

Criação musical generativa baseada na consonância entre ondas: uma abordagem não-psicoacústica

Gabriel D’Incao², Artemis Moroni¹

¹Divisão de Sistemas Ciberfísicos – DISCF
CTI Renato Archer – Campinas/SP

²Universidade Estadual de Campinas - UNICAMP

g253628@dac.unicamp.br, artemis.moroni@cti.gov.br

Abstract. *This article presents part of the conceptual development of a generative algorithm for musical creation, grounded in an approach to consonance that is not anchored in psychoacoustic principles but in a simple methodology based on the periodicity between waves. Following a critical historical review of the subject, the author details the process of generating the system’s melodic paths, which is structured around the filtering of sets according to the proximity between frequencies.*

Resumo. *Este artigo apresenta parte do desenvolvimento conceitual de um algoritmo generativo para criação musical fundamentado em uma abordagem de consonância que não está ancorada a princípios psicoacústicos, mas a uma simples metodologia baseada na periodicidade entre ondas. Após uma crítica revisão histórica do tema, o autor dissecou o processo de geração dos percursos melódicos do sistema, que é estruturado na filtragem de conjuntos a partir da proximidade entre frequências.*

1. Uma breve contextualização histórica

A relação entre ondas despertou o interesse de muitas grandes mentes da ciência ao longo da história. Essas pessoas estudaram o fenômeno físico/psicoacústico da consonância através da linguagem matemática, utilizando o som como a materialização de suas pesquisas [Sethares 2005]. Tais investigações influenciaram diretamente o desenvolvimento da cultura musical, uma vez que grande parte dos instrumentos harmônicos construídos utilizaram sistemas de afinação que, em algum nível, se apoiaram em fundamentos teóricos resultantes dessas pesquisas. Um exemplo fundamental desse fenômeno foi a contribuição de Pitágoras, um filósofo grego da região da Jônia que, em algum momento próximo a 500 A.C., construiu um sistema de afinação baseado na geometria de cordas tensionadas.

Em seus estudos, percebeu que um segmento de comprimento λ , e outro com $\frac{1}{2}\lambda$, quando excitados, soavam de alguma forma semelhantes, mas com a menor sendo percebida como aguda em relação à maior. Esse fenômeno, mais tarde denominado como o intervalo perceptivo de uma "oitava", se repete também quando o comprimento é reduzido a $\frac{1}{4}\lambda$, que soa ainda mais agudo. Isso nos revela que a altura percebida de uma onda periódica é inversamente proporcional ao tamanho do corpo (com tensão e massa constantes): $f(\lambda) = \frac{1}{\lambda}$, o que, em termos atuais, evidencia a relação logarítmica do processo perceptivo do fenômeno sonoro.

Seu sistema de afinação foi construído à partir da manipulação de dois intervalos: $\frac{2f}{f}$ e $\frac{3f}{2f}$, ou seja, tamanhos de corda $\frac{1}{2\lambda}$ e $\frac{2}{3\lambda}$, tipicamente reconhecidos como "oitava" e "quinta justa". Seu resultado sonoro é tipicamente percebido como "consonante", uma vez que o material de sua construção parte de razões de números pequenos. No entanto, esse fenômeno perceptivo ocorre em sua forma máxima somente quando o centro da abordagem é o f_0 inicial gerador do ciclo de quintas [Tsuji and Müller 2021].

Muitos outros sistemas de afinações foram propostos e aplicados para determinar as possíveis frequências emitidas pelos instrumentos musicais, e apesar de cada tentativa carregar suas distintas características e qualidades, todos - até recentemente - possuíam o mesmo defeito de não se manterem em sua forma ideal (com intervalos mais consonantes) em diferentes centros tonais, o que dificultava modulações muito distantes. Isso ficou mais evidente quando compositores como Bach passaram a expandir o tonalismo, e exigir mais mudanças de centro tonal em suas peças. O resultado era que cada tonalidade possuía uma certa cor, o que não era uma característica necessariamente vista como negativa na época, mas uma que restringia o processo criativo a notas com relações intervalares não-simétricas.

