

sistema de multimídia para avaliação do desenvolvimento de rotas de leitura em escolares e de seu comprometimento em disléxicos

*elizeu coutinho de macedo¹; fernando César capovilla²;
marcelo duduchi³; roberto amilton bernardes sória⁴*

(1) Doutorando em Psicologia Experimental do Instituto de Psicologia da Universidade de São Paulo; (2) Ph.D., Professor do Instituto de Psicologia da Universidade de São Paulo; (3) Mestrando em Neurociências e Comportamento do Instituto de Psicologia da Universidade de São Paulo e Professor do Departamento de Processamento de Dados da FATEC-SP; (4) Mestrando da Escola Politécnica da Universidade de São Paulo

O sistema de multimídia *CronoFonos 2.0* permite caracterizar o grau de desenvolvimento das rotas de leitura fonológica e lexical em escolares de pré-escola a quarta série do primeiro grau, bem como de seu comprometimento em portadores de dislexias adquiridas ou do desenvolvimento. Numa tarefa de leitura em voz alta de itens isolados, apresenta itens psicolinguísticos e analisa os tipos de erros cometidos em sua pronúncia bem como os parâmetros temporais da locução. Tais parâmetros incluem o tempo de reação locucional, a duração locucional, a frequência de segmentos locucionais (formantes) juntamente com as respectivas durações. Ao comparar o número de segmentos locucionais (picos de energia) presentes na pronúncia com o número de sílabas presentes na ortografia, obtém os padrões de segmentação (PS) da leitura. Quanto maior a segmentação da leitura, tanto maior é a dependência da criança no uso da rota fonológica para proceder à leitura. As distribuições relativas dos tipos de erros e os parâmetros temporais são analisados como função de variáveis como a lexicalidade (palavra *versus* quase-palavra *versus* pseudopalavra), frequência de ocorrência na língua (alta *versus* baixa), regularidade das relações grafo-fonêmicas

envolvidas (alta nos itens de pronúncia regular *versus* média naqueles cuja pronúncia envolve regras de posição *versus* baixa naqueles cuja pronúncia é irregular ou excepcional), comprimento (número de letras ou de sílabas), tipo de letra (de fôrma *versus* cursiva), além de variáveis de natureza articulatória.

Pode ser usado psicometricamente como teste de habilidade e desempenho de leitura em voz alta de itens isolados, ou como teste de exame neuropsicolinguístico cognitivo, uma vez que, ao avaliar a distribuição de tipos de erro e os padrões temporais da locução, permite avaliar a integridade das rotas de leitura bem como das unidades de processamento que as compõem, tais como o sistema de reconhecimento visual de palavras, o sistema de produção fonológica das palavras etc. De qualquer modo, trata-se de um teste construído por referência a um construto teórico. Os traços latentes avaliados pelo teste consistem nos construtos rotas de leitura. A construção do teste foi feita a partir da abordagem teórica de *processamento de informação* em psicologia cognitiva (Eysenck e Keane, 1990), mais especificamente da *teoria de duplo processo* na

leitura (Ellis e Young, 1988; Morton, 1969, 1979, 1989). Tal teoria propõe a existência de duas rotas principais para a leitura: a fonológica e a lexical. Na fonológica a leitura ocorre por decodificação de acordo com regras de correspondência grafema-fonema, e a pronúncia é construída segmento a segmento. Na lexical a leitura ocorre por acesso direto ao léxico, sem a intermediação de decodificação fonológica, e a pronúncia é emitida como um todo após sua recuperação por localização ou endereçamento lexical.

Como a pronúncia pela rota lexical tende a ser emitida após recuperação lexical como um todo não-segmentado, enquanto que a pronúncia pela rota fonológica tende a ser construída segmento a segmento por decodificação fonológica, considerando-se um dado comprimento (medido em termos de número de grafemas ou de sílabas componentes), a pronúncia de palavras lidas fonologicamente tem usualmente maior duração locucional do que a de palavras lidas lexicalmente. Além disso, as locuções de palavras lidas fonologicamente também apresentam usualmente um maior número de segmentos locucionais discerníveis do que as de palavras lidas lexicalmente. Portanto, normalmente a correspondência entre o número de segmentos locucionais discerníveis no sinal vocálico (isto é, registro da locução) e o número de segmentos silábicos presentes na ortografia é maior para itens lidos fonologicamente do que para aqueles lidos lexicalmente.

