

Bias Tape: a multitrack tape machine model

Brian M. Gilmore, Symbia Labs · September 2026 · [PDF](#)

Abstract

Bias Tape is a real-time model of an analog tape machine built for a live insert host, where up to 22 tracks of a hardware mixer pass through it at once. Each track runs a physically motivated chain: standards-based record emphasis, a Jiles–Atherton magnetization model solved at twice the sample rate, complementary playback de-emphasis, playback-head losses, head bump, wow and flutter, and hiss. The contribution is the treatment of many tracks as one machine rather than many machines. Speed deviation is a pure function of the shared transport clock, so every track drifts sample-identically. Tape speed, formulation, format and noise belong to the machine and follow across tracks. A shared headstack lets each track pick up its neighbours at a level set by the format's track width.

The model was evaluated against predictions registered before measurement. Third-harmonic distortion reaches 3% within 0.2 dB of the target level for three tape formulations. One mono channel costs 0.57% of real time on a laptop CPU, and a silent channel costs 0.03%. Two predictions about level-dependent treble loss did not hold as registered and are reported as such. The model has not been compared with a physical machine.

1. Introduction

Tape emulation is usually built and judged one channel at a time. A plugin models one track of one machine, and a session that wants tape on every channel loads the plugin many times. The result is many independent machines. Their speed errors are uncorrelated, their settings can disagree, and nothing passes between them. A real multitrack recorder is the opposite on all three counts: one capstan moves one ribbon of tape past one headstack.

Bias Tape was written for Bias, a macOS insert host for the TASCAM Model 2400, a 22-input analog mixer and 24-track recorder with a multichannel USB interface. Bias takes each channel's send over USB, runs it through a chain of Audio Units, and returns it to the desk. Each strip has a dedicated tape

slot. A model in that slot runs on up to 17 strips (12 mono, 5 stereo) inside one audio callback, with a round trip of about 13 ms at a 128-sample buffer. That setting fixes the constraints: bounded cost per channel, no allocation or locks on the audio thread, a stated latency, and nothing that can leave a filter ringing on a strip with nothing plugged in.

The work was also a licensing decision. The best-known open tape model, CHOW Tape Model, is distributed under the GPLv3, which does not suit a closed commercial host. Bias Tape was therefore written from the published physics and from measurement papers, without reference to any existing implementation's source.

This paper makes three contributions:

1. A per-track signal path in which treble headroom follows from the playback equalization standards rather than from a tuned filter, using an exactly complementary emphasis pair around the magnetization model.
2. A multitrack formulation: speed deviation as a pure function of the transport's sample clock, machine-level settings shared across linked tracks, and a headstack that carries adjacent-track bleed at a level set by track width, with a result that does not depend on the order tracks are rendered in.
3. An evaluation in which every claim was registered as a prediction before it was measured, with the predictions that failed reported beside the ones that held.

2. Related work

Work on tape falls into four groups, and Bias Tape draws on all four while sitting in none of them.

Hysteresis models. Jiles and Atherton's 1986 model of ferromagnetic hysteresis describes magnetization as an anhysteretic curve that the material follows imperfectly, with domain-wall pinning producing the loop. It is the standard physical account and the basis of most analog tape models in audio. Chowdhury's 2019 DAFx paper applies it to tape with fourth-order Runge–Kutta at 16 times oversampling, simulating the high-frequency bias signal explicitly. Bias Tape uses the same equation, solved differently, and was written from the 1986 formulation rather than from any implementation of it.

Head and medium losses. Wallace's 1951 analysis of magnetically recorded signals gives the losses that dominate the treble: separation loss falling exponentially with the ratio of head-to-tape spacing to wavelength, thickness loss from the depth of the recorded layer, and gap loss from averaging across the reproduce gap. These are wavelength effects, so the same head is brighter at higher tape speeds. Bias Tape uses the textbook expressions directly.