Neste contexto, surge no século XVI um sistema de afinações como uma solução teórica para o problema da simetria dos intervalos musicais: o sistema do temperamento igual, este que é construído à partir de uma única unidade, o semi-tom = $\sqrt[12]{2}$. A proposta parte do princípio de dividir a "oitava" ($2f$) em 12 sons que soam perceptivelmente equidistantes. Dessa forma, seria teoricamente possível transpor qualquer material sonoro para qualquer outro centro tonal, mantendo as relações intervalares intactas. O desenvolvimento desse trabalho ao longo dos séculos permitiu com que instrumentos com alturas fixas, como o cravo e posteriormente o piano, pudessem tocar em conjunto com grandes grupos de instrumentos de altura indefinida, como o violino e a voz. Pode-se dizer que formação da orquestra sinfônica como conhecemos hoje, assim como o amplo desenvolvimento do tonalismo, e consequentemente sua negação no século XX, se devem, entre muitas outras questões, ao surgimento desse imperfeito sistema.

1.1. Números físicos e digitais

No entanto, com seu estabelecimento, e com o gradual desaparecimento dos outros sistemas, acabamos culturalmente acostumando com a sonoridade de seus intervalos, que são aparentemente próximos dos ideais propostos por Pitágoras, mas que na verdade são incalculavelmente distantes, uma vez que $\sqrt[12]{2}$ é um número irracional. Ou seja, nascemos em um contexto cultural acostumado com a sonoridade de intervalos que jamais serão alcançados, sempre aproximados, uma vez que é impossível afinar perfeitamente dois corpos de forma que emitam um intervalo irracional. E sim, isso também é o caso para qualquer outra afinação, já que corpos reais não são fisicamente perfeitos, portanto jamais poderiam atingir tamanha perfeição intervalar, mas o sistema do temperamento igual busca aproximar-se não de razões simples, mas de um intervalo praticamente incomputável.

Pode-se dizer que este não é um problema tão grave no contexto prático do fenômeno musical, uma vez que não somos capazes de facilmente perceber diferenças tão pequenas de frequência, e que mais importa a praticidade e a conveniência da flexibilidade tonal, a se limitar a um arbitrário e igualmente inacessível rigor matemático para

construção de um sistema de afinação. No entanto, quando analisamos a questão considerando a aplicabilidade destes sistemas em ambientes digitais, o problema se torna mais palpável.

Em um computador, o sistema temperado soa particularmente mal. Nele, não existem as inflexões naturais resultantes da intuição humana para compensar suas imperfeições [Duffin 2007], ele apenas aproxima $\sqrt[12]{2}$ da melhor forma possível, cálculo este que possui precisão limitada à quantidade de bits. O resultado é a aproximação imperfeita de um intervalo imperfeito, e seus defeitos são facilmente percebidos quando comparados a frequências geradas por outros métodos.

O mesmo não é verdade para a afinação pitagórica, por exemplo, que é constituída de números e operações matemáticas facilmente computáveis com perfeição usando a tecnologia atual. Ao mesmo tempo, no contexto digital, não estamos mais amarrados às limitações práticas da afinação de instrumentos com um número finito de cordas, portanto é possível construir um sistema que gera material harmônico que mantém as qualidades intervalares mesmo em distintos centros tonais, garantindo o rigor matemático sem sacrificar a possibilidade de jogar com processos modulatórios na prática composicional.

O presente artigo explica o funcionamento de um algoritmo generativo que gera caminhos de frequências de forma a explorar este princípio de colocar as relações harmônicas entre as vozes em primeiro lugar, mas partindo de uma abordagem não-psicoacústica do fenômeno da consonância. Nele, será possível construir caminhos melódicos que não estão limitados a um conjunto fixo de frequências, mas que consideram as propriedades energéticas e periódicas do fenômeno das ondas, de forma objetiva, e desconsiderando as limitações físicas do aparelho auditivo.

2. "Consonancia"

Apesar do desenvolvimento milenar do conceito de consonância, este é um tópico que nunca foi resolvido pela ciência, ou seja, não há um modelo que é aceito universalmente. Isso se deve principalmente ao fato de que este termo está tipicamente relacionado ao fenômeno perceptivo da relação entre ondas sonoras, o que dificulta uma quantificação objetiva do conceito, por parecer, à primeira vista, de se tratar de um campo subjetivo. Ainda assim, muita pesquisa foi desenvolvida nos últimos séculos com o intuito de formalizar a relação entre a sensação de dissonância e as proporções de frequências, à partir de análises laboratoriais sob uma abordagem psicoacústica.