O padrão de segmentação (PS), que é uma medida da correspondência entre a ortografia e o sinal vocálico, é obtido subtraindo-se o número de segmentos silábicos presentes na ortografia do número de segmentos locucionais discerníveis no sinal vocálico. Portanto, PS pode ser classificado como *igualação* (PS=), *subsegmentação* (PS-), ou *supersegmentação* (PS+). A classificação de um PS como de *igualação* (PS=) indica que o número de segmentos locucionais discerníveis no sinal vocálico corresponde àquele da ortografia, sugerindo que a leitura é fonológica. A classificação de um PS como de *subsegmentação* (PS-) indica que o número de segmentos locucionais discerníveis no

sinal vocálico é inferior àquele da ortografia, sugerindo que a leitura é lexical. Já a classificação de um PS como de *supersegmentação* (PS+) indica que o número de segmentos locucionais discerníveis no sinal vocálico é superior àquele da ortografia, o que sugere gagueira.

Um dos estudos que emprega *software* específico para análise da dinâmica temporal de repetições em crianças pré-escolares com gagueira é o de Throneburg e Yairi (1994). Este estudo empregava o *software Cspeech* (Milenkovic, 1987) para medir, nos *displays* visuais dos espectrogramas, as durações das unidades de repetições faladas, os intervalos de silêncio entre as unidades e a disfluência total. Aqueles autores descobriram que a duração total das disfluências das crianças que gaguejavam era significativamente mais curta devido a os seus intervalos silenciosos serem mais curtos quando comparados às disfluências de unidades de repetição iguais produzidas pelos sujeitos do grupo controle. Conforme Yairi e Lewis (1984), uma unidade de repetição é definida como a produção de um segmento extra (como, por exemplo, *b-but* e *and-and*) e uma unidade dupla de repetição é definida como duas produções extras de segmento (como *b-b-but*, e *and-and-and*). Para a operação desse *software*, os autores precisavam selecionar as unidades de repetição marcando-as, eles próprios, por meio de colocação de cursores no início e final da curva de energia espectral. *CronoFonos 2.0* poderia ser empregado para automatizar por completo tal processo, pois detecta, em tempo real, o início e o término não apenas de locuções inteiras como também de seus segmentos.

CronoFonos 2.0 pode ser customizado para diferentes *corpora*. Isto é feito por meio da manipulação de até cinco parâmetros independentes que, uma vez estabelecidos, operam em todo o *corpus* fornecendo uma medida documentada, precisa e objetiva do efeito de cada uma das categorias psicolinguísticas de interesse do examinador para aquele *corpus* específico.

A seleção de itens para a composição do teste deve ser feita de modo a cobrir as principais variáveis

psicolinguísticas (lexicalidade, regularidade, frequência de ocorrência, comprimento, tipo de letra) em seus níveis representativos (palavra *versus* pseudopalavra, regular *versus* regra *versus* irregular, alta *versus* baixa, bi *versus* trissilábicas, de fôrma *versus* cursiva) que permitem distinguir as duas rotas. Para tanto, podem ser empregadas formas paralelas de listas já publicadas de itens psicolinguísticos feitas para crianças (Pinheiro, 1994, 1996). Listas de frequência de palavras para adultos também encontram-se disponíveis (Françozo, não publicado). *CronoFonos* pode ser executado em *desktops* ou *notebooks Pentium 100 MHz* com kit multimídia 10x. Grava os arquivos de som no formato WAV, com 8 bits, mono, e 11 Hz.

CRONOFONOS 2.0 apresenta ao sujeito uma lista de itens psicolinguísticos, um a um, e registra a leitura em voz alta que o sujeito faz de cada item. A Figura 1 contém o *layout* de uma das centenas de telas do *software*. Nesta tela em particular encontra-se apresentada a pseudopalavra “calafra”. Abaixo do item, na parte inferior da tela, encontram-se os botões *Pausa* e *Seguinte*. O botão *Seguinte* deve ser pressionado pelo examinador a cada término de locução pelo sujeito. Isto fecha o arquivo de som, mas não afeta quaisquer parâmetros temporais da locução. O botão *Pausa* permite a interrupção momentânea da apresentação sequencial dos estímulos escritos.