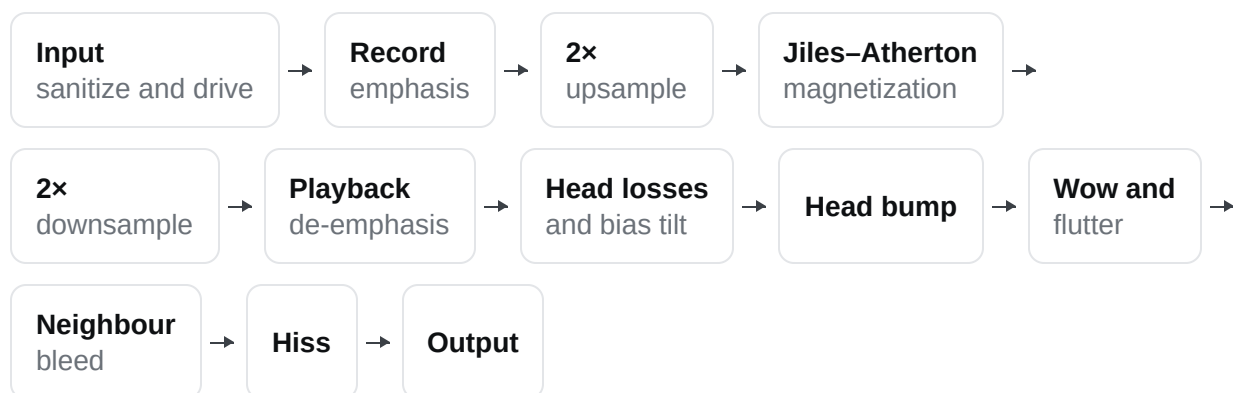
Measured and data-driven models. Mikkonen and colleagues' 2023 DAFx work measures a consumer reel-to-reel machine and fits a neural model to it, and reports the transport's real behaviour: delay-modulation depth, the frequency band the modulation occupies, and a noise floor that rises toward the low end. Those measurements informed the wow, flutter and hiss here. Arnardottir, Abel and Smith's 2008 AES model of the Echoplex tape delay is an earlier reference for treating capstan and pinch-wheel irregularity as a quasi-periodic modulation of delay. Bias Tape is not fitted to any measurement set; it borrows the reported magnitudes and shapes.

Equalization standards and bias. Adriaensen's 2025 Linux Audio Conference paper on tape simulation argues that the treble behaviour of a tape machine follows from the recording and playback equalization standards and a nonlinearity, and notes the playback time constants in use: NAB's 50 μs , IEC's 70 μs at 19 cm/s and 35 μs at 38 cm/s, AES's 17.5 μs at 76 cm/s, and the 120/70 μs pair for cassette. McKnight's work on biasing describes what AC bias does to linearity. This is the piece Bias Tape leans on most: the treble headroom is not a tuned effect but a consequence of pre-emphasis chosen by the standard for that speed and tape.

What is missing from all of these is the machine. Each models a channel. The multitrack formulation in Section 7 is, as far as we are aware, not treated in the audio literature, though the underlying physics — common transport, common headstack, crosstalk set by track width — is ordinary recording-engineering knowledge.

3. The signal path

One channel runs eleven stages, in the order a signal meets them on a real machine.



The order matters in three places. Emphasis comes before the nonlinearity and de-emphasis after it, so treble meets the tape hotter than the bass does and saturates sooner; this is what makes slow tape run out of treble headroom first, and it is a consequence of the standard rather than a separate effect. Oversampling brackets only the nonlinearity, because that is the only stage that creates new

frequencies. Wow and flutter come after the filters and before the noise, because a real machine's speed error modulates what has already been recorded but not the electronics' own hiss.

Drive and output are ramped across each block, so a moved control does not step. Input is clamped to ± 12 dBFS and anything that is not a finite number is replaced by zero before it reaches the tape, since the feedback in the hysteresis model and the recursive filters would otherwise carry a bad sample forever.

A strip with nothing plugged into it costs almost nothing. After one second below -90 dBFS — the noise floor of a desk with an open input — every filter in the chain has run out, so the channel parks: it clears its delay line and emits hiss alone until signal returns. Section 9 reports what that saves.

4. Magnetization

The tape itself is one Jiles–Atherton state per channel: magnetization M driven by field H , both dimensionless. Saturation magnetization and the Langevin shape parameter are both set to 1, so M lies in -1 to 1 and H is measured in units of the anhysteretic knee. The effective field is $H_e = H + \alpha M$ with $\alpha = 0.02$, and the anhysteretic curve is the Langevin function.

$$\frac{dM}{dH} = \frac{(1-c) \delta (M_{an} - M)}{(1-c)k - \alpha \delta (M_{an} - M) + c L'(H_e)} + c L'(H_e), \quad M_{an} = L(H_e) = \coth(H_e) - \frac{1}{H_e}$$

δ is the sign of dH/dt . The irreversible term is dropped whenever $\delta(M_{an} - M)$ is negative, the usual correction that keeps susceptibility from going negative at a turning point, and its denominator is floored at a quarter of $(1-c)k$ so the slope stays finite. The Langevin function and its derivative are tabulated over 0 to 16 in 4096 points with linear interpolation, with a series expansion near zero and the $1 - 1/x$ asymptote beyond the table; anything that is not a number falls into the asymptote branch and cannot reach the table index.