Experimentos descritos no artigo "Tonal Consonance and Critical Bandwidth" nos revelam que a sensação de dissonância entre sons puros (senoides) está diretamente relacionada com a distância linear entre suas frequências [Plomp and Levelt 1965]. Os resultados da pesquisa, fortemente influenciada pelas ideias de Helmholtz [Helmholtz 1863], mostram que a dissonância cresce até atingir um pico quando a separação está dentro da largura de banda crítica em torno da frequência média. Esse ponto de máxima dissonância foi posteriormente interpretado como ocorrendo quando a separação entre as frequências corresponde a aproximadamente um quarto da largura de banda crítica do sistema auricular em torno da frequência média:

$$D_{\text{sensorial máximo}} \approx \frac{1}{4} CBW\left(\frac{f_0+f}{2}\right)$$

Abaixo desse valor, a sensação de "rugosidade" decresce até o ponto em que a diferença

é percebida como um batimento lento, descrito como “agradável”. O decaimento da dissonância também ocorre à medida que Δf cresce além desse valor, até que as duas ondas passam a ser percebidas como notas distintas.

William A. Sethares, em seu livro “Tuning, Timbre, Spectrum, Scale”, propõe um método para a quantificação da dissonância, baseando-se no processo de análise descrito acima. Este cálculo é feito ao somar a “rugosidade” de todas as relações intervalares em um grupo de senóides, e usa desse processo para argumentar que qualquer conjunto de parciais, incluindo timbres inarmônicos, soam “consonantes” se a diferença linear entre cada frequência não estiver próxima a este ponto de máxima dissonância sensorial. Ele parte deste princípio para propor a construção de escalas à partir do espectro, e espectros à partir de escalas [Sethares 2005].

Embora essa perspectiva psicoacústica seja valiosa para entendermos o que nos soa “agradável”, penso que a análise do fenômeno Consonância não deveria ser fundamentada apenas nas limitações biológicas de uma única espécie, uma vez que este é um assunto que engloba áreas do conhecimento que se sustentam independentemente da ínfima existência humana. Por essa razão, e pelo ingênuo desejo que tenho de contribuir com o desenvolvimento da ciência fora da bolha musical, utilizarei uma metodologia alternativa à visão psicoacústica para analisar qualitativamente a relação intervalar entre frequências, baseando-me em um conceito muito mais antigo.

2.1. A periodicidade entre ondas

O método para escolha de frequências utilizado pelo algoritmo descrito neste trabalho parte de uma visão minimalista, que considera o nível de dissonância entre duas ou mais ondas como o período da onda resultante. Ou seja, quanto maior a quantidade de ciclos necessários para que a onda com menor frequência coincida em fase com as outras, maior a dissonância.

Para visualizar este fenômeno, imagine um lago absolutamente calmo, sem vento ou ondulação própria, um plano d’água perfeito. Nesse espaço ideal, duas torneiras a e b , suspensas acima da superfície, pingam gotas de água gerando ondas circulares que se espalham indefinidamente. Ambas ocupam o mesmo ponto no espaço, mas cada uma libera suas gotas em uma frequência distinta. Se a razão entre a e b for de 1:2, e ambas começarem a pingar no mesmo instante, o ponto máximo de interferência construtiva será na crista de a , ou seja, o período T da soma dessas ondas é idêntico ao de $a \rightarrow T_{a,b} = T_a$. O mesmo ocorre em todas as proporções onde B possui frequência múltipla de A , portanto o conjunto: [1:1, 1:2, 1:3, 1:4, ...] representa intervalos com o mesmo “nível” de dissonância $D = 1$, que é o mínimo possível.

Agora, se a razão entre as frequências de a e b for 2:3, o ponto culminante das cristas ocorrerá a cada dois ciclos de a , portanto o período será $T_{a,b} = 2T_a$. Note que 2:4 é na verdade 1:2, então o conjunto que contém as proporções de $D = 2$ é: [2:3, 2:5, 2:7, ...], e consequentemente $D = 3$: [3:4, 3:5, 3:7, 3:8, 3:10...]. Portanto, sob esse método, o nível de dissonância D de qualquer relação intervalar é simplesmente o menor número da proporção em sua forma mais simplificada. Neste caso, perceba que há uma hierarquia de consonância construída à partir da distância entre os picos da onda resultante, e o intervalo 3:2, tipicamente conhecido como a quinta justa, é considerada mais dissonante que 3:1, uma décima segunda (quinta justa oitava acima).

