't reviewed this
 37 enough to say if there's like a specific thing that will apply to us.

38 127. **Segall:** And I think so much of the fear around AI and who has the power is
 39 based in — a grounded, you know, people worried that their jobs might not exist,
 40 their livelihood might not exist. I'd love to play a question from a mom who lives in
 41 Rhode Island. A big part of this show and what I want to do is be able to give
 42 people like that access to people like you to ask real questions.

1 [voicemail from Rhode Island mom plays]

2 128. **Rhode Island Mom:** Hi, I'm the mom of two middle schoolers in Rhode Island.
 3 And my kids are kind of coming around to the idea of having a backup plan for
 4 their careers. As most kids, they have aspirational careers of being professional
 5 sports players and fighter jet pilots. And we've been chatting about other options
 6 for careers. And I'm curious to see what careers you might recommend that would
 7 be AI-proof. My husband and I have been talking a lot about skilled trades. How,
 8 as a homeowner, it's difficult to find skilled trades people to come and help you out
 9 with electrical problems, with heating and cooling options. Maybe that's a great
 10 path for kids these days, and more traditional great career paths in engineering
 11 might take a backseat to these skilled trades. Obviously, kids will do what they
 12 gravitate toward and we encourage that. But if you have thoughts on AI-proof
 13 careers, I'd be interested to hear them. Thanks so much.

14 [voicemail ends]

15 129. **Altman:** One thing I will point out though in terms of things that I think are truly
 16 AI-proof is — I don't think anyone cares about the AI version of Lorie Segall. I
 17 don't think anyone wants to watch an AI come interview people like me. We are
 18 obsessed with other people. I don't think anyone — you know, she talked about her
 19 kids wanting to be sports stars and fighter jet pilots. Maybe fighter jet pilots society
 20 is happy to say, "Yeah, put the AI in there. That's safer and better." Okay. The
 21 sports stars — are we really going to watch like 11 robots play sports against each
 22 other? Some people will. I mean, a lot of that stuff goes viral. I saw like a robot
 23 falling, dancing. It has novelty, but like, are we going to get caught up in a sports
 24 league of just robots without the personal interest? Are we going to read AI-written
 25 novels and not have an author that we could understand their life story went into
 26 that? I would bet there's a huge cluster of — people care about people. And in the
 27 future, people will care even more about people. But this is so deep in evolutionary
 28 biology that those things are all quite AI-proof.

29 130. **Altman:** Someone said to me recently — I thought it was at the time, it seemed
 30 not that profound, more insightful the more I thought about it — in a world with
 31 true abundance, there is still one limited thing, which is human attention. And we
 32 care a lot about human attention. So, there's a whole category of stuff there. On the
 33 specifics, I'd go into it if we had more time, but like yes, clearly there will be a
 34 huge amount of short-term demand for electricians and skilled trades people.
 35 Eventually, I think the robots will do some of those things, but all of these things
 36 are just sort of like periods of time. If you just look at how much jobs have shifted
 37 every last decade, that will keep going.

38 131. **Segall:** You talk about human attention and so much of the last decade of tech
 39 was kind of really a business model that preyed on human attention. I'm curious
 40 how you think about the lessons of the last decade because now you're a for-profit
 41 company. And there is this idea that eyeballs on screen, more people using the
 42 product, is what makes the company successful. How do you not fall into the same
 43 business model trap that I would say the last decade did?

132. **Altman:** We decided to just totally shut down Sora. We were thinking about other versions of keeping it before the compute crunch came — like we were talking about putting it into the ChatGPT app, really focusing on generation and creativity. But one thing that we had realized is that to succeed with it as the product was currently conceptualized and sort of this way you could watch a lot of videos — that would have put a series of incentives on us and would have led to a bunch of decisions to win that we just didn't want to make. So, there was some version of that — and there still is some version of letting people generate video that I think is awesome. To compete in short-term video feeds and where that's going to go in a world of AI was not something I wanted to have to go do.

133. **Segall:** Right. It's an interesting decision of what you don't do and what that business model could demand. And there can be other versions of this with chat, like we could make a version of ChatGPT that was very addicting to talk to and that would have a bunch of complications too.

134. **Segall:** I want to play another voice memo from a guy named Joseph.

[voicemail from Joseph plays]

135. **Joseph:** Hey, my name is Joseph. I'm a recent college grad trying to work in film in New York City. And I have a couple questions for you. Number one, what do you dream about? Do you ever have recurring nightmares about AI? Do you often have dreams about AI? Number two, if you were in my shoes as like a young person trying to find themselves in a career or just trying to live in this world, what do you think is important to read and learn about right now and what skills are most useful? And then finally, what specific social causes do you care about? Thank you so much.

[voicemail ends]

136. **Segall:** It's a lot of them, so you can pick a couple. Pick your favorite one.