Two parameters carry the user-facing controls.

Formulation sets k , the pinning constant, which is the loop's width: 0.5, 0.65 and 0.9 for the three tape types, brighter and more forgiving to harder and cleaner.

Bias sets c , the reversible fraction. This is the paper's one interpretive move. AC bias is what makes tape linear, and in Jiles–Atherton terms a well-biased tape is one that follows its anhysteretic curve almost exactly, saturating only at the top. Aligned bias is therefore $c = 0.96$. Under-biasing opens the loop steeply (c falls to 0.40 at the extreme) so quiet signals go gritty; over-biasing has little

room left and moves only to 0.99, with the audible change coming from the treble tilt described in Section 5. Bias is not modelled as a high-frequency carrier added to the signal, which would demand far higher oversampling for no audible gain at these rates.

The equation is solved with midpoint Runge–Kutta at twice the strip's sample rate. A two-times half-band stage with a 47-tap Kaiser-windowed sinc brackets it. That is far less than the 16 times Chowdhury uses, which is needed there because the bias carrier itself is simulated; representing bias as reversibility removes that need. The choice was made for cost and checked by measurement rather than assumed.

Absolute field has no meaning until something fixes it, so each configuration is calibrated once, off the audio thread. A bisection of 18 geometric steps finds the field at which third-harmonic distortion of a 1 kHz tone reaches 3%, and full scale is then placed at the formulation's rated level for that figure: –5, –4 and –1 dBFS for the three types. The test is judged where a listener would hear it, after playback, so the probe tone is lifted by the record emphasis at 1 kHz and its third harmonic is weighted by the emphasis ratio between 1 kHz and 3 kHz. A second probe 30 dB down then sets a makeup gain so that quiet signals come back at the level they went in. Calibration is skipped when a setting changes that cannot move it, such as wow depth or output level.

5. Equalization and the playback head

Treble headroom is not tuned here. It follows from the playback equalization standard for the speed and tape in use, which sets how much hotter the treble goes onto the tape.

Speed	Standard	Type I	Type II / IV
1⅞ IPS (4.76 cm/s)	IEC cassette	120 μs	70 μs
3¾ IPS (9.5 cm/s)	double-speed cassette	70 μs	50 μs
7½ IPS (19 cm/s)	NAB	50 μs	50 μs
15 IPS (38 cm/s)	IEC	35 μs	35 μs
30 IPS (76 cm/s)	AES	17.5 μs	17.5 μs

The record side applies a first-order shelf with that time constant, limited to a gain of 6.3 (16 dB) so it plateaus rather than rising without bound:

$$H(s) = \frac{1 + sT}{1 + sT/G}, \quad G = 6.3$$

bilinear-transformed to a one-pole, one-zero section. Playback applies the same section with numerator and denominator exchanged. The pair is exactly complementary by construction, so a quiet signal comes back unaltered whatever the time constant; only what the tape did in between survives. An earlier version designed both sides as truncated FIR filters and they failed to cancel, leaving 15 IPS drooping 2.3 dB at 18 kHz. That was the reason for the change.

What the playback head loses is Wallace's separation and gap terms, with head geometry switched by speed: 0.2 μm spacing and a 1.0 μm gap for the cassette speeds, 0.5 μm and 2.5 or 4.0 μm for the studio speeds. Coating thickness loss is set to zero, because that is the loss the equalization standard exists to make up for; including it as well would take the treble out twice.

$$L(f) = \underbrace{e^{-kd}}_{\text{spacing}} \cdot \underbrace{\left| \frac{\sin(kg/2)}{kg/2} \right|}_{\text{gap}}, \quad k = \frac{2\pi f}{v}$$

Since k depends on tape speed v , one head is brighter at 30 IPS than at $1\frac{7}{8}$ IPS without any parameter changing. Two more terms ride the same filter: a tilt of -4 dB per unit of bias above 10 kHz, so under-biased tape is brighter and over-biased tape erases its own treble; and a correction for the half-band filters' own droop near the top of the band. The three are multiplied and realised as one 64-tap minimum-phase FIR, designed by folding the real cepstrum of the wanted magnitude through a 1024-point DFT and fading the last quarter of the taps with a raised cosine.