Para calcular o valor D em função de qualquer grupo de frequências racionais com tamanho n , devemos primeiramente representar cada frequência como fração em sua forma mais simples. Em seguida, deve-se transformar o conjunto em um com apenas números naturais, multiplicando os valores pelo mínimo múltiplo comum dos denominadores. Por fim, divide-se a menor frequência pelo máximo divisor comum de todos os elementos [Euclides et al. 2009]:

$$D(f_1, \dots, f_n) = \frac{\min\{f_1, \dots, f_n\}}{mdc(f_1, \dots, f_n)}$$

Sob essa ótica, o intervalo gerador do sistema temperado, $\sqrt[12]{2}$, apresentaria nível de dissonância essencialmente infinito, pois, sendo um número irracional, não pode ser expresso como razão de inteiros e, conseqüentemente, não é suscetível a simplificação. No lago ideal, uma proporção que utilizasse esse valor resultaria em um movimento eternamente caótico que nunca se repetiria. Esse tipo de problemática não é considerada pela visão estritamente psicoacústica, e por essa razão, o sistema descrito abaixo não o considera. No entanto, uma combinação das duas abordagens é possível, e será incorporada em um futuro desenvolvimento, de forma a tornar o algoritmo perceptivamente inclusivo.

3. Relação entre vozes

A relação de consonância entre as vozes é o motivo central da abordagem generativa do presente trabalho. Ela ocorre individualmente, ou seja, em um sistema com 3 vozes A, B, C, haverá relações únicas para $[A \rightarrow A]$, $[A \rightarrow B]$, $[A \rightarrow C]$, $[B \rightarrow A]$, $[B \rightarrow B]$, $[B \rightarrow C]$, $[C \rightarrow A]$, $[C \rightarrow B]$, $[C \rightarrow C]$. Cada relação contém variáveis que descrevem um processo que gera material seguido por um outro que o filtra, assim definindo um conjunto de frequências possíveis para o movimento da voz à esquerda.

O material inicial do processo é construído ao reunir todas as frequências dentro de um alcance predeterminado (relativo a voz à esquerda) que estabelecem níveis de dissonância D com a voz à direita da relação, onde $D = [d_1, d_2, \dots]$ é uma lista que determina os níveis de dissonância a serem considerados. Por exemplo, imagine a relação $[A \rightarrow B]$, com os estados iniciais $A = 100\text{Hz}$ e $B = 200\text{Hz}$. Se $D = [2, 3]$, e o alcance de A for de $A/2$ a $3A$ (50 a 300), o conjunto resultante será a união entre todas as frequências que estabelecem níveis de dissonância 2 e 3 em relação a 200Hz, dentro do alcance:

$$(200 \cdot [2/7, 2/5, 2/3, 3/2]) \cup (200 \cdot [3/11, 3/10, 3/8, 3/7, 3/5, 3/4, 3/4]) \\ \approx [57.1, 80, 133.3, 300] \cup [54.5, 60, 75, 85, 120, 150] \text{ Hz}$$

Frequências disponíveis: $[54.5, 57.1, 60, 75, 80, 85.7, 120, 133.3, 150, 300] \text{ Hz}$

Em seguida, o conjunto é filtrado por um processo baseado na proximidade P , onde $P = [p_1, p_2, \dots]$ é uma lista de índices que indica as posições das frequências a serem selecionadas dentro de uma reorganização do conjunto baseada na ordem crescente de proximidade em relação a voz da esquerda na relação. Por exemplo, $P(1, 2, -2)$ indica a escolha da frequência mais próxima, da segunda mais próxima e da segunda mais distante em relação à voz da esquerda (100Hz). No entanto, é possível adotar distintos critérios para a determinação da proximidade entre as frequências, dependendo da abordagem. Atualmente o algoritmo possibilita a utilização de 3 métodos:

- Linear: Proximidade absoluta entre as frequências $\rightarrow P_{\text{lin}}(f_0, f) = |f_0 - f|$
- Logarítmica: Proximidade perceptiva entre as frequências $\rightarrow P_{\text{log}}(f_0, f) = \frac{\max(f_0, f)}{\min(f_0, f)}$
- Energética: ΔE entre sistemas harmônicos ideais com amplitude e massa constantes, mas frequências distintas $\rightarrow P_{\text{E}}(f_0, f) = |f_0^2 - f^2|$