calafra



FIGURA 1 - LAYOUT DE TELA DE APRESENTAÇÃO DE ITEM A SER LIDO EM CRONOFONOS 2.0

A Figura 2 contém o *layout* da tela de programação, que contém o item “palavra” na janela de leitura e os botões de opção *Adicionar*, *Inserir*, *Excluir*, *Arquivo*, *Copiar como*, *Fonte* e *Voltar*. A janela de leitura contém a lista de itens a serem apresentados ao sujeito para leitura em voz alta numa dada sessão. O botão *Adicionar* permite acrescentar itens à lista. Os botões *Inserir* e *Excluir* possibilitam incluir um novo item em qualquer posição da lista ou excluí-lo, respectivamente. Os botões *Arquivo* e *Copiar como* permitem criar diferentes tipos de listas e salvá-las com diferentes nomes para acesso ulterior. O botão *Fonte* permite determinar o tipo, tamanho e cor do caractere (letra) dos itens escritos a serem apresentados para leitura. O botão *Voltar* permite sair do modo de configuração da sessão para o modo de aplicação propriamente dito.

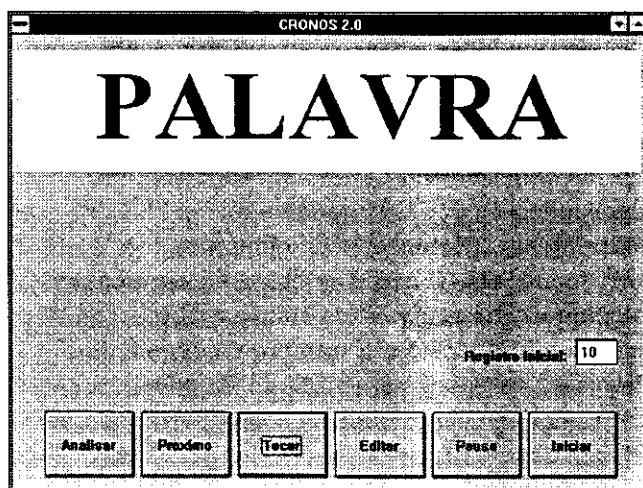


FIGURA 2 - LAYOUT DA TELA PRINCIPAL DE CRONOFONOS 2.0 COM OS BOTÕES ANALISAR, PRÓXIMO, TOCAR, EDITAR E INICIAR

As Figuras 3, 4 e 5 contêm os registros da forma de onda (*waveforms*) dos itens “palavra”, “calafra” e “táxi”, respectivamente, lidos por um adulto. Cada *waveform* consiste num registro de amplitude de sinal (ordenada) em função de duração da locução numa escala de ms. (abscissa). No *waveform* o ponto N0 indica o momento de apresentação do estímulo escrito, N1 indica o momento do início da locução (ou primeiro segmento locucional), e N2, o de seu término. O *tempo de*

reação locucional (TRL) é calculado subtraindo N0 de N1. A duração locucional total (DLT) é calculada subtraindo N1 do último N registrado. Por exemplo, nos waveforms abaixo, a duração locucional total do item *palavra* é N2-N1 e do item *calaфра* é N6-N1. O número de segmentos locucionais (NSL) corresponde ao índice do último N dividido por 2. Assim, o NSL de *palavra* é 1 (isto é, 2/2), e o de *calaфра* é 3 (isto é, 6/2). A duração do segmento locucional (DSL) (isto é, do primeiro segmento locucional) é calculada subtraindo N1 de N2. Do mesmo modo, N3 indica o momento do início do segundo segmento locucional e N4 de seu fim, e a duração deste segundo segmento é calculada subtraindo N3 de N4. E assim por diante.