Head bump is two biquads: a $+2$ dB peak at 4.5 times the speed in IPS — 17 Hz at $3\frac{3}{4}$, 135 Hz at 30 — and a second-order high-pass at 45% of that frequency, which also removes the DC offset a magnetized tape leaves behind.

6. Transport and noise

Speed error is a displacement, not a modulation applied per channel. The model computes how far the tape is from where it should be, in seconds, as a closed-form function of absolute transport time:

$$e(t) = - \sum_i \frac{w_i W \cos(2\pi f_i t + \phi_i)}{2\pi f_i} - \sum_j \frac{v_j F \cos(2\pi g_j t + \psi_j)}{2\pi g_j} + (\text{smooth noise terms})$$

W and F are the peak speed deviations the Wow and Flutter controls set. Wow uses components at 0.55 Hz and 1.37 Hz plus a slow random walk; flutter uses the capstan rotation rate, a partial at 2.63 times it, and a faster random term. Capstan rate is taken as 5.8 to 24.2 Hz across the five speeds. Each division by $2\pi f$ is the integration from speed error to displacement, which is why wow moves pitch far more than flutter at the same percentage. The random terms are a hash of the sample index with smoothstep interpolation, so they too are a function of t alone and not of a running state.

That last property is what makes the transport shared. Every instance on the machine evaluates the same $e(t)$ from the same host timestamp, so all tracks drift together, sample for sample, with no communication between them. There is no leader and no shared buffer of speed values. Measurement T5 confirms two stages on one clock produce bit-identical output.

The excursion is read from a delay line by four-point Lagrange interpolation, with a base delay of the peak excursion plus three samples so the read point can swing both ways. That base delay is added to the reported latency, and it glides rather than steps when the depth control moves.

Hiss is white noise from an xorshift generator plus the same noise through a one-pole low-pass at 150 Hz at half the amplitude, the sum normalized to unit RMS; the result is flat with a rise at the bottom, which is where a measured machine's noise sits. Its level falls 3 dB for each doubling of tape speed and 3 dB for each doubling of track width, so the Hiss control reads as a figure for one cassette track at $3\frac{3}{4}$ IPS and the model puts a 2-inch 24-track machine at 30 IPS about 12 dB quieter for the same setting.

7. One machine, many tracks

Every instance in the host is a track of a machine, not a machine of its own. Two machines exist: the multitrack that the channel strips record to, and the mixdown deck the multitrack is mixed to. A track joins one of them when the host loads it, and carries a Link switch that takes it off the machine and gives it its own settings.

Settings divide by what they physically belong to. Format, Speed, Tape, Wow, Flutter and Hiss belong to the machine: set one on any linked track and every linked track follows, because one transport cannot run at two speeds and one reel cannot hold two formulations. Input, Bias trim and Output belong to the track, because those are alignment and gain staging, set per channel on a real machine too. A track arriving with saved state sets the machine from it; a track arriving fresh takes the machine as it finds it. The machine keeps its settings with no tracks on it, so the next track to load finds the machine as it was left.

Format sets what the tape is and how wide a track on it is, which in turn sets two things the user never adjusts directly.

Format	Track width	Neighbour bleed at 1 kHz
Cassette 4-track	0.6 mm	−45 dB
¼" 8-track	0.5 mm	−45 dB
½" 8-track	1.0 mm	−50 dB
1" 16-track	1.0 mm	−50 dB
2" 16-track	1.8 mm	−60 dB
2" 24-track	1.1 mm	−55 dB
¼" 2-track	2.0 mm	−55 dB
½" 2-track	5.0 mm	−60 dB

Narrow tracks are noisier and leak more. The 2-inch 24-track is the interesting case: it is a wide tape, but 24 tracks across it are narrower than 16 on the same tape, so it is noisier and leakier than the 2-inch 16-track despite the format sounding grander.

The headstack

Tracks hear their neighbours through a shared structure indexed by the transport's sample count: 24 lines of one block each, plus the sample position each line has been written up to. A track writes its finished block at position t and reads the blocks of tracks above and below it covering $[t - n, t)$, one block old.

That one-block lag is the mechanism that makes the result independent of render order. The host does not guarantee which strip's callback runs first, and a same-block read would give a track its neighbour's current block or its previous one depending on that order, making output non-deterministic. Reading strictly one block behind means every track sees the same thing however they are scheduled. The cost is one buffer of delay on the bleed path, which at 128 samples is under 3 ms and inaudible against a signal already 45 dB down.