Vale destacar a dedução deste terceiro método para não parecer arbitrário. Partimos da expressão clássica da energia de um oscilador harmônico ideal:

$$E = \frac{1}{2}m\omega^2 A^2$$

Substituindo a frequência angular ω por $2\pi f$, temos:

$$E = 2\pi^2 m A^2 f^2$$

Considerando amplitude e massa constantes, tem-se que a energia depende apenas do quadrado da frequência, ou seja, $E \propto f^2$. Portanto a diferença energética entre duas ondas com frequências distintas assume a forma:

$$\Delta E = m2\pi^2 A^2 |f_0^2 - f^2|$$

Por essa razão a proximidade energética é representada como:

$$P_{\text{E}}(f_0, f) = |f_0^2 - f^2|$$

A reorganização da lista que descreve as possíveis frequências em nosso exemplo fica da seguinte forma:

$$P_{\text{lin}} = [85.7, 120, 80, 75, 133.3, 60, 57.1, 54.5, 150, 300] \text{ Hz}$$

$$P_{\text{log}} = [85.7, 120, 80, 75, 133.3, 150, 60, 57.1, 54.5, 300] \text{ Hz}$$

$$P_{\text{E}} = [85.7, 80, 75, 120, 133.3, 150, 60, 57.1, 54.5, 300] \text{ Hz}$$

Vale ressaltar que é possível encontrar situações onde duas frequências possuem o mesmo valor de proximidade para cima ou para baixo, como 120 e 80 em relação a 100 sob o método linear. Neste caso, o algoritmo escolhe “arbitrariamente” se utiliza a opção acima ou abaixo na definição da ordem dessa lista. Por fim, o processo de filtragem descrito anteriormente nos gera os seguintes conjuntos:

$$P_{\text{lin}}(1, 2, -2) = [85.7, 120, 150] \text{ Hz}$$

$$P_{\text{log}}(1, 2, -2) = [85.7, 120, 54.5] \text{ Hz}$$

$$P_{\text{E}}(1, 2, -2) = [85.7, 80, 54.5] \text{ Hz}$$

3.1. Filtragem lógica

O método descrito acima explica o processo de geração das possíveis frequências para uma única relação de A ($[A \rightarrow B]$), no entanto sua frequência final é escolhida de forma a considerar todas as suas relações com as demais vozes. Isso é feito por meio de um procedimento lógico no qual as listas resultantes de cada relação interagem, comparando seus elementos, eliminando alguns e mantendo outros, de forma a gerar uma lista final que representa todas as possíveis frequências para A.

O circuito lógico utilizado é serial: a primeira lista é comparada à segunda, o conjunto resultante é então comparado à terceira, e assim sucessivamente. Isso configura um sistema onde a ordem das operações determina uma certa hierarquia entre as relações harmônica das vozes. A Figura 1 permite a visualização desse sistema:

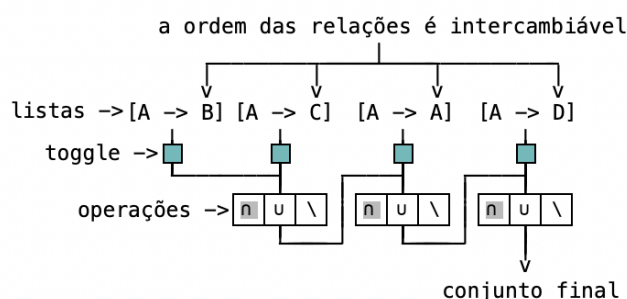


Figura 1. Trecho simplificado da interface gráfica do algoritmo construído em Pure Data. Nele observamos o caráter serial do cálculo, onde a ordem das listas resultantes de cada relação influencia diretamente no resultado do conjunto final.

Cada etapa do cálculo computa as listas utilizando um de 3 possíveis critérios, baseados na teoria dos conjuntos [Cantor 1895]: interseção ($\cap$), união ($\cup$), ou exclusão ($\setminus$). O "toggle" determina se a lista será ou não computada, ou seja, caso o caminho esteja travado (com x), a lista será transformada em um conjunto vazio $\emptyset$ antes de atravessar a função.