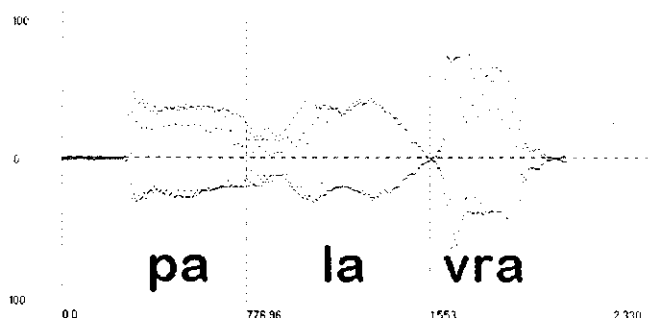


FIGURA 3 - REPRESENTAÇÃO GRÁFICA DO SINAL VOCÁLICO DA PALAVRA "PALAVRA" LIDA POR UM ADULTO

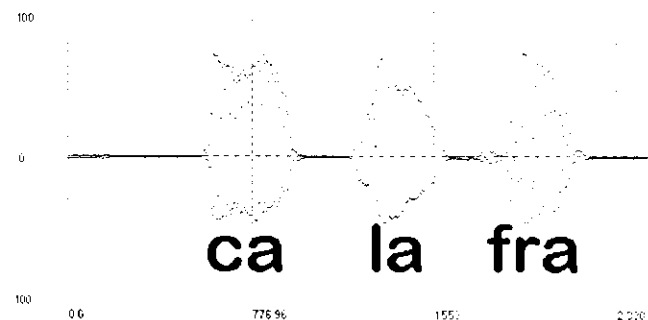


FIGURA 4 - REPRESENTAÇÃO GRÁFICA DO SINAL VOCÁLICO DA PSEUDOPALAVRA "CALAFRA" LIDA POR UM ADULTO

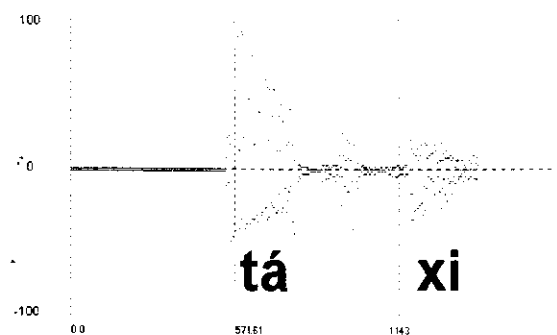


FIGURA 5 - REPRESENTAÇÃO GRÁFICA DOS SINAIS VOCÁLICOS DA PALAVRA "TÁXI" LIDA POR UM ADULTO

Como explicado anteriormente, os itens escritos apresentados para leitura em voz alta podem ser palavras ou pseudopalavras, e podem ter diferentes números de sílabas escritas ou NSO (número de segmentos ortográficos). Como pseudopalavras são lidas fonologicamente, é esperado um NSL mais próximo ao número de segmentos ortográficos (NSO). Assim, é esperado um *matching* ($PS=$), ou seja, uma *igualação* ou concordância entre NSL e NSO para pseudopalavras. Logo, para pseudopalavras, $NSO - NSL \approx 0$. Por outro lado, como palavras de alta frequência são lidas lexicalmente, é esperado um NSL mais baixo do que o NSO. Logo, para palavras, $NSO > NSL$. A disparidade do NSL registrado em relação ao NSO esperado pode ser de dois tipos: *subsegmentação* ($PS-$) quando $NSL < NSO$; e *supersegmentação* ($PS+$) quando $NSL > NSO$. Os termos *subsegmentação* e *supersegmentação* resultam de uma expectativa teórica de que num determinado ponto do desenvolvimento da leitura o número de segmentos locucionais emitidos durante a leitura em voz alta tende a coincidir (isto é, tende a haver *matching* ou *igualação*) com o número de segmentos ortográficos (isto é, $NSL = NSO$).

Conforme o arrazoado acima, as disparidades entre NSL e NSO resultam em *supersegmentação*, ou seja, um excesso de segmentos locucionais em relação aos ortográficos; ou *subsegmentação*, ou seja, um *déficit* de segmentos locucionais em relação aos ortográficos. A frequência de segmentos em excesso (*super*) ou *déficit* (*sub*) é obtida

subtraindo *NSO* de *NSL*. Quando o resto de *NSL - NSO* for negativo, ele é chamado *PS-n*; quando for positivo, é chamado de *PS+n*, sendo que *n* designa o número de segmentos a menos ou a mais. A presença de *PS-* (isto é, *NSL < NSO*) pode indicar leitura lexical fluente ou mesmo omissão de segmentos na leitura; já a presença de *PS+* (isto é, *NSL > NSO*) pode indicar gagueira ou leitura escandida de determinados caracteres. Na ortografia, o caractere *x* é um dos caracteres que tende a evocar tal leitura escandida.