The bleed itself is the neighbours summed, then a gain plus three times a one-pole low-passed copy at 200 Hz, normalized so the table's figure is the level at 1 kHz. Crosstalk on a real headstack rises

toward the bottom, where wavelengths are long enough to reach across the guard band; here it is about 10 dB stronger at 60 Hz than at 1 kHz. A track with nothing on it still plays its neighbours' bleed, as a real one does. An unlinked track neither sends nor receives.

Stereo strips get this for free: their two channels are adjacent tracks, so left and right bleed into each other at the format's level.

8. Implementation

The model is about 900 lines of Swift with Accelerate, in two files: the DSP, which depends on nothing but Accelerate and can be measured offline, and an Audio Unit wrapper.

The wrapper is an `AUAudioUnit` subclass registered into the engine's own process at startup, with a component type of its own. It is not installed on the Mac and no other application can see it. The reason is uniformity: the host already drives plugins through the version-2 render interface and already handles parameter trees, factory presets, full-state save and restore, metering and editors through that one path. Registering in-process puts the built-in model on exactly that path, so nothing in the host has a special case for it. The model appears in the plugin list beside 300-odd installed Audio Units, loads into a slot the same way, and its state travels in a saved song the same way.

Real-time safety is handled by keeping everything the render thread touches fixed. Buffers are allocated when the unit is prepared, sized for the largest block. Anything derived from a setting — the calibration bisection, the minimum-phase FIR design, biquad coefficients — is computed off the audio thread into an immutable configuration object, which the render block picks up by swapping an unmanaged pointer at the top of a block. Superseded configurations are held in a list for several seconds before release, so the audio thread cannot dereference freed memory even if it is mid-block when a control moves. Parameters that cannot affect calibration reuse the previous calibration outright.

Parameters are presented in two groups, Track and Machine, matching the division in Section 7, which lets the host's generic editor draw the distinction without knowing anything about tape. Eight parameters, six factory presets, and a state dictionary that carries the settings plus the current preset.

Reported latency is 23 samples from the two half-band filters, plus the wow-and-flutter base delay when either is non-zero. With the default 0.06% wow and 0.04% flutter at 44.1 kHz that comes to about 0.8 ms.

9. Evaluation

Every claim below was written down as a prediction, with the figure it had to hit and what would count as failing, and lodged in a signed record before the measuring code was run. Results are reported against those predictions, including the two that did not hold.

Two harnesses measure the model. One drives the DSP class directly, offline and deterministically. The other loads it as the host loads any plugin, through the version-2 render interface, and checks that path. Figures below are from a single run of both on an M-series laptop at 44.1 kHz with a 128-sample block.