O caráter hierárquico do sistema depende da ordem disposta das relações e dos cálculos para cada estágio. Por exemplo, se todas as etapas do procedimento lógico estiverem selecionadas como interseção ($\cap$), a relação mais importante se torna a primeira, pois o conjunto final será vazio caso qualquer uma das outras listas não compartilhe ao menos um de seus elementos. O caráter dessa hierarquia assume outra forma quando todas as operações são de exclusão ($\setminus$), uma vez que o conjunto resultante contém apenas frequências da primeira relação, mas não necessita que as outras relações possuam seus elementos. A hierarquia deixa de existir caso todos os cálculos sejam de união ($\cup$), já que neste caso o conjunto final contém todos os elementos.

Por fim, o resultado dessas operações passa por mais um processo de filtragem idêntico ao descrito no capítulo anterior. No entanto, nessa fase é permitida apenas a entrada de um valor de proximidade, portanto ele define, dentre as opções do conjunto final, qual será a frequência da voz em questão.

3.2. Equilíbrio

Para mim, o processo de balancear a hierarquia entre as vozes com o conteúdo espectral de suas relações é onde está o principal potencial poético dentro deste algoritmo. Nem sempre é possível deslocar-se de forma a estar em consonância com os outros indivíduos. Construímos uma rede de hierarquias, e buscamos o menor movimento disponível para satisfazer o máximo de necessidades dentro de nosso alcance. No entanto, a sociedade muitas vezes nos exige uma movimentação que não está em consonância com nossa própria voz $[A \rightarrow A]$, e ficamos doentes. Ao mesmo tempo, se agirmos apenas de forma a satisfazer as nossas necessidades, condenamos a sociedade ao adoecimento e ao caos. O equilíbrio é necessário.

4. Frustrações

Apesar de já estar totalmente teorizado, os módulos do algoritmo responsáveis por temporalizar a mudança das frequências - o ritmo - terão que ser descritos com profundidade em um futuro trabalho, uma vez que seria impossível finalizar sua descrição com o devido rigor que necessita dentro do limite de páginas definido para a entrega deste relatório, sem sacrificar uma boa quantidade de material descrito anteriormente - alguém poderia argumentar que seria possível reduzir a revisão histórica, ou essa mesma divagação, mas eu replicaria que ela é essencial, uma vez que as ideias presentes no artigo parecem arbitrárias sem o contexto com o qual elas contrastam.

No entanto gostaria de destacar que os conceitos e técnicas descritos acima representam apenas a dimensão vertical do som dentro do algoritmo, mas o sistema permanece incompleto sem o tempo. O ritmo merece o mesmo tratamento baseado em consonância que a altura, já que a distinção entre ambos reside unicamente na forma como percebemos os fenômenos ondulatórios. Embora não possamos perceber frequências muito baixas como tons, elas inegavelmente o são, e por essa razão, está totalmente relacionada com os conceitos de consonância e lógica descritos acima.

Outro elemento do algoritmo que gostaria de ter incorporado no presente trabalho era a descrição do funcionamento da interface gráfica, e a influência de sua construção no desenvolvimento do rigor matemático e metodológico do sistema. No entanto deixo abaixo seu resultado para satisfazer a curiosidade do leitor interessado (Figura 2).

5. Conclusão

Este algoritmo, que está em fase inicial de implementação, parte de uma metodologia não-psicoacústica para gerar caminhos melódicos ancorados em um rigoroso processo matemático. Suas características permitem um processo composicional capaz de sonorizar relações de todo tipo sem se apoiar em nenhum processo estocástico, ou seja, não utiliza pseudo-aleatoriedades em nenhum estágio de seu processo generativo. Ele é uma proposta metodológica para a criação sonora que desconsidera a relação sensorial e cultural que os humanos tem com o fenômeno sonoro.