Um exemplo de *igualação* (*PS=*) é apresentado no *waveform* da Figura 4 referente à pseudopalavra “calafra”. Como pode ser observado, como se trata de uma pseudopalavra, a leitura foi feita por decodificação fonológica dos segmentos correspondentes às sílabas. Assim, o número de segmentos locucionais coincidiu com o número de segmentos ortográficos (*NSL = NSO = 3*). Um exemplo de *subsegmentação* (*PS-*) é apresentado no *waveform* da Figura 3 referente à palavra de alta frequência “palavra”. Como pode ser observado, como se trata de uma palavra real de alta frequência, a leitura em voz alta foi fluente e não-interrompida por pausas ou intervalos locucionais. Conseqüente-mente, o número de segmentos locucionais (*NSL*) resultante foi 1, embora o número de segmentos ortográficos (*NSO*) fosse 3. Isto resultou num padrão de *subsegmentação* de 2, ou seja, em dois segmentos locucionais a menos que ortográficos (isto é, *PS-2*) (ou seja, *PS = NSL - NSO = 1 - 3 = -2*). Um exemplo de *supersegmentação* (*PS+*) é apresentado no *waveform* da Figura 5 referente à palavra “táxi”. Como pode ser observado, como o item contém o grafema “x” que deve ser lido como /cs/, a leitura em voz alta tendeu a ser escandida a tal ponto que intervalos locucionais “artificiais” foram introduzidos no registro. Conseqüentemente, o número de segmentos locucionais (*NSL*) resultante foi 3, embora o número de segmentos ortográficos (*NSO*) fosse 2. Isto resultou num padrão de *supersegmentação* igual a 1, ou seja, de um segmento locucional a mais que ortográfico (isto é, *PS = NSL - NSO = 3 - 2 = +1*).

A leitura fonológica é obtida pela decodificação serial dos segmentos ortográficos (por exemplo, sílabas) e tende a ser escandida e, portanto, mais claramente

segmentada do que a lexical. Assim, para um mesmo número de sílabas escritas, como a leitura fonológica é usualmente feita quando os itens a serem lidos são pseudopalavras e a lexical, palavras de alta frequência, tenderá a haver um maior número de segmentos locucionais na leitura (fonológica) de pseudopalavras do que na leitura (lexical) de palavras de alta frequência.

Num estudo preliminar de validação (Capovilla e colaboradores, no prelo) do *software* como instrumento para avaliação das estratégias de leitura fonológica e lexical, 35 universitários foram expostos a *CronoFonos 2.0* programado com a lista de Pinheiro (1994). Nesse estudo havia quatro variáveis independentes: lexicalidade (pseudopalavra, palavra real), frequência (baixa, alta), regularidade (regular, regra, irregular), e comprimento (4 a 7 letras), e três variáveis dependentes: padrão de *subsegmentação PS-* (isto é, frequência de segmentos locucionais identificados a menos), tempo de reação e duração locucional. De modo geral, os dados adequaram-se bastante bem às expectativas baseadas na bibliografia. Os efeitos estatisticamente significantes encontram-se sumariados a seguir.

Em termos de lexicalidade, a segmentação média foi maior nas pseudopalavras do que nas palavras. O tempo de reação das pseudopalavras foi também maior que o das palavras. Finalmente, a duração locucional de palavras de alta frequência foi menor do que a de palavras de baixa frequência e de pseudopalavras. Assim, pseudopalavras produziram maior tempo de reação locucional, maior duração locucional e maior frequência de segmentação do que palavras reais, o que é plenamente compatível com a interpretação de que tais pseudopalavras são lidas fonologicamente. Também, a duração locucional de palavras de alta frequência foi menor do que a das de baixa e do que a das de pseudopalavras, o que é compatível com a interpretação de que elas são lidas lexicalmente. Em termos de regularidade, a segmentação média foi maior nos itens regulares do que nos regra, e nesses do que nos irregulares. Além disso, os itens regulares (em que não havia discrepância entre as pronúncias obtidas por decodificação e por acesso lexical) produziram tempo de reação menor do que

os itens regra e irregulares. Em termos de comprimento, a segmentação nos itens de 7 letras foi maior do que nos demais, e nos de 6, maior do que nos de 5. Consistente com isto, o tempo de reação nos itens com 7 letras foi maior do que nos demais. Finalmente, a duração da locução nos itens de 7 letras foi maior do que nos de 6, nos de 6, do

que nos de 5, e nos de 5, do que nos de 4.