Single-value claims

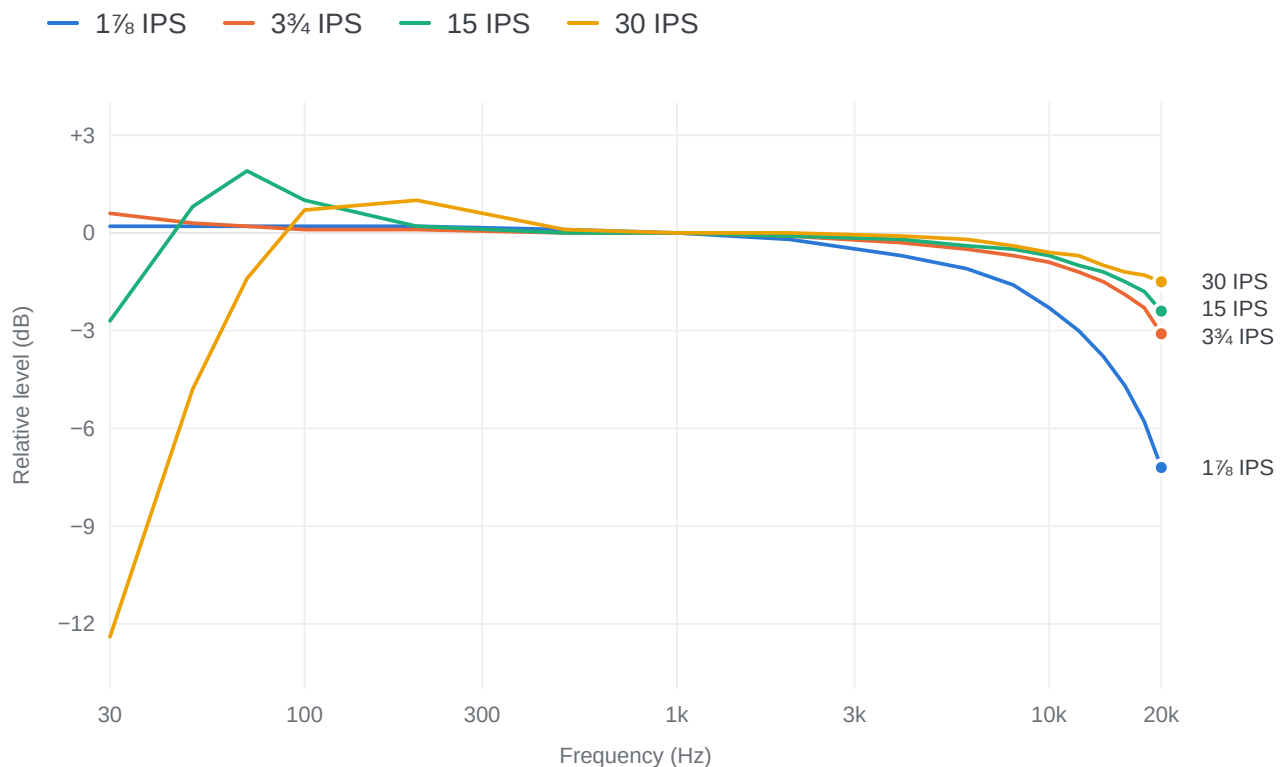
Measurement	Predicted	Measured	Held
Gain at 1 kHz, -30 dBFS	within 1 dB	-0.12 dB	yes
3% third harmonic, Type I	-5 dBFS $\pm$ 1.5	-4.88 dBFS	yes
3% third harmonic, Type II	-4 dBFS $\pm$ 1.5	-3.84 dBFS	yes
3% third harmonic, Type IV	-1 dBFS $\pm$ 1.5	-0.80 dBFS	yes
Wow depth at 0.2% setting	0.15–0.25%	0.188%	yes
Two stages on one clock	bit-identical	identical	yes
Hiss at -60 setting	-60 dBFS $\pm$ 1.5	-59.95 dBFS	yes
Output under abuse	finite, under +12 dBFS	finite, +5.29 dBFS peak	yes
Settled DC after a burst	under -80 dBFS	-107.5 dBFS	yes
Latency, transport steady	23 samples $\pm$ 2	23 samples	yes
Cost, one mono channel, 10 s	under 60 ms	61.8 ms	no
Host render path vs direct	identical samples	0 of 44,032 differ	yes
State round trip	values and preset equal	equal, 591 bytes	yes

The cost claim failed by 3%, on a machine that was compiling two applications at the time; earlier runs on an idle machine measured 57 to 58 ms. It is recorded as a failure because that is what the run produced. At 0.62% of real time per mono channel, 27 channels of tape would take about 17% of one core.

Frequency response

Response by tape speed

Type I at -30 dBFS, relative to 1 kHz. Head bump is the rise below 200 Hz.



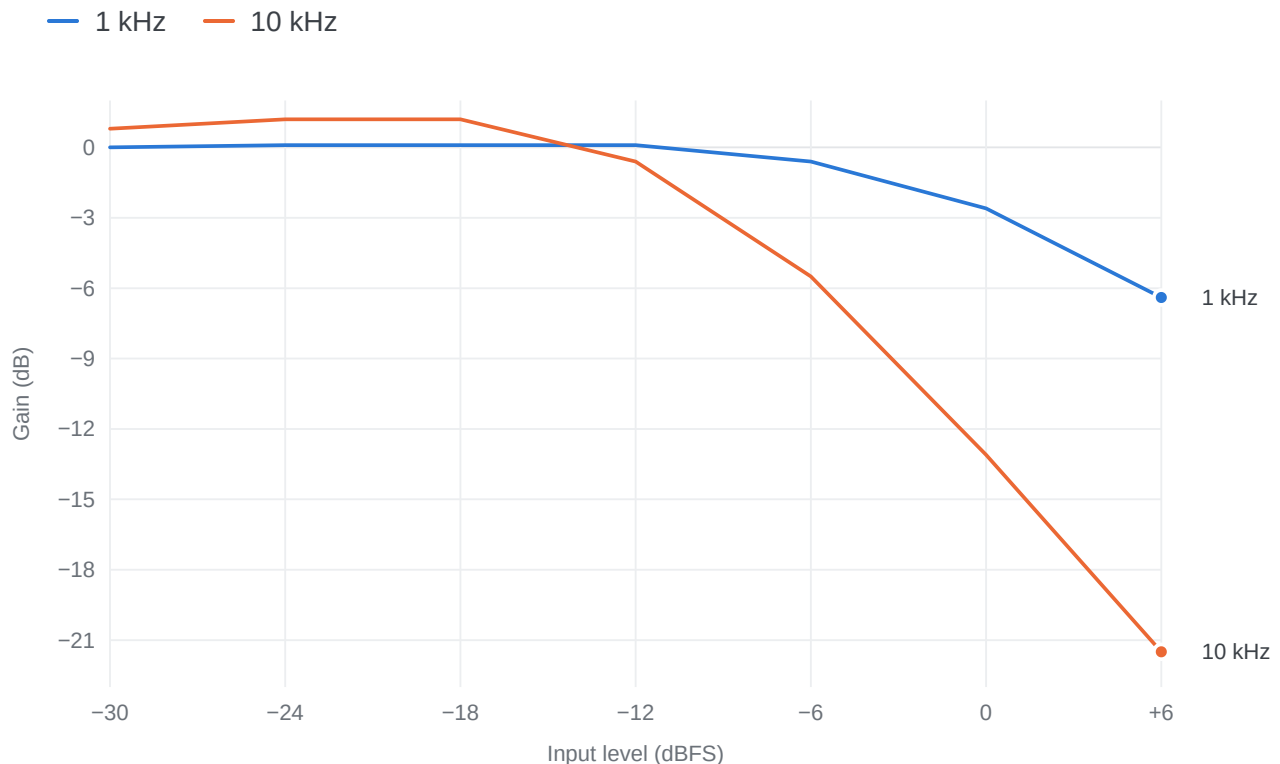
tapecheck T3 · 4 speeds × 17 frequencies