137. **Altman:** Thanks, Joseph. I can do them all fast. When I — until my early 30s, I had this recurring kind of bad dream of right when I dropped out of college to start my company. I would be in that period, but I was — in my dream, I was still in school and I was always missing something and I was always, you know, late for an exam or forgot about a class. And this went on till like deep into OpenAI time and by far the most frequent dream of my life. And it was so vivid. It was a little different every time, but it was like a kind of crazy, most intense dream. So I don't know what to make of that.

138. **Altman:** What I would study — I think the answer is just get really good at using AI tools more than any specific thing to study. And then the kind of skills that I expect to matter in the future — and I think these are learnable — like resilience, adaptability, kind of like calmness, being comfortable with a lot of change.

139. **Altman:** Social causes — I am very interested in this question of what the new version of the social contract looks like. You know, I don't think it's enough to say, "Hey, this jobs transition is going to be really hard but there's going to be a great time on the other side." I think we can say that and then we have to say, "And here

1 is our proposal for a new tax system and a new way to think about collective
2 ownership."

3 140. **Segall:** That makes sense. What would you propose?

4 141. **Altman:** I am interested in ideas where everybody is an owner in the magic of
5 capitalism. I think that the thing that happened with these new Trump accounts for
6 kids — that direction of stuff I really like. So, a world where we kind of preserve as
7 much of what's worked about capitalism of the last couple of centuries as we can,
8 but we say, "Hey, everybody just by virtue of being a citizen is going to own some
9 and we're going to have a tax system that enables that." I feel very excited about
10 that. I think as long as people can afford a good life and the price of admission,
11 they will outperform expectations.

13 XIII. Deepfake Pornography and the Regulatory Gap

15 142. **Segall:** Being in front of you and having spent so much time over the last four
16 years — we discovered there was this site on the internet where a man who was
17 anonymous for seven years prior had created hyperrealistic sexually explicit
18 deepfakes of women using artificial intelligence. It was horrific, looking at
19 deepfake porn and the way that these impacted women. And so much of that — I
20 think when we look at agency and control — I've talked to so many women, and it's
21 not just women, all sorts of victims who are victims of deepfake pornography —
22 and they felt like they lost their sense of agency. They had no control. They went to
23 — the laws hadn't caught up, so they went to the police. They said, "We don't even
24 know what you're talking about." I had a woman talk about wanting to end her life.
25 And so the human impact of this is so incredibly visceral. And I think about this
26 through the lens of: whose job is this? And granted, I say this with a caveat that
27 OpenAI has gotten in front of a lot of this stuff. We're not seeing a lot of these —
28 these were open source models. But I think that you're an important voice to be
29 able to weigh in on this type of thing.

30 143. **Altman:** I mean, that one seems like such a — that's not a hard decision to make,
31 right? Like that's so clear that we're not going to allow our models to be used for
32 that. Open source models, yes, people are going to use a lot of that.

33 144. **Altman:** I would kind of go back to the principles that I outlined earlier. You can
34 say a lot of bad things about the republic, but I think it has done a good job of
35 representing the will of the citizens and protecting against the tyranny of the
36 majority in that process. I think the Constitution remains one of the most amazing
37 governing documents the world has ever produced and I believe in it as a
38 governance system. So like with other technologies that have changed the world, I
39 hope that our systems, our governments, figure out how to run the process to decide
40 how to make these trade-offs about, say, freedom and protection and what's in the
41 best interest of society. Now, what I think you might say — and I would agree with
42 — is maybe that worked really well in a different time. Our leaders didn't do a great

1 job with the internet or with social media or whatever. And that is fair, which is
2 part of why we've been trying to talk about this more and earlier.

3 145. **Segall:** But if I could say specifically — so many of the folks, victims I've talked
4 to, finally had recourse when there were state laws. There had yet to be a federal
5 law. It takes a long time. I know OpenAI has looked at — supported curbing state
6 legislation in the name of innovation and "we have to move fast." So I'm curious if
7 you can explain.

8 146. **Altman:** This may be one we disagree on. I don't think state legislation will be
9 helpful here. I would much rather see the system urgently try to get a federal
10 framework right now. Maybe we try that for a couple of years and we can't and we
11 say, "Hey, it's just taking too long and we need something." But I don't think — I
12 think we need one stance on this. I mean, really there's an even trickier thing, which
13 is I think we need a global stance on a lot of this. A lot of these things are going to
14 affect the whole world. We've talked about things like bioterror, but certainly what
15 happens in one state will affect a lot of other states.

16 147. **Segall:** But in theory — but in reality sometimes the change happens quicker at
17 the state level.

18 148. **Altman:** I think even if it does it doesn't mean that that will have the impact on
19 the industry or the technology you would hope for. You know, if we take an area
20 like social where I think we'd both agree the government didn't really do their job,
21 public pressure did a lot. And the companies — you know, I think most people that
22 work at these companies probably want to do the right thing too. And the
23 companies, not every decision I would have made the same way, but I think they
24 made a lot of good ones. So it's not like we only have governments to rely on here.