Apesar de não ter sido mencionado no corpo do artigo, a fim de não interromper o percurso retórico, é importante destacar o impulso inicial que deu origem a este trabalho: o projeto Gaia Senses [Mammoli et al. 2025]. Idealizado e coordenado por Artemis Moroni, o projeto tem como objetivo gerar composições audiovisuais generativas a partir

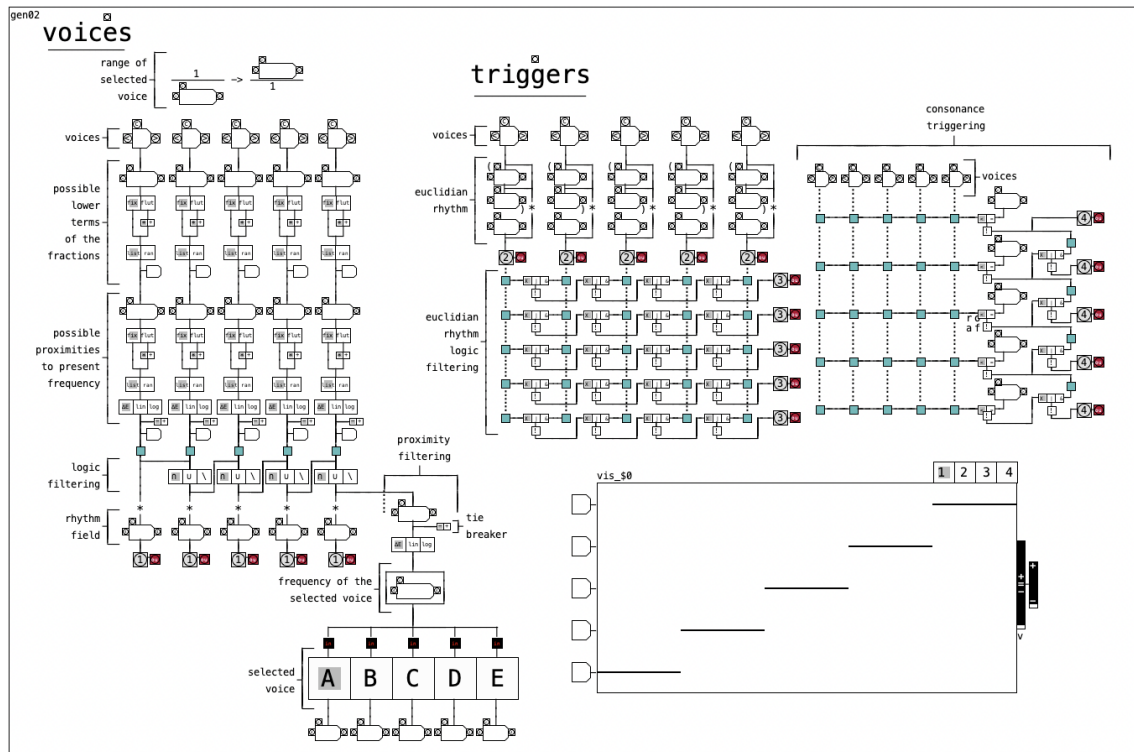


Figura 2. Atual estado da interface para o algoritmo generativo descrito no presente artigo

de dados de satélite, coletados em tempo real do usuário por meio de um site e de um aplicativo, disponível em:

<https://gaiasenses-web.vercel.app>.

O algoritmo será futuramente incorporado nas composições do projeto assim que finalizado, com o intuito de sonorizar as mudanças climáticas à partir da relação temporal entre dados de distintos períodos históricos.

Referências

- Cantor, G. (1895). Beiträge zur begründung der transfiniten mengenlehre. *Mathematische Annalen*, 46:481–512.
- Duffin, R. W. (2007). *How equal temperament ruined harmony (and why you should care)*. WW Norton & Company.
- Euclides, E., Bicudo, I., et al. (2009). *Os elementos*. Unesp.
- Helmholtz, H. L. (1863). *On the Sensations of Tone as a Physiological Basis for the Theory of Music*. Cambridge University Press.
- Mammoli, F., Zanchetta, G., Fontolan, L., and Moroni, A. (2025). Integrating computational creativity and climate data for environmental awareness: Design and analysis of an ongoing project.
- Plomp, R. and Levelt, W. J. M. (1965). Tonal consonance and critical bandwidth. *The journal of the Acoustical Society of America*, 38(4):548–560.

Sethares, W. A. (2005). *Tuning, timbre, spectrum, scale*. Springer.

Tsuji, K. and Müller, S. C. (2021). Tunings—from pythagoras to equal temperament. In *Physics and Music: Essential Connections and Illuminating Excursions*, pages 71–90. Springer.