The four curves differ only by tape speed; no filter parameter changes between them. Slow tape loses treble and fast tape does not, because head losses are wavelength effects and the equalization standard for each speed sets how much was put on in the first place. At 1 $\frac{7}{8}$ IPS the response is 3 dB down at 14 kHz; at 15 and 30 IPS it holds to 20 kHz. The low end shows the other side of the same mechanism: head bump sits at 4.5 times the speed, so it is a gentle lift around 17 Hz at 3 $\frac{3}{4}$ IPS and a pronounced 1 dB shelf near 200 Hz at 30 IPS, with the machine's own high-pass taking the bottom out below it.

Level dependence, and what broke

Compression by level

Type I at $3\frac{3}{4}$ IPS, aligned bias, gain relative to the same tone at -40 dBFS.



tapecheck level sweep · 7 levels × 2 frequencies

The registered prediction was that at $3\frac{3}{4}$ IPS a 10 kHz tone at -10 dBFS would compress at least 3 dB more than a 1 kHz tone at the same level, and that at 30 IPS the two would stay within 1.5 dB of each other. Neither held. Measured at -10 dBFS the difference at $3\frac{3}{4}$ IPS is 1.97 dB, short of the 3 dB claimed; at 30 IPS it is 1.67 dB, just outside the 1.5 dB bound and in the wrong direction.

The chart shows why. Above -6 dBFS the shape is right and strong: at 0 dBFS the 10 kHz tone is 13.1 dB down while 1 kHz is 2.6 dB down, a 10 dB spread. The problem is below that. Between -30 and -18 dBFS the hysteresis model expands treble by about 1 dB rather than leaving it alone, so the crossover from expansion to compression happens later than the prediction assumed, and a measurement taken at -10 dBFS catches the model mid-turn. This is a real artefact of the Jiles–Atherton parameters at this reversibility, not a measurement error, and it is the clearest open problem in the model.

The machine

A second track loading onto a machine took all six machine settings from the track already there and kept its own input trim. Setting speed on either track changed both. With a tone on one track and silence on the other, the silent track's output measured 42.8 dB below the loud one at 1 kHz on Cassette 4-track against a table figure of 45 dB, the difference being the loud track's own -2 dB output trim, and 33.3 dB below at 60 Hz, showing the low-end rise. Switching the machine to 2-inch 24-track from either track moved both, and bleed fell to -52.8 dB. Unlinking a track dropped its bleed to -330 dB, which is silence, and left it holding its own speed while the other track's changed.

10. Limitations

No reference machine. Nothing here has been compared against a physical tape recorder. The targets the model is calibrated to — distortion levels by formulation, equalization time constants, head spacings and gap widths, crosstalk figures by format — come from published standards and textbook values, not from measuring a deck. Every figure in Section 9 says the model does what the model was told to do. That is a weaker claim than it may appear.

The mid-level treble artefact. Between -30 and -18 dBFS the model expands treble by about 1 dB instead of leaving it flat, which broke both halves of the level-dependence prediction. The likely cause is the interaction of the reversibility parameter with the floor on the irreversible term's denominator, which was added for numerical stability and has an audible cost. Fixing it without losing the distortion calibration is the main open problem.

Bias is an abstraction. AC bias appears as a reversibility parameter and a treble tilt, not as an actual high-frequency carrier. The model therefore misses what bias does to short wavelengths directly, and the trade-off between bias level, sensitivity and maximum output is only approximated by the two controls' combined effect.

Absent effects. Print-through, dropouts, azimuth error, scrape flutter, and the difference between the record and playback heads are all unmodelled. The last of these matters most for the machine: without a sync and repro distinction there is no way to represent monitoring off the record head during a bounce, which is the thing a multitrack machine is actually for.

Crosstalk is simplified. Bleed reaches only immediate neighbours, is symmetric, and has one frequency shape. A real headstack couples further than one track, and couples differently in record than in playback.

The mixdown machine is unreachable. The model provides for a second machine for the stereo mix, but the hardware's main mix has no USB return, so no track can be on it today.

11. Future work

The near work is in three parts. First, resolve the mid-level treble artefact and re-register the level-dependence predictions rather than restating them. Second, a sync and repro head switch, which would give bounce and generation loss an honest representation, since a signal already passes through the model twice when a track is bounced. Third, dropouts and creases shared across tracks, which the headstack already makes possible — a crease hits every track at the same instant, and that is a multitrack effect no single-channel plugin can produce.

The measurement work is to get access to a real machine. A recorded sweep, a distortion series and a flutter capture from one deck would turn most of Section 9's claims from self-consistency checks into comparisons.

12. Conclusion

The per-channel model here is conventional in its parts: Jiles–Atherton magnetization, Wallace's head losses, a delay-modulated transport. Two choices in how those parts are arranged are worth carrying elsewhere. Taking treble headroom from the playback equalization standard, through an exactly complementary emphasis pair around the nonlinearity, gives the level- and speed-dependent treble behaviour without a filter tuned to taste. And calibrating absolute field by bisecting to a published distortion figure makes the tape formulation a measured target rather than a set of dials.

The multitrack formulation is the part that does not exist elsewhere. Making speed deviation a closed-form function of the shared transport clock costs nothing and buys correlated drift across every track with no communication. Indexing the shared headstack by transport sample count, and reading one block behind, makes crosstalk deterministic under any render order. Tying track width to the tape format ties both noise and crosstalk to a single choice a user already understands. Together these turn a rack of identical plugins into one machine, which is what a multitrack recorder is.

What the model has not yet earned is a comparison with a real deck. Until then its measurements say only that it is self-consistent.

Disclosure

The model, its measurement harnesses and this paper were produced with Claude, an AI system made by Anthropic, working under the author's direction in a single extended session. The author set the goals and made the design decisions, including treating many tracks as one machine; the code and text were written by the AI and reviewed by the author. The prediction-before-measurement

records were kept in a signed Symbia Labs catalog. No part of any existing tape plugin's source was consulted.

References

1. D. C. Jiles and D. L. Atherton, "Theory of ferromagnetic hysteresis," *Journal of Magnetism and Magnetic Materials*, vol. 61, pp. 48–60, 1986. doi:[10.1016/0304-8853\(86\)90066-1](https://doi.org/10.1016/0304-8853(86)90066-1)
2. R. L. Wallace, Jr., "The reproduction of magnetically recorded signals," *Bell System Technical Journal*, vol. 30, no. 4, pp. 1145–1173, October 1951. [index](#)
3. J. Chowdhury, "Real-time physical modelling for analog tape machines," *Proc. 22nd International Conference on Digital Audio Effects (DAFx-19)*, 2019. [pdf](#)
4. O. Mikkonen, A. Wright, E. Moliner and V. Välimäki, "Neural modeling of magnetic tape recorders," *Proc. 26th International Conference on Digital Audio Effects (DAFx 2023)*, 2023. [arXiv:2305.16862](https://arxiv.org/abs/2305.16862)
5. S. Arnardottir, J. S. Abel and J. O. Smith III, "A digital model of the Echoplex tape delay," *Audio Engineering Society 125th Convention*, paper 7649, 2008. [AES](#)
6. F. Adriaensen, "Tape simulation," *Linux Audio Conference (LAC 2025)*, INSA Lyon, June 2025. Read from a copy of the paper; no public link found.
7. J. G. McKnight, "Biasing in magnetic tape recording," 1967. Publication details not verified.