26 XIV. Lightning Round

28 149. **Segall:** I want to do a quick lightning round. What will be the most valuable
29 human skillset that's AI-proof?

30 150. **Altman:** Caring about other people.

31 151. **Segall:** OpenAI's biggest opportunity and biggest missed opportunity.

32 152. **Altman:** Biggest opportunity in front of us is this sort of cluster of automating
33 researchers and companies and the super personal assistant for people's lives.
34 Biggest missed opportunity is we passed on a compute deal once we shouldn't have
35 passed on. Like a big one.

36 153. **Segall:** Like what?

37 154. **Altman:** I can't say exactly what, but there was like a very large compute deal we
38 could have taken.

39 155. **Segall:** How likely is OpenAI to acquire an entertainment company?

40 156. **Altman:** Not a top consideration right now.

- 1 157. **Segall:** How likely is OpenAI to go public this year?
- 2 158. **Altman:** Could happen. I don't know.
- 3 159. **Segall:** How likely is it that you could one day be replaced as CEO by an AI?
- 4 160. **Altman:** On the skills of my job? Very likely. On "will the world want a human
5 responsible for the decisions of a company like OpenAI?" I think the world will
6 demand that.
- 7 161. **Segall:** Do you say please and thank you to ChatGPT?
- 8 162. **Altman:** I do.
- 9 163. **Segall:** Why?
- 10 164. **Altman:** Force of habit.
- 11 165. **Segall:** Not because you're expecting, if it all goes bad, you want to be on the
12 other side?
- 13 166. **Altman:** No.
- 14 167. **Segall:** Last time you spoke to Elon Musk and what was the conversation?
- 15 168. **Altman:** It was a while and it was just some emojis back and forth.
- 16 169. **Segall:** Have you seen the fruit AI videos online?
- 17 170. **Altman:** No.
- 18 171. **Segall:** Oh, they represent AI slop at its finest.
- 19 172. **Segall:** When will we see AI wearables mainstream?
- 20 173. **Altman:** Two to four years. The reason I give such a broad range is right now if
21 you're talking to someone that's got like a forward-facing camera on their glasses,
22 it's incredibly off-putting. Maybe that gets more comfortable quickly or maybe we
23 just figure out a very different kind of wearable.
- 24 174. **Segall:** Can you give any sense of how you believe that will look? Is it like a pin
25 people are wearing?
- 26 175. **Altman:** Pin doesn't feel quite right to me. I mean, I hope — more as a general
27 thing, not just about wearables — I hope that the whole concept of products in the
28 age of AI fade into the background. Like the ideal AI assistant I want — there's
29 almost no product at all. It's an extremely smart model that's got full access to all
30 the context in my life that I can just trust. I can say one thing to and it's going to
31 happen. It'll be proactive with me only exactly when I want. And so when you think
32 about software products in that sense and also wearables, they could be quite subtle
33 or quite restrained.
- 34 176. **Segall:** The craziest thing coming down the pipeline when it comes to innovation
35 that you don't think we're talking enough about?
- 36 177. **Altman:** Probably automated researchers. Like I know we're talking about it
37 some, but if we weighted every conversation in this room relative to its importance

1 to the future, we probably would have talked way more about automating research
2 and way less about everything else.

3 178. **Segall:** You're saying one of the most important conversations we're not having is
4 about automated research? Like what does that mean for humanity?

5 179. **Altman:** You know, like 10 years of the whole world's scientific progress in one
6 year and then a hundred years of progress in one year. The transformation that will
7 have to quality of life, to the economy, to new risks. Like this is just — there's no
8 analog for this.

9 180. **Segall:** Do you think that's a possibility?

10 181. **Altman:** I think it's a possibility. Yeah.

11 182. **Segall:** What has to happen for it to happen?

12 183. **Altman:** The models need to get smarter. We need to eventually figure out how
13 to connect physical labs to them, things like that.

14
15 **XV. Closing — The Ethical Question of Our Lifetime**
16

17 184. **Segall:** I want to play one thing you said and then we'll wrap up. I'm excited for
18 this. What do you think is the single most important ethical question we need to ask
19 ourselves when it comes to the future of tech and humans in a world where we're
20 going to have computers that can think like humans?

21 185. **Altman:** What is the society we want to design? This question of what do we
22 want the role of humans to be in the world and how do we make sure the world is
23 good for humans in the broadest sense. I think that is the biggest ethical question of
24 our lifetime.

25 186. **Segall:** Do you agree with that?

26 187. **Altman:** I have nothing to add or change to that.

27 188. **Segall:** How do you think we're doing on answering that question?

28 189. **Altman:** Better than I — much better than I feared we would be doing. You
29 know, obviously not as well as I hoped, but I'm always like a little too optimistic.

30
31 Source: <https://www.youtube.com/watch?v=mJSnn0GZmls>

32 Date: April 2, 2026
33