O sistema foi recentemente aplicado a uma amostra de cerca de 140 escolares de primeira à quarta série do primeiro grau. Os dados de tal estudo serão relatados em artigo ulterior.

referências bibliográficas

1. Capovilla, F.C.; Macedo, E.C.; Duduchi, M.; Sória, R.A.B. - Avaliação das estratégias de leitura fonológica e lexical II: Aplicação do software *CronoFonos 2.0* a universitários. *Reflexões e Pesquisa em Psicologia*, no prelo.
2. Ellis, A.; Young, A.W. - *Human cognitive neuropsychology*. London, Lawrence Erlbaum, 1988.
3. Eysenck, M.W.; Keane, M.T. - *Psicologia cognitiva: Um manual introdutório*. Porto Alegre, Editoria Artes Médicas Sul, 1990.
4. Françoze, E. - *Lista de frequência de palavras*. Campinas, Universidade Estadual de Campinas, não publicado.
5. Millenkovic, P. - Least mean square measures of voice perturbation. *J. Speech Hear. Res.*, v.30, p.529-38, 1987.
6. Morton, J. - The interaction of information in word recognition. *Psychol. Rev.*, v.76, p.165-78, 1969.
7. Morton, J. - Facilitation in word recognition: Experiments causing change in the Logogen Model. In: Kohlers, P.A.; Wrolstand, M.E. e Bouma, H. (eds.) - *Processing of visible language*. V.1. New York, MIT Press, 1979. p.259-68.
8. Morton, J. - An information-processing account of reading acquisition. In: Galaburda, A.M. (org.) - *From Reading to Neurons*. Cambridge, MIT Press, 1989.
9. Pinheiro, A.M.V. - *Leitura e escrita: Uma abordagem cognitiva*. Campinas, Editorial Psy II, 1994.
10. Pinheiro, A.M.V. - *Contagem de frequência de ocorrência de palavras expostas a crianças na faixa pré-escolar e séries iniciais do primeiro grau*. São Paulo, Associação Brasileira de Dislexia, 1996.
11. Throneburg, R.N.; Yairi, E. - Temporal dynamics of repetitions during the early stage of childhood stuttering: An acoustic study. *J. Speech Hear. Res.*, v.37, p.1067-75, 1994.
12. Yairi, E.; Lewis, B. - Disfluencies at the onset of stuttering. *J. Speech Hear. Res.*, v.27, p.154-9, 1984.

errata

Desculpamo-nos por alguns erros de publicação ocorridos em **Temas sobre Desenvolvimento, v.6, n.34, p.14-9, 1997**, em artigo intitulado "*Estudo de follow-up em crianças autistas*", de autoria da Psicóloga Celiane Ferreira Secunho, cujas correções fazemos agora:

- página 14, primeiro parágrafo, terceira linha: texto correto: "... colaboração com Lean Eisenberg, em 1943. Foram ..."
- página 17, tabela 1: tabela correta

TABELA 1 - FOLLOW-UP DE 12 AUTISTAS
COM RELAÇÃO À ADAPTAÇÃO SOCIAL

CARACTERÍSTICA	NÚMERO DE CLIENTES	PORCENTAGEM DE GRUPO
Desadaptados/alheios	3	27.3%
Participam de atividades em casa/escola	8	72.7%
Dependentes da família	11	100.0%

- página 17, tabela 3: tabela correta:

TABELA 3 - FOLLOW-UP DE 12 AUTISTAS
COM RELAÇÃO À ALIMENTAÇÃO

CARACTERÍSTICA	NÚMERO DE CLIENTES	PORCENTAGEM DE GRUPO
Alimenta normalmente	4	36.4%
Depende para comer	3	27.3%
Precisa de alguma ajuda para comer	4	36.4%

- página 18, tabela 7: tabela correta:

TABELA 7 - FOLLOW-UP DE 12 AUTISTAS COM RELAÇÃO
AO NÍVEL DE INTELIGÊNCIA

CARACTERÍSTICA	NÚMERO DE CLIENTES	PORCENTAGEM DE GRUPO
Média	3	25.0%
Inferior	4	33.4%
Deficiente	5	45.